

Indian Institute of Information Technology, Allahabad



Image Categorization

By

Dr. Shiv Ram Dubey

Computer Vision and Biometrics Lab

Department of Information Technology

Indian Institute of Information Technology, Allahabad



DISCLAIMER

The content (text, image, and graphics) used in this slide are adopted from many sources for Academic purposes. Broadly, the sources have been given due credit appropriately. However, there is a chance of missing out some original primary sources. The authors of this material do not claim any copyright of such material.



Visual Recognition and Learning

- Image Categorization
- Image Features
- Classifiers
- Neural Networks
- Convolutional Neural Networks
- Object Detection
- Segmentation
- Image Generation
- Etc.



WHAT DO YOU SEE IN THIS IMAGE?



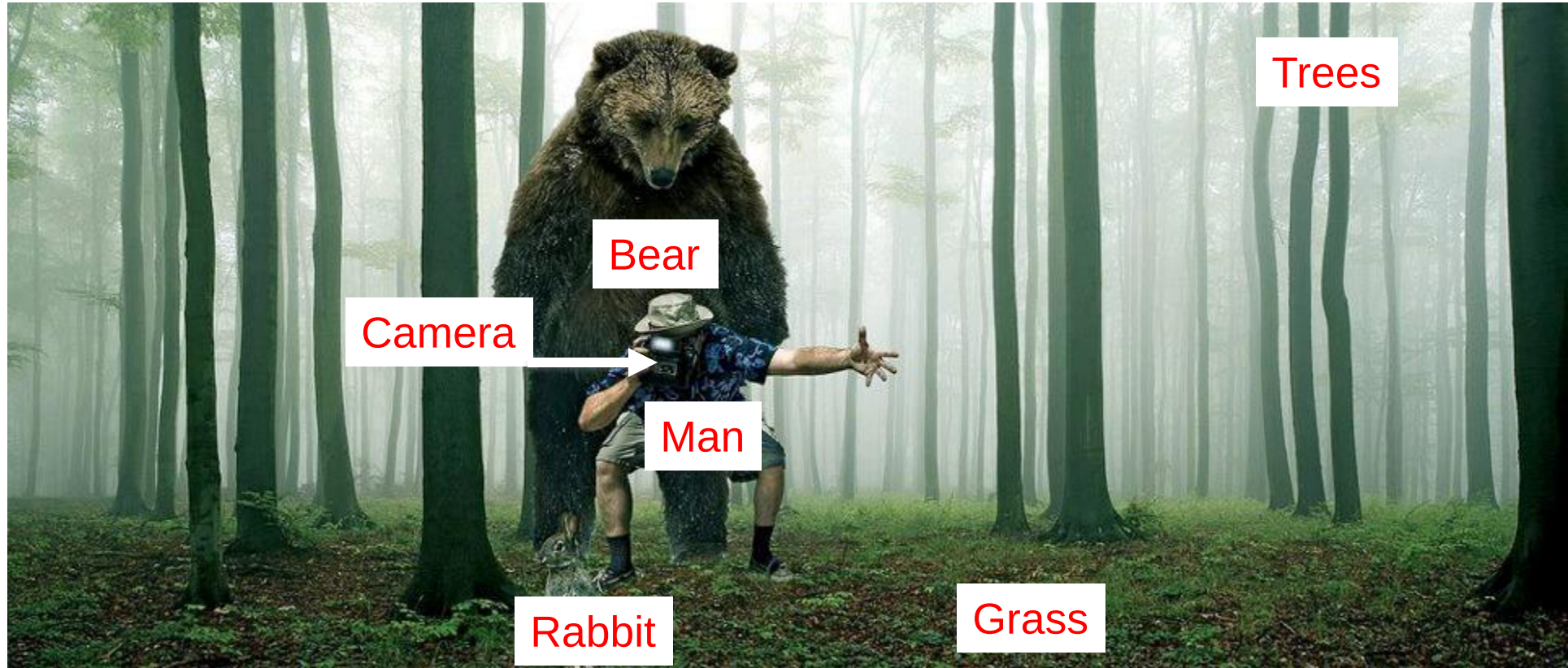
WHAT DO YOU SEE IN THIS IMAGE?



Forest



WHAT DO YOU SEE IN THIS IMAGE?



Forest



DESCRIBE, PREDICT, OR INTERACT WITH THE OBJECT BASED ON VISUAL CUES



Is it **dangerous**?

Is it **alive**?

Is it **soft**?

How **fast** does it run?

Does it have a **tail**?

Can I **poke** with it?



WHY DO WE CARE ABOUT CATEGORIES?

- From an object's category, we can make predictions about its behavior in the future, beyond of what is immediately perceived.
- Pointers to knowledge
 - Help to understand individual cases not previously encountered

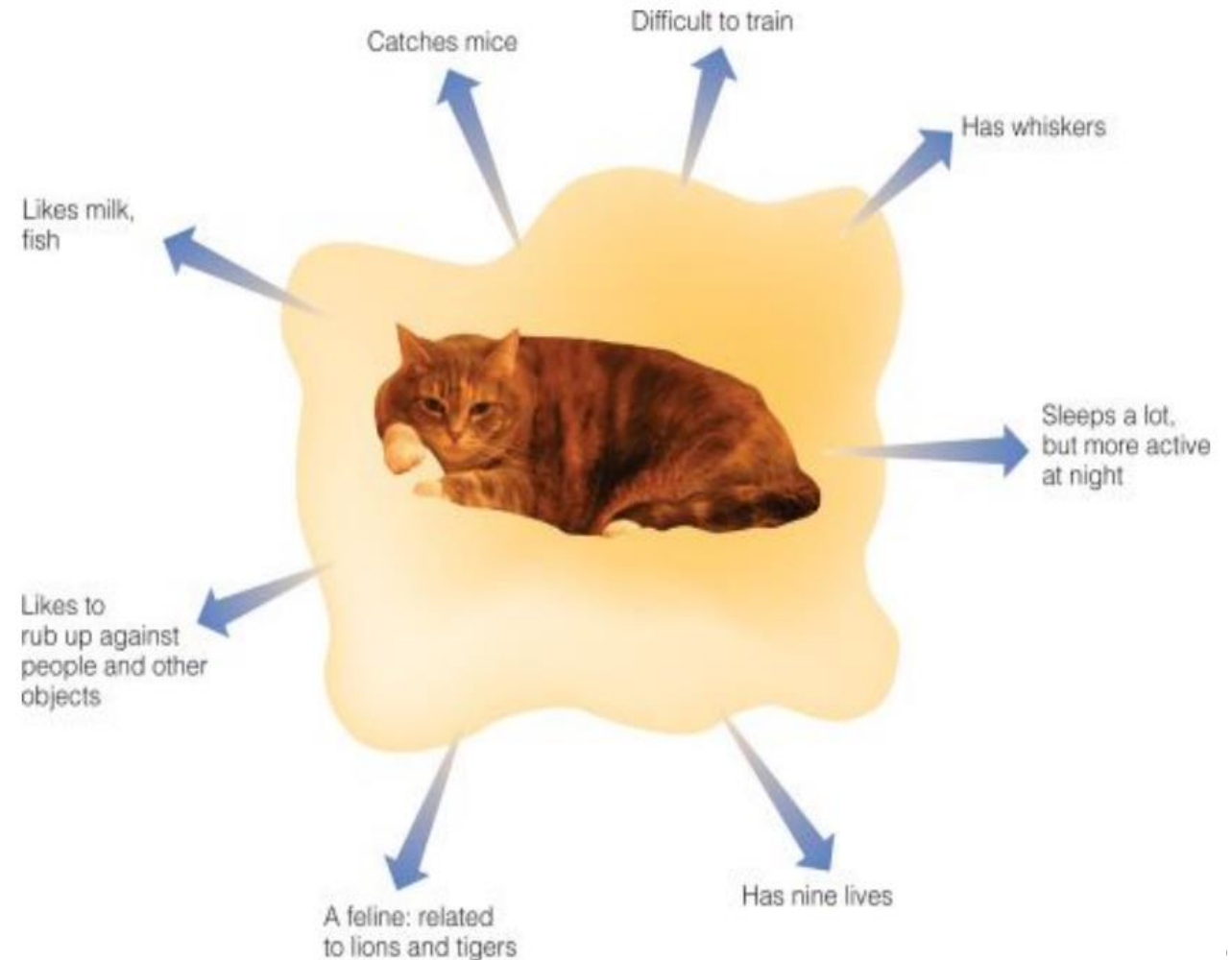


IMAGE CATEGORIZATION

- Cat vs Dog



IMAGE CATEGORIZATION

- Object recognition



Caltech 101 Average Object Images



IMAGE CATEGORIZATION

- Place recognition



Places Database [[Zhou et al. NIPS 2014](#)]



IMAGE CATEGORIZATION

- Image style recognition



HDR



Macro



Baroque



Rococo



Vintage



Noir



Northern Renaissance



Cubism



Minimal



Hazy



Impressionism



Post-Impressionism



Long Exposure



Romantic



Abs. Expressionism



Color Field Painting

Flickr Style: 80K images covering 20 styles.

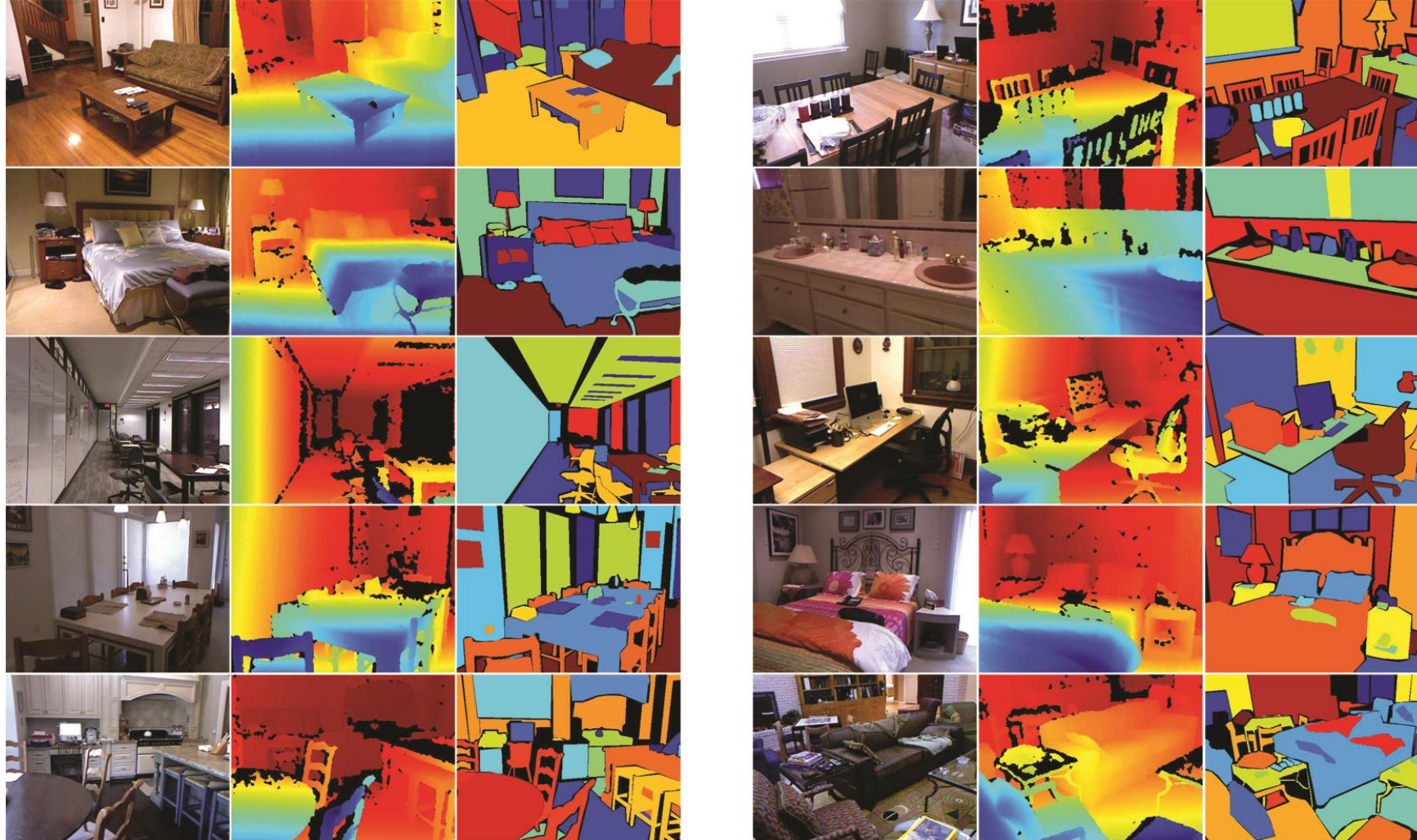
Wikipaintings: 85K images for 25 art genres.

[[Karayev et al. BMVC 2014](#)]



REGION CATEGORIZATION

- Semantic segmentation from RGBD images



[Silberman et al.,
ECCV 2012]



REGION CATEGORIZATION

- Material recognition



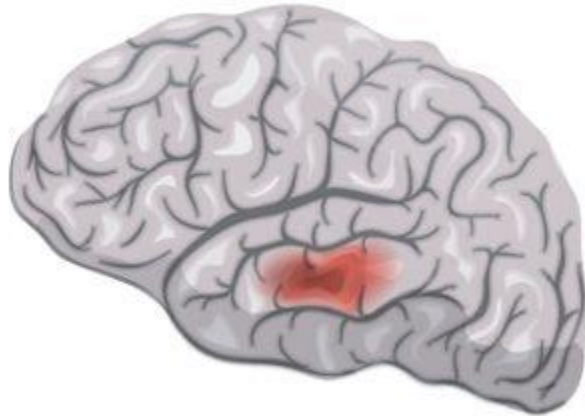
[Bell et al.
CVPR 2015]



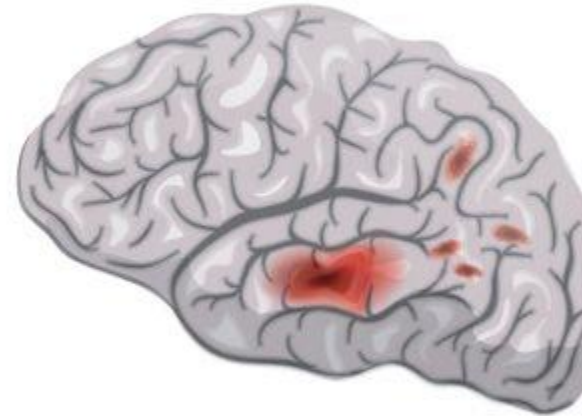
REGION CATEGORIZATION

- Primary Tumor vs Metastasis

Brain Cancer



the primary tumor

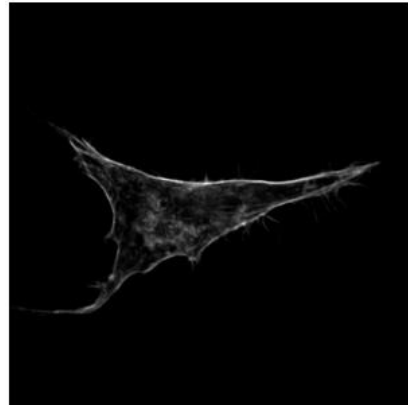


metastasis

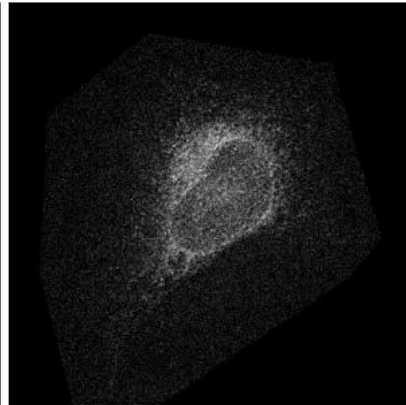


REGION CATEGORIZATION

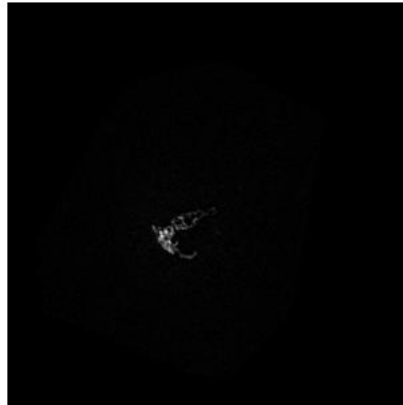
- Identification of sub-cellular organelles



DNA (Nuclei)



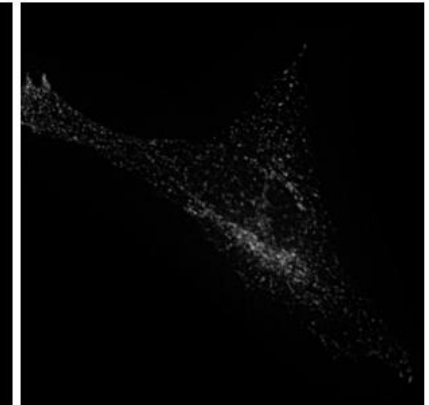
ER (Endoplasmic reticulum)



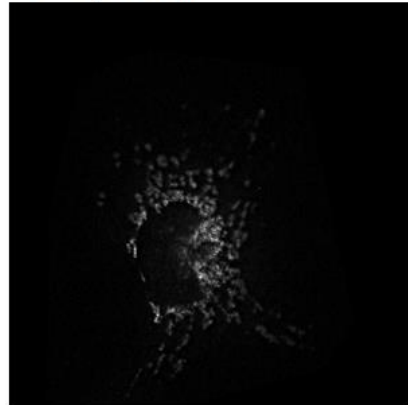
Giantin (cis/medial Golgi)



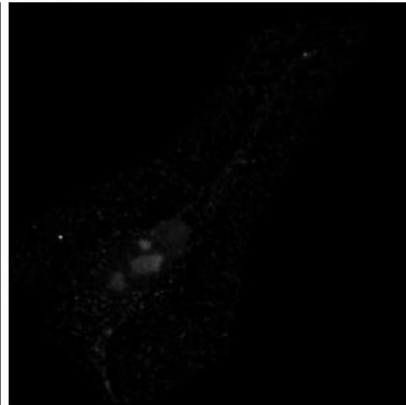
GPP130 (cis Golgi)



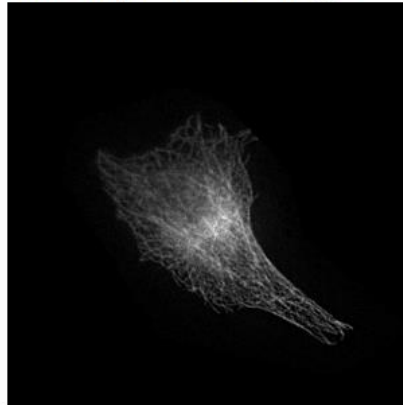
Lamp2 (Lysosomes)



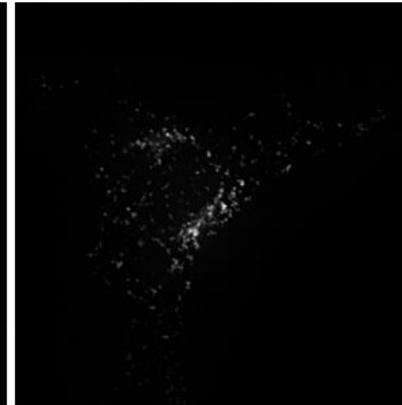
Mitochondria



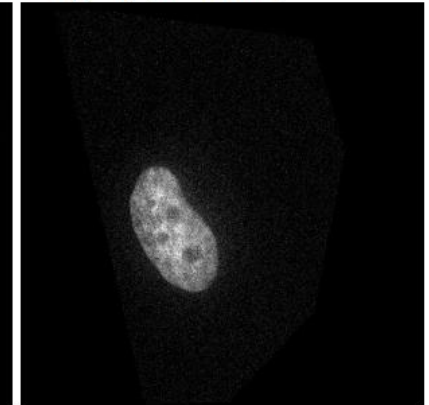
Nucleolin (Nucleoli)



Actin



TfR (Endosomes)



Tubulin

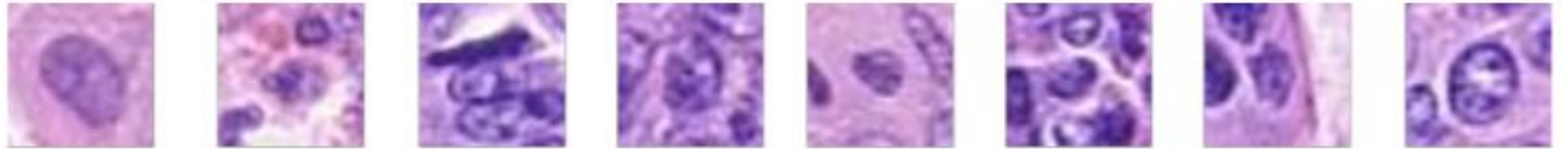
Fluorescence microscopy images of HeLa cells



REGION CATEGORIZATION

- Colon Cancer Nuclei Classification

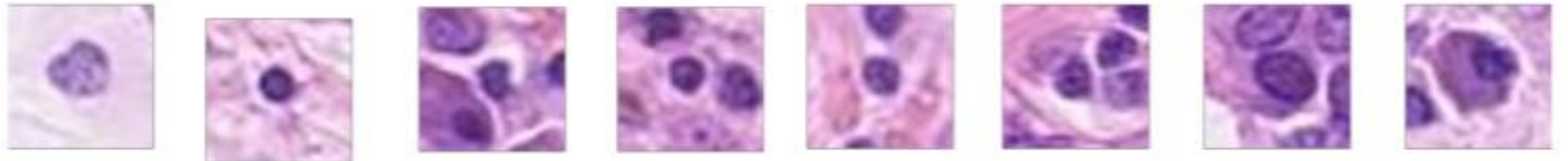
‘Epithelial’



‘Fibroblast’



‘Inflammatory’



‘Miscellaneous’



“CRCHistoPhenotypes” dataset images



IMAGE CATEGORIZATION / CLASSIFICATION



THE STATISTICAL LEARNING FRAMEWORK

- Apply a prediction function to a feature representation of the image to get the desired output:

$f(\text{apple image}) = \text{"apple"}$

$f(\text{tomato image}) = \text{"tomato"}$

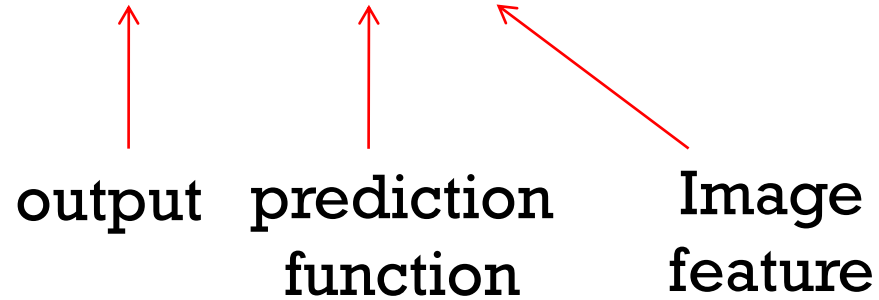
$f(\text{cow image}) = \text{"cow"}$



THE STATISTICAL LEARNING FRAMEWORK

$$\mathbf{y} = f(\mathbf{x})$$

output prediction
 function Image
 feature



The diagram illustrates the statistical learning framework equation $\mathbf{y} = f(\mathbf{x})$. Three red arrows point from the labels below to the components of the equation: one from 'output' to $\mathbf{y}$, one from 'prediction function' to f , and one from 'Image feature' to $\mathbf{x}$.



THE STATISTICAL LEARNING FRAMEWORK

$$\mathbf{y} = f(\mathbf{x})$$

output prediction Image
 function feature

The diagram illustrates the statistical learning framework equation $\mathbf{y} = f(\mathbf{x})$. Three red arrows point from the labels below to the corresponding parts of the equation: one from 'output' to $\mathbf{y}$, one from 'prediction function' to f , and one from 'Image feature' to $\mathbf{x}$.

- **Training:** given a *training set* of labeled examples $\{(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_N, \mathbf{y}_N)\}$, estimate the prediction function f by minimizing the prediction error on the training set



THE STATISTICAL LEARNING FRAMEWORK

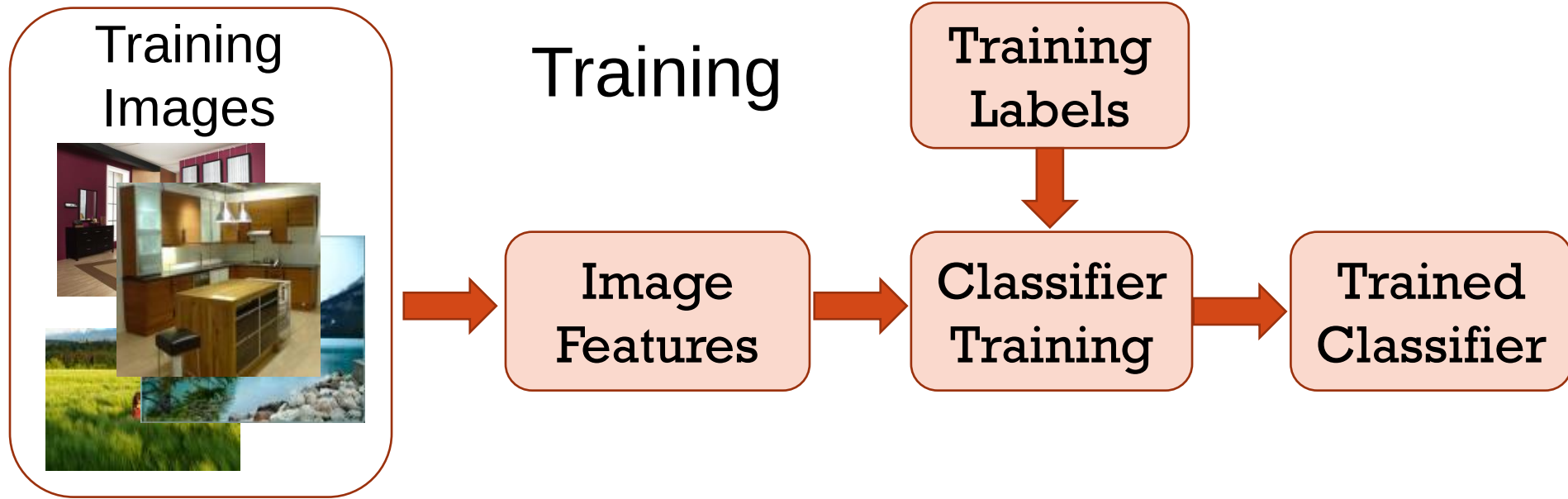
$$\mathbf{y} = \mathbf{f}(\mathbf{x})$$

The diagram illustrates the equation $\mathbf{y} = \mathbf{f}(\mathbf{x})$. Below the equation, three labels are positioned: 'output' under $\mathbf{y}$, 'prediction function' under $\mathbf{f}$, and 'Image feature' under $\mathbf{x}$. Red arrows point from each label to its corresponding symbol in the equation: an arrow from 'output' to $\mathbf{y}$, an arrow from 'prediction function' to $\mathbf{f}$, and an arrow from 'Image feature' to $\mathbf{x}$.

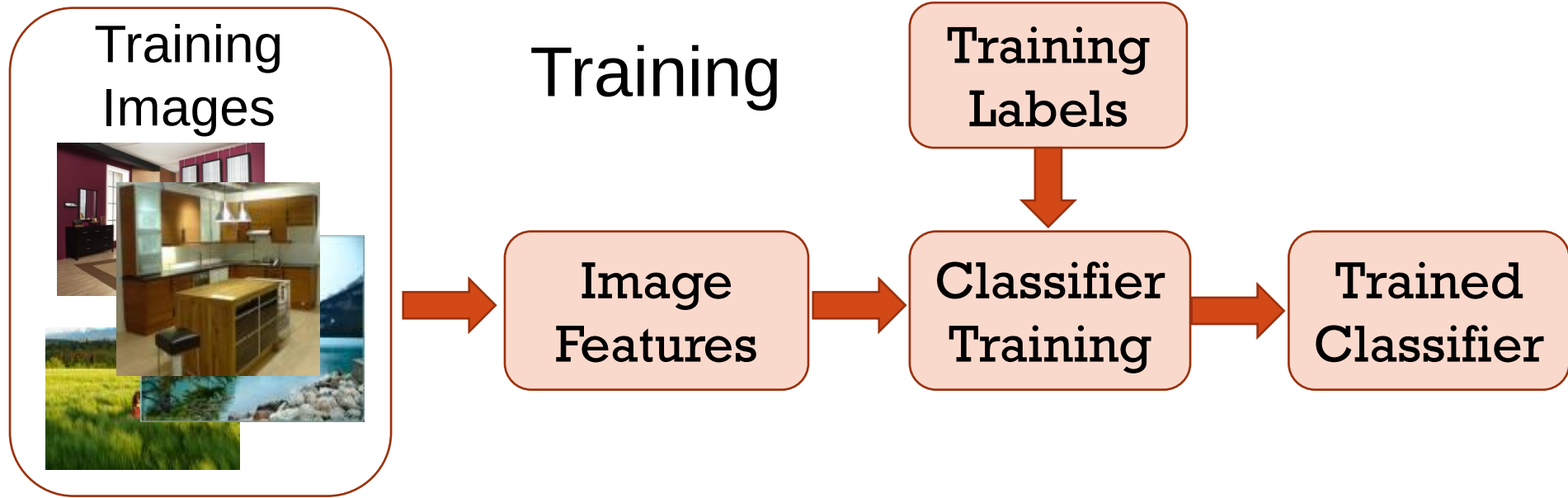
- **Training:** given a *training set* of labeled examples $\{(\mathbf{x}_1, \mathbf{y}_1), \dots, (\mathbf{x}_N, \mathbf{y}_N)\}$, estimate the prediction function $\mathbf{f}$ by minimizing the prediction error on the training set
- **Testing:** apply $\mathbf{f}$ to a never before seen *test example* $\mathbf{x}$ and output the predicted value $\mathbf{y} = \mathbf{f}(\mathbf{x})$



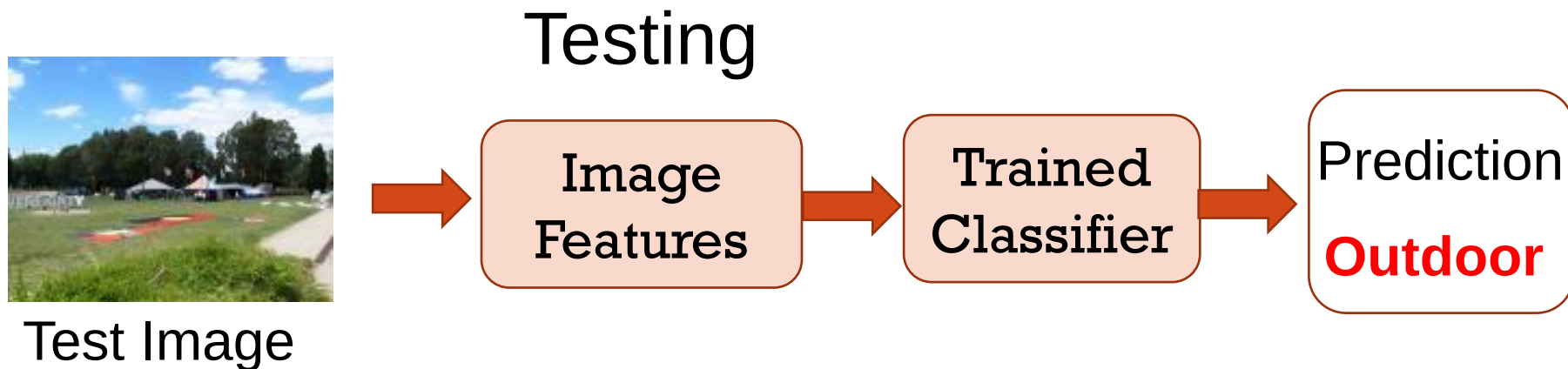
TRAINING PHASE



TRAINING PHASE



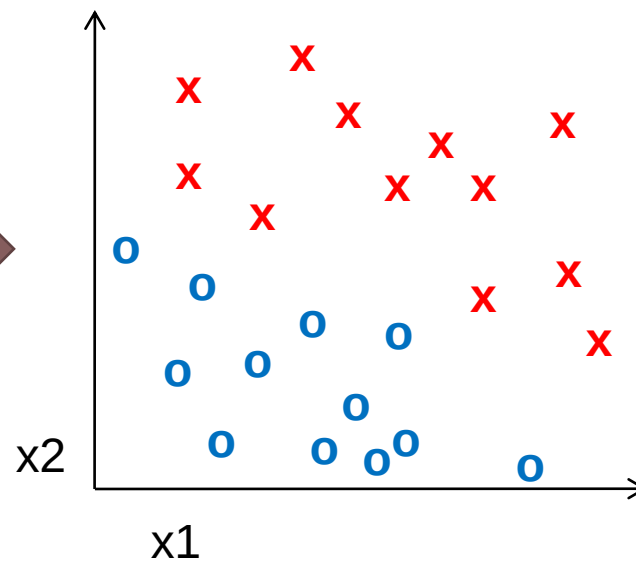
TESTING PHASE



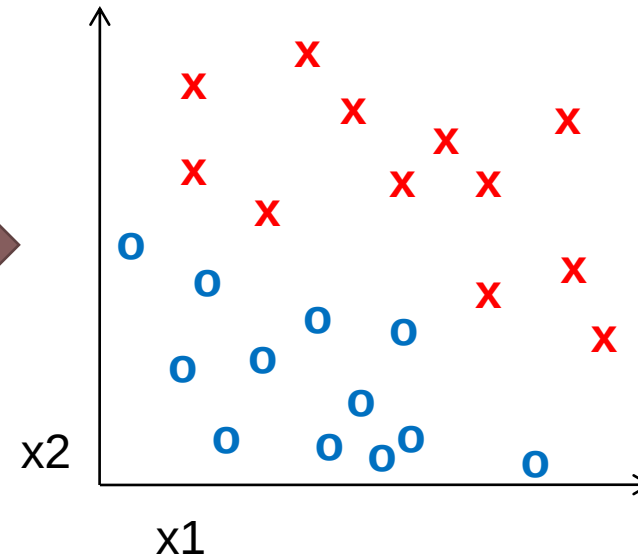
- **Image features:** map images to feature space



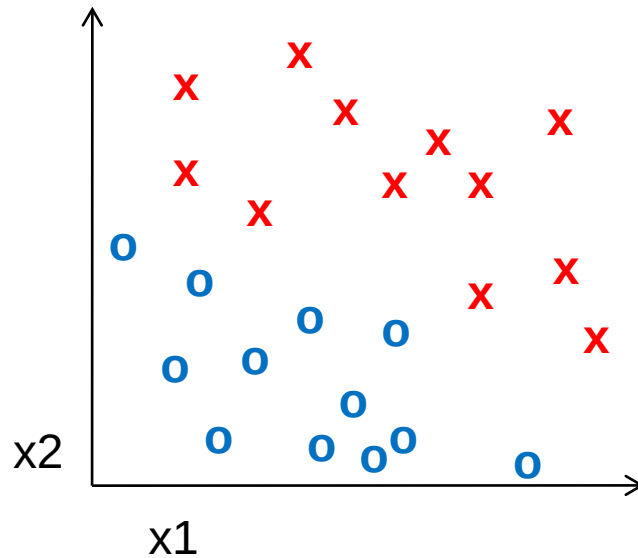
- **Image features:** map images to feature space



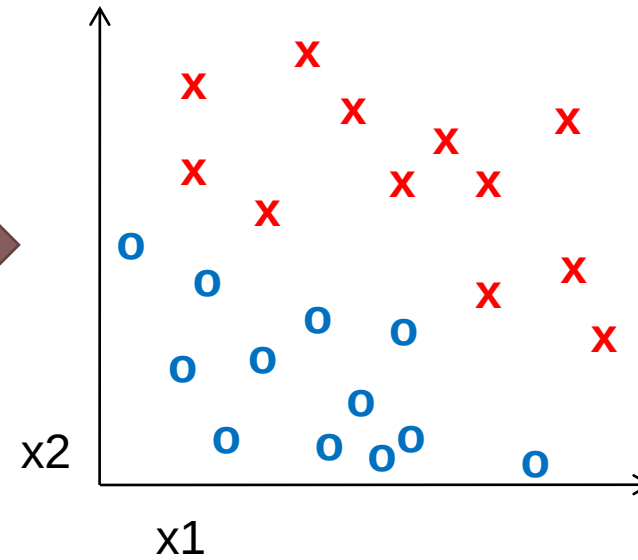
- **Image features:** map images to feature space



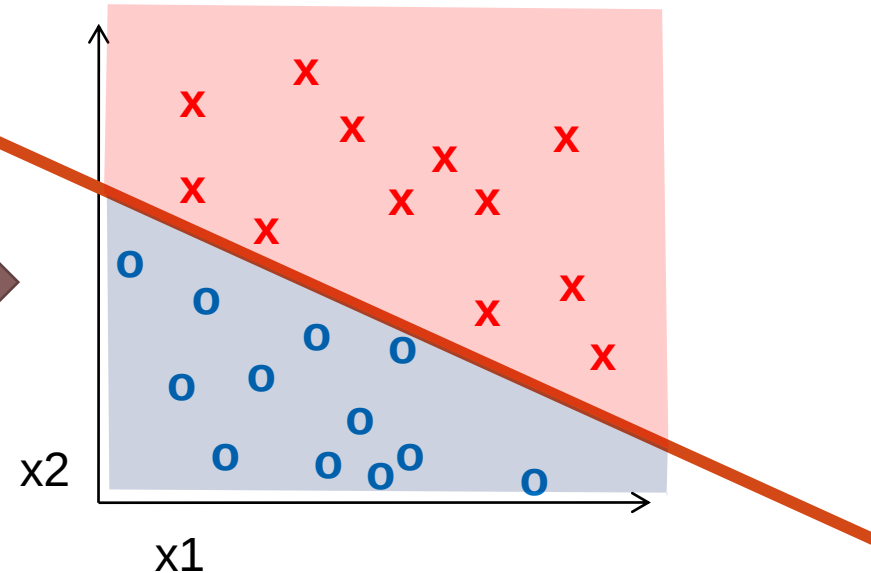
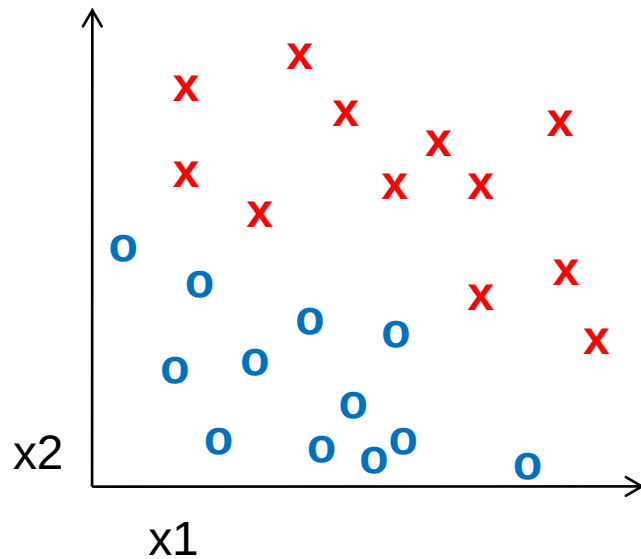
- **Classifiers:** map feature space to label space



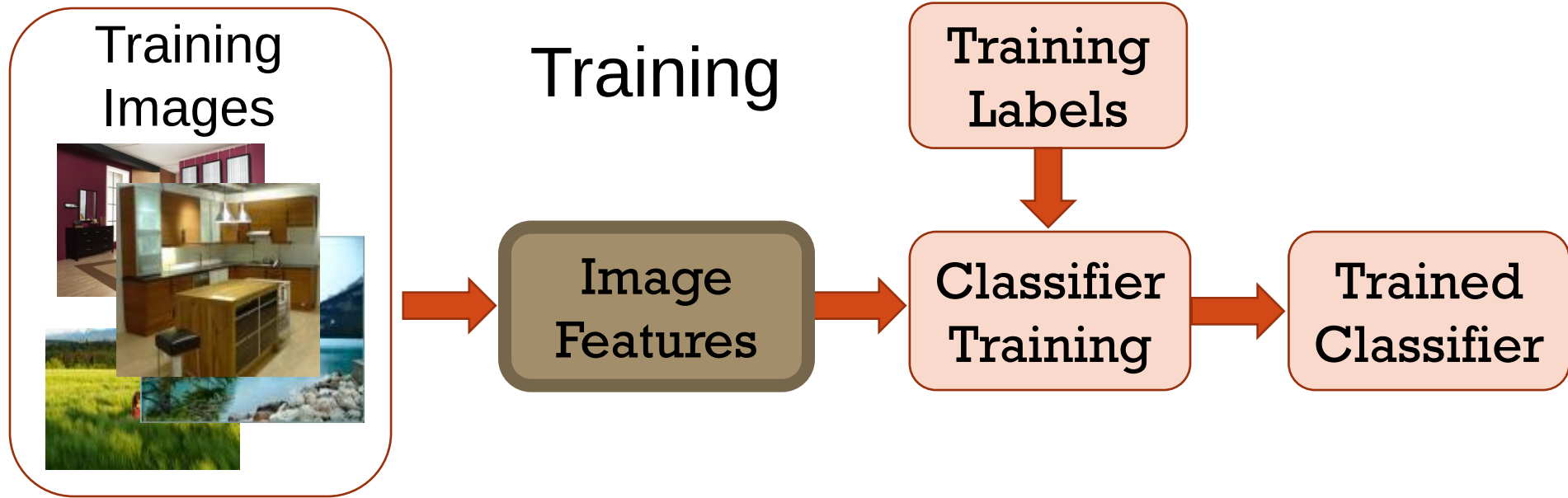
- **Image features:** map images to feature space



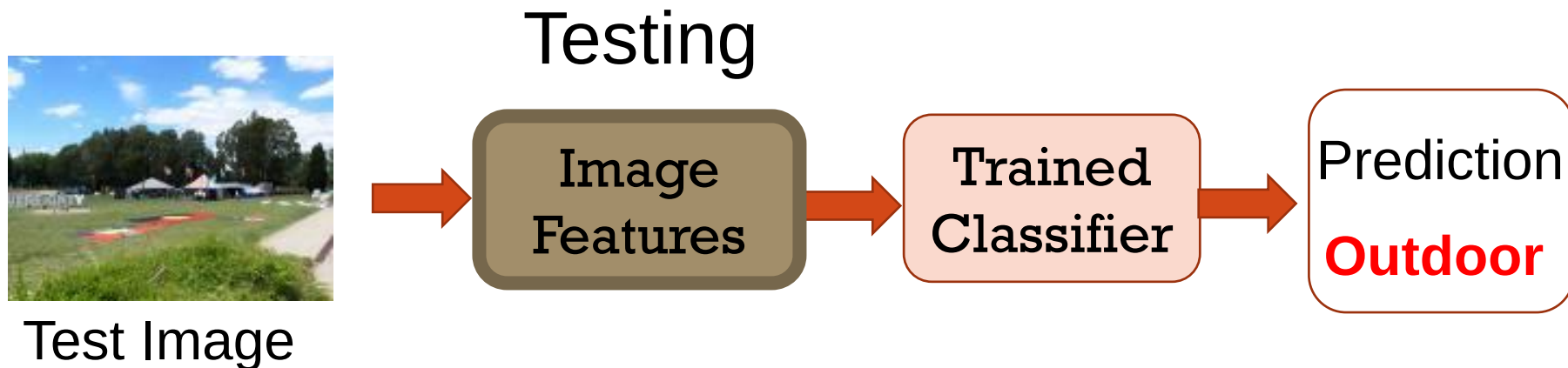
- **Classifiers:** map feature space to label space



TRAINING PHASE



Testing phase



Q: WHAT ARE GOOD FEATURES FOR...

- Recognizing a beach?



Q: WHAT ARE GOOD FEATURES FOR...

- Recognizing cloth fabric?



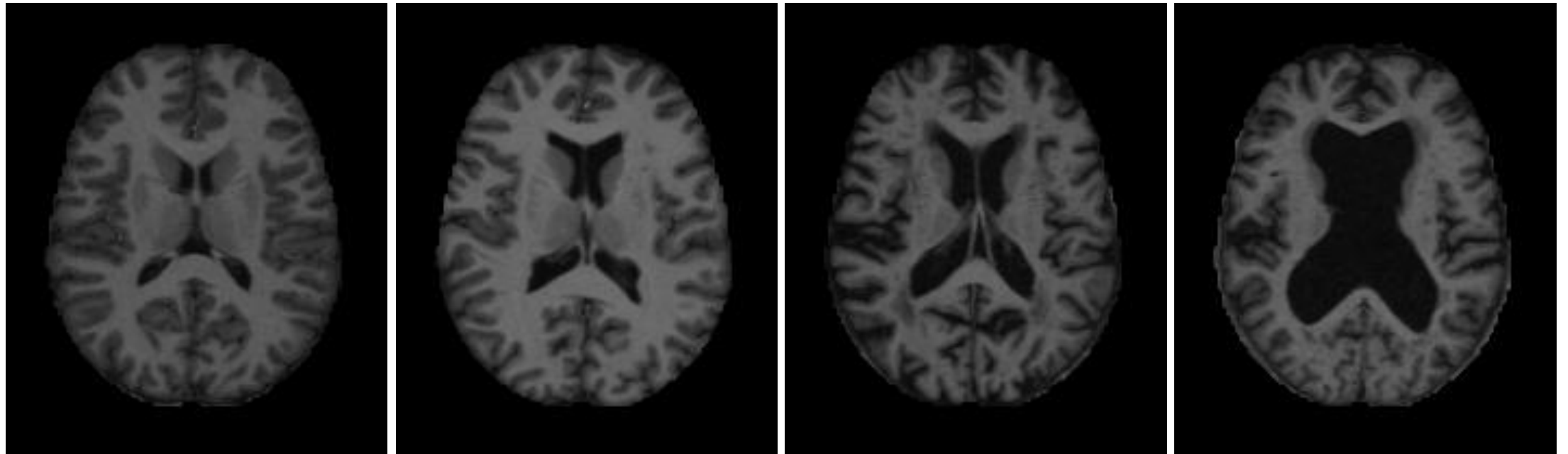
Q: WHAT ARE GOOD FEATURES FOR...

- Recognizing a mug?



Q: WHAT ARE GOOD FEATURES FOR...

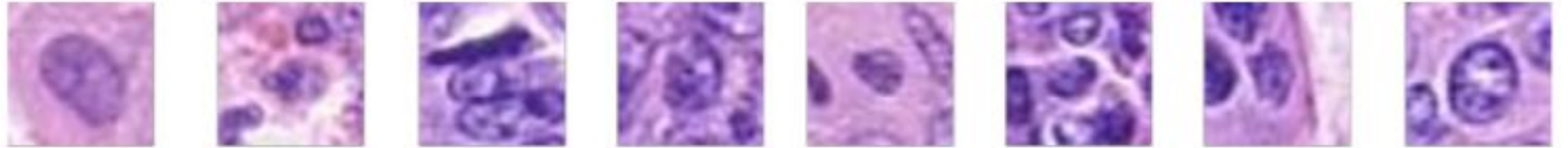
- Recognizing the nodule in MRI data?



Q: WHAT ARE GOOD FEATURES FOR...

- Recognizing the type of colon cancer?

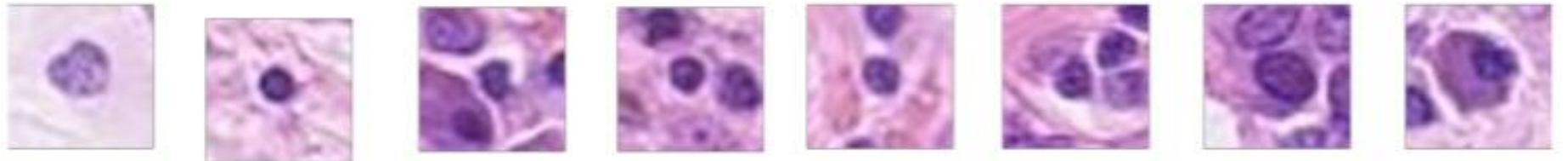
‘Epithelial’



‘Fibroblast’



‘Inflammatory’



‘Miscellaneous’



“CRCHistoPhenotypes” dataset images



WHAT ARE THE RIGHT FEATURES?

Depend on what you want to know!

- **Object: shape**
 - Local shape info, shading, shadows, texture
- **Scene: geometric layout**
 - Linear perspective, gradients, line segments
- **Material properties: albedo, feel, hardness**
 - Color, texture
- **Action: motion**
 - Optical flow, tracked points



IMAGE REPRESENTATIONS

- Templates

- Intensity, gradients, etc.



Image
Intensity



Gradient
template

- Histograms

- Color, Texture, SIFT, LBP descriptors, etc.

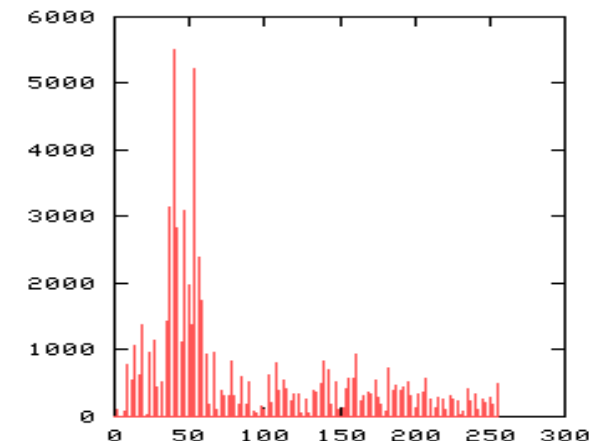
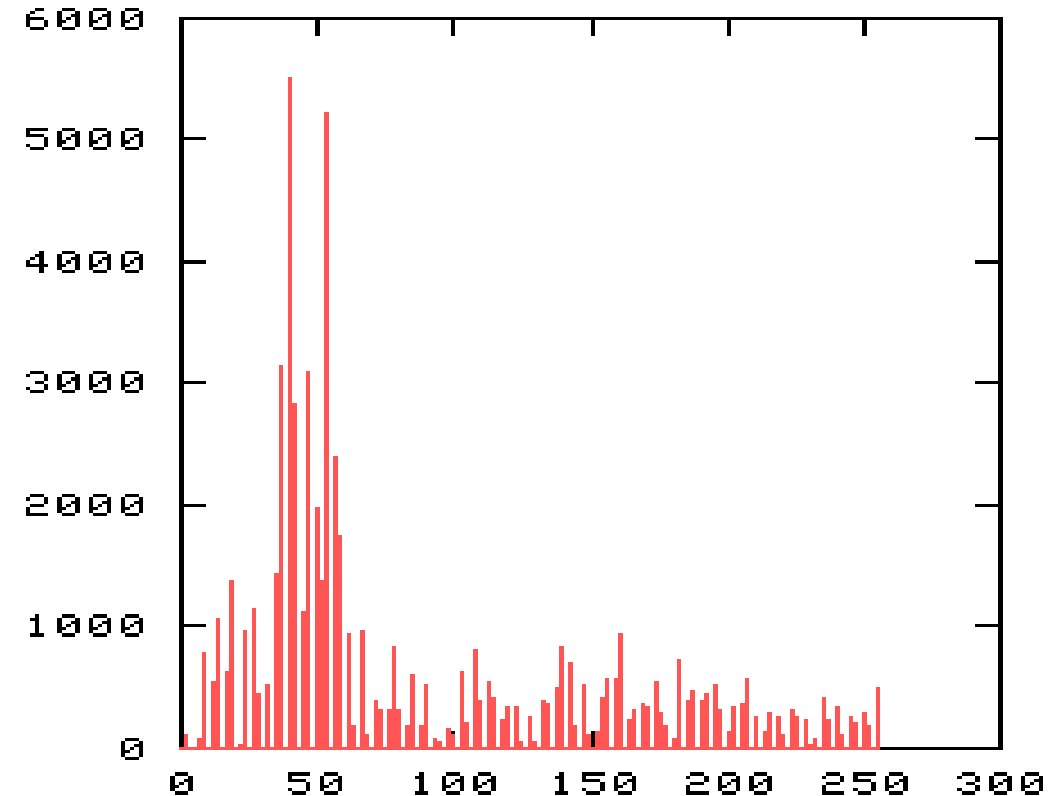


IMAGE REPRESENTATIONS: HISTOGRAMS



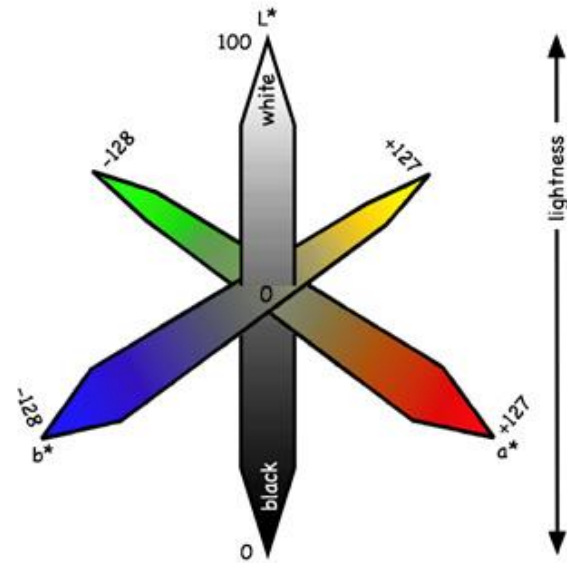
Global histogram

- Represent distribution of features
 - Color, texture, depth, ...

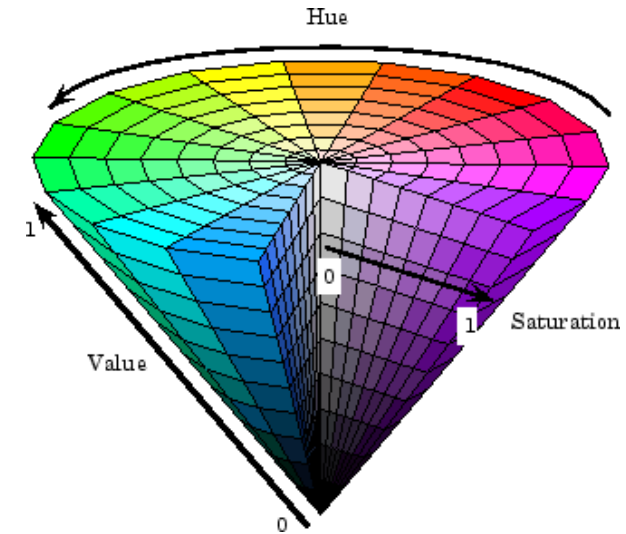


WHAT KIND OF THINGS DO WE COMPUTE HISTOGRAMS OF?

- **Color**

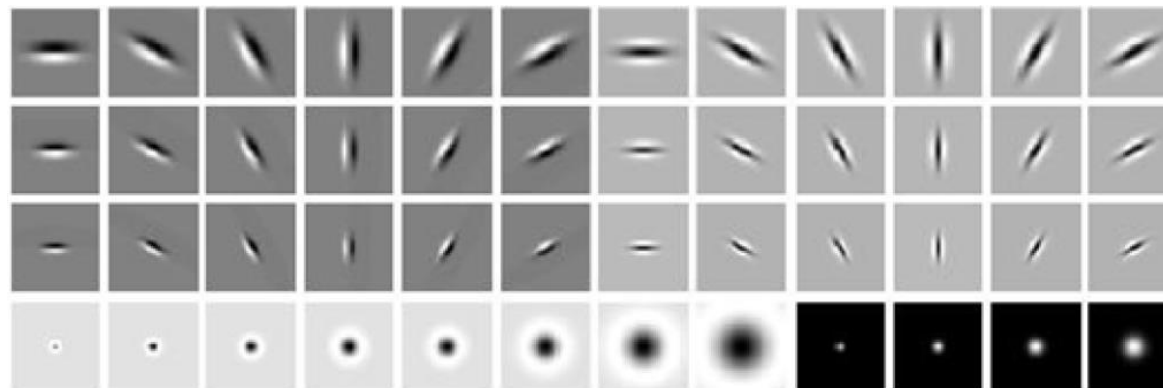


L*a*b* color space

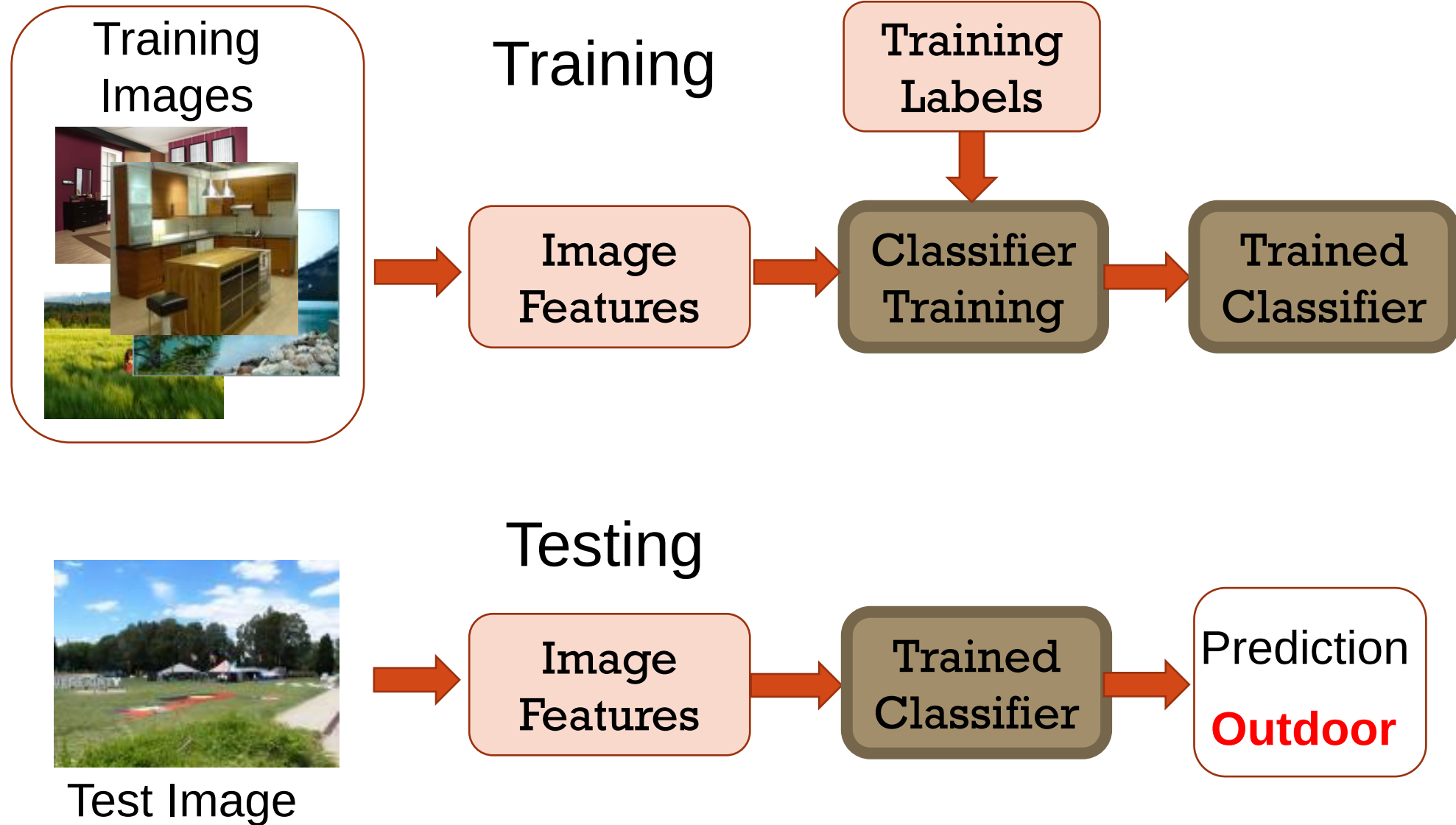


HSV color space

- **Texture** (filter banks or HOG over regions)



CLASSIFIER



CLASSIFIER

K - Nearest Neighbor Classifier

Linear Classifier

Support Vector Machine

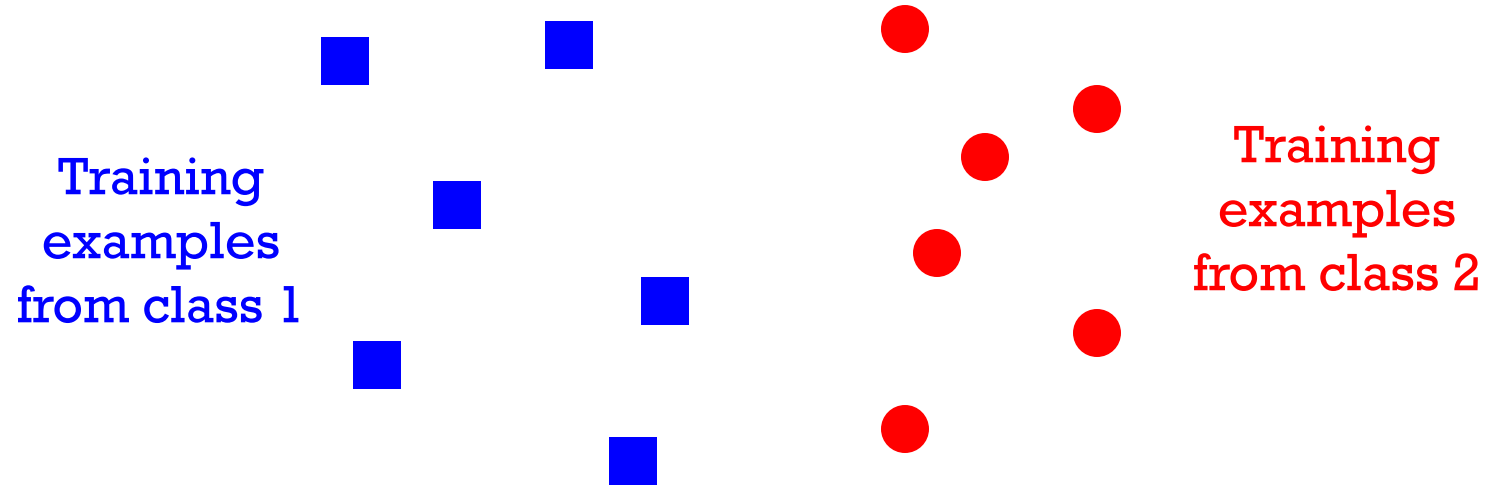
Non-linear SVM

Multi-class SVM

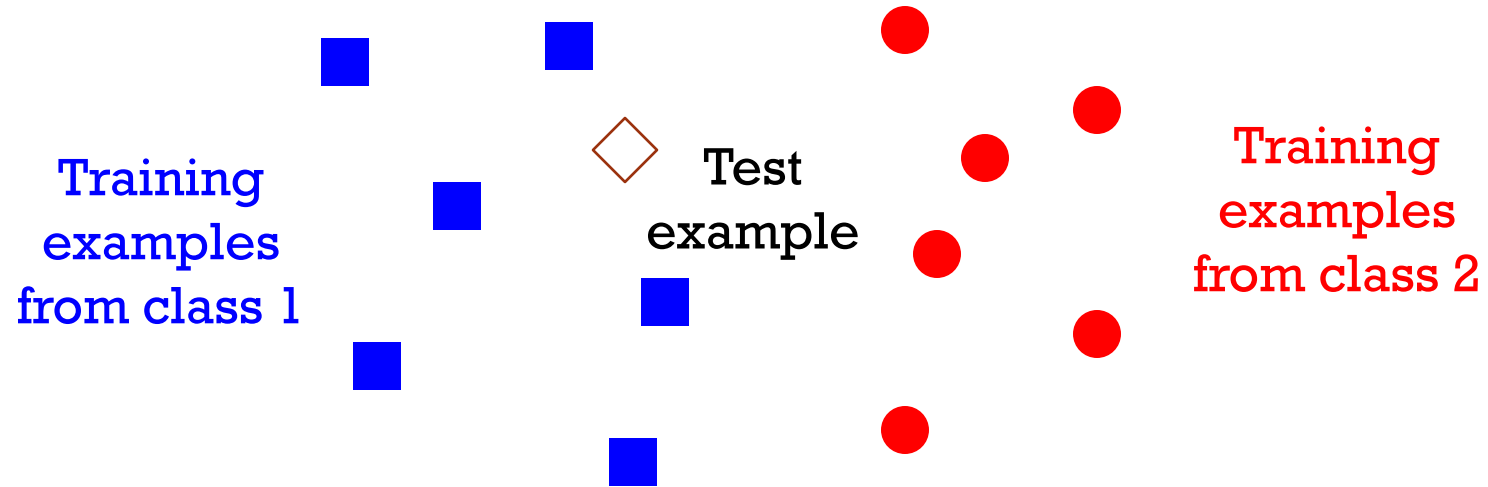
Softmax Classifier



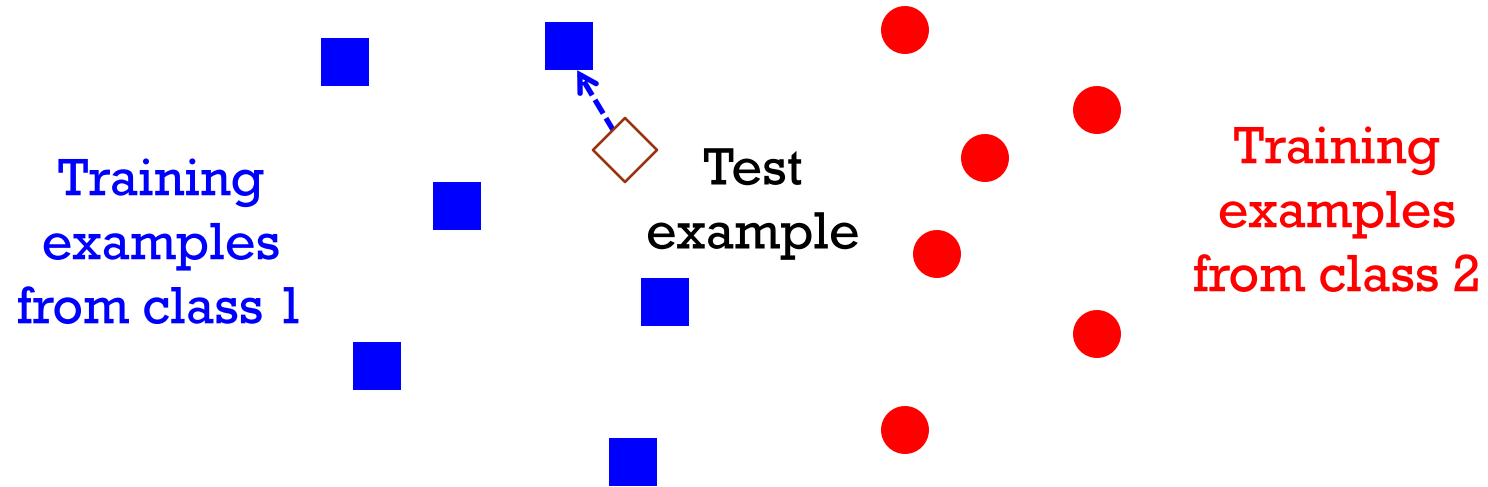
CLASSIFIERS: NEAREST NEIGHBOR



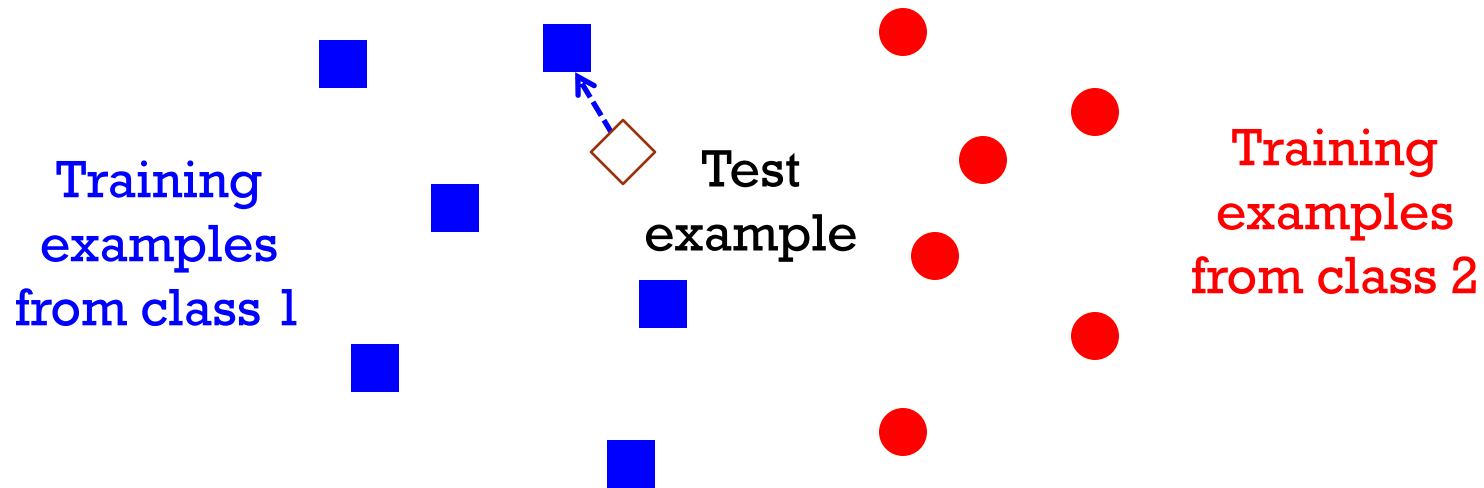
CLASSIFIERS: NEAREST NEIGHBOR



CLASSIFIERS: NEAREST NEIGHBOR



CLASSIFIERS: NEAREST NEIGHBOR



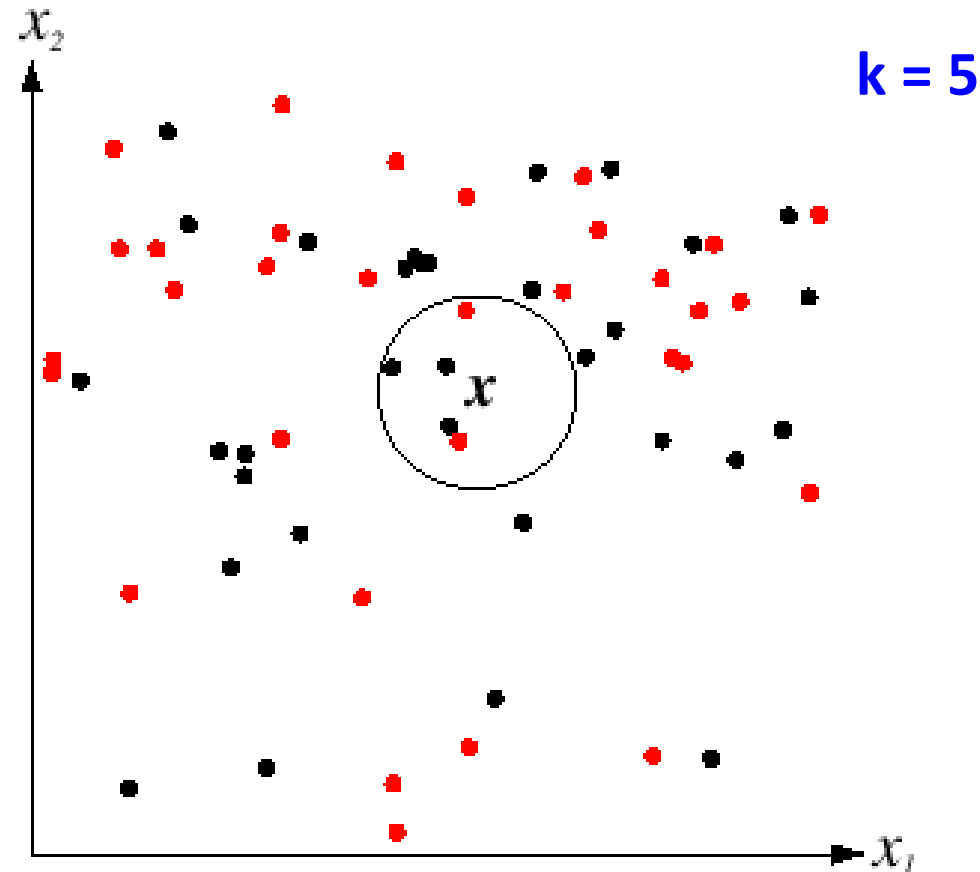
$f(\mathbf{x}) = \text{label of the training example nearest to } \mathbf{x}$

- All we need is a distance function for our inputs
- No training required!



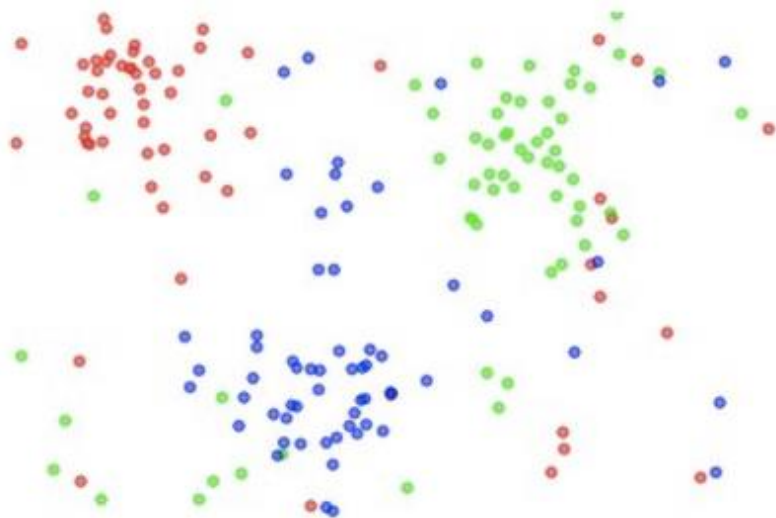
K-NEAREST NEIGHBOR CLASSIFIER

- For a new point, find the k closest points from training data
- Vote for class label with labels of the k points

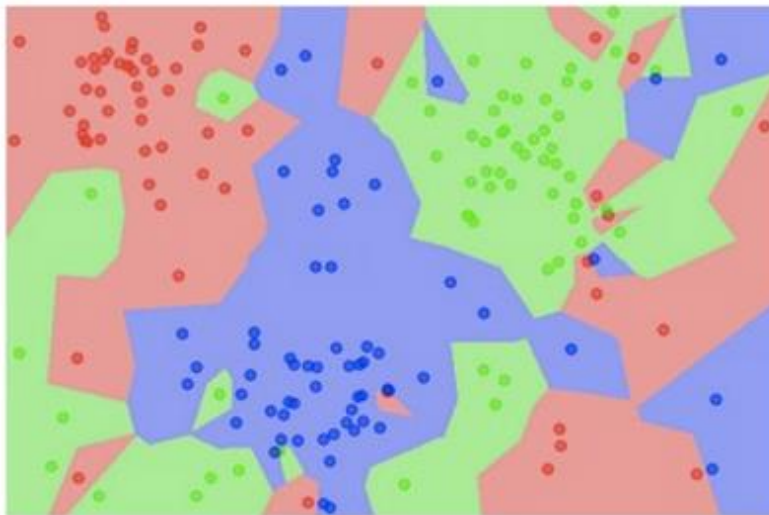


K-NEAREST NEIGHBOR CLASSIFIER

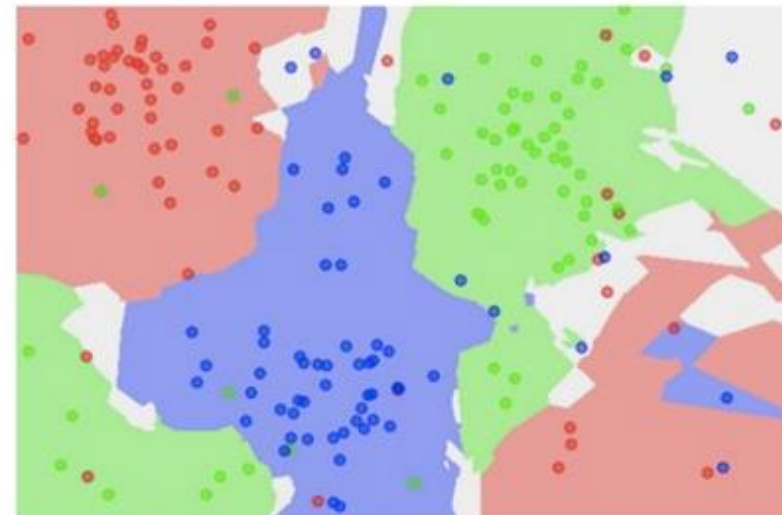
the data



NN classifier

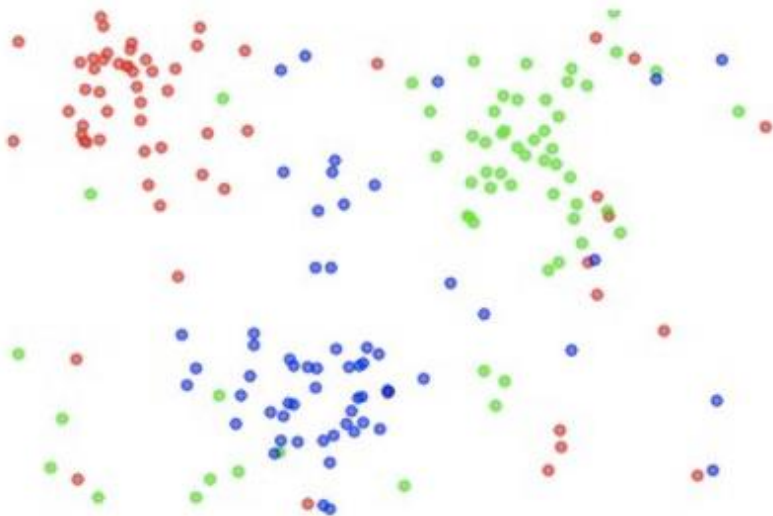


5-NN classifier

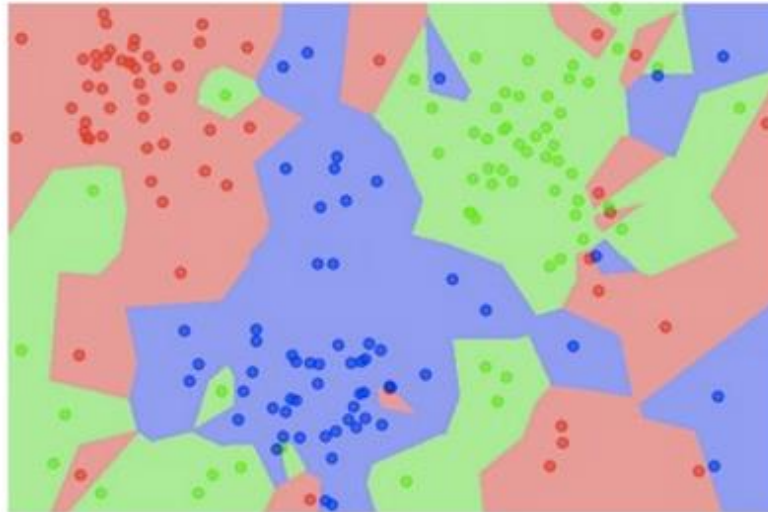


K-NEAREST NEIGHBOR CLASSIFIER

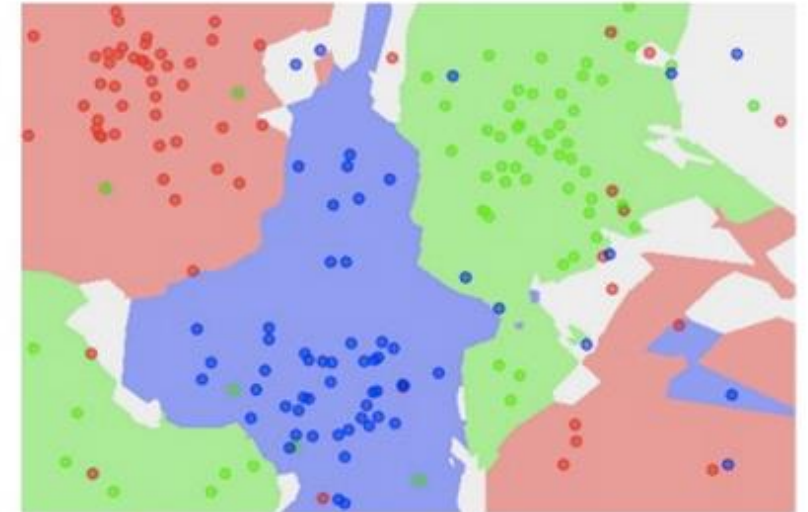
the data



NN classifier



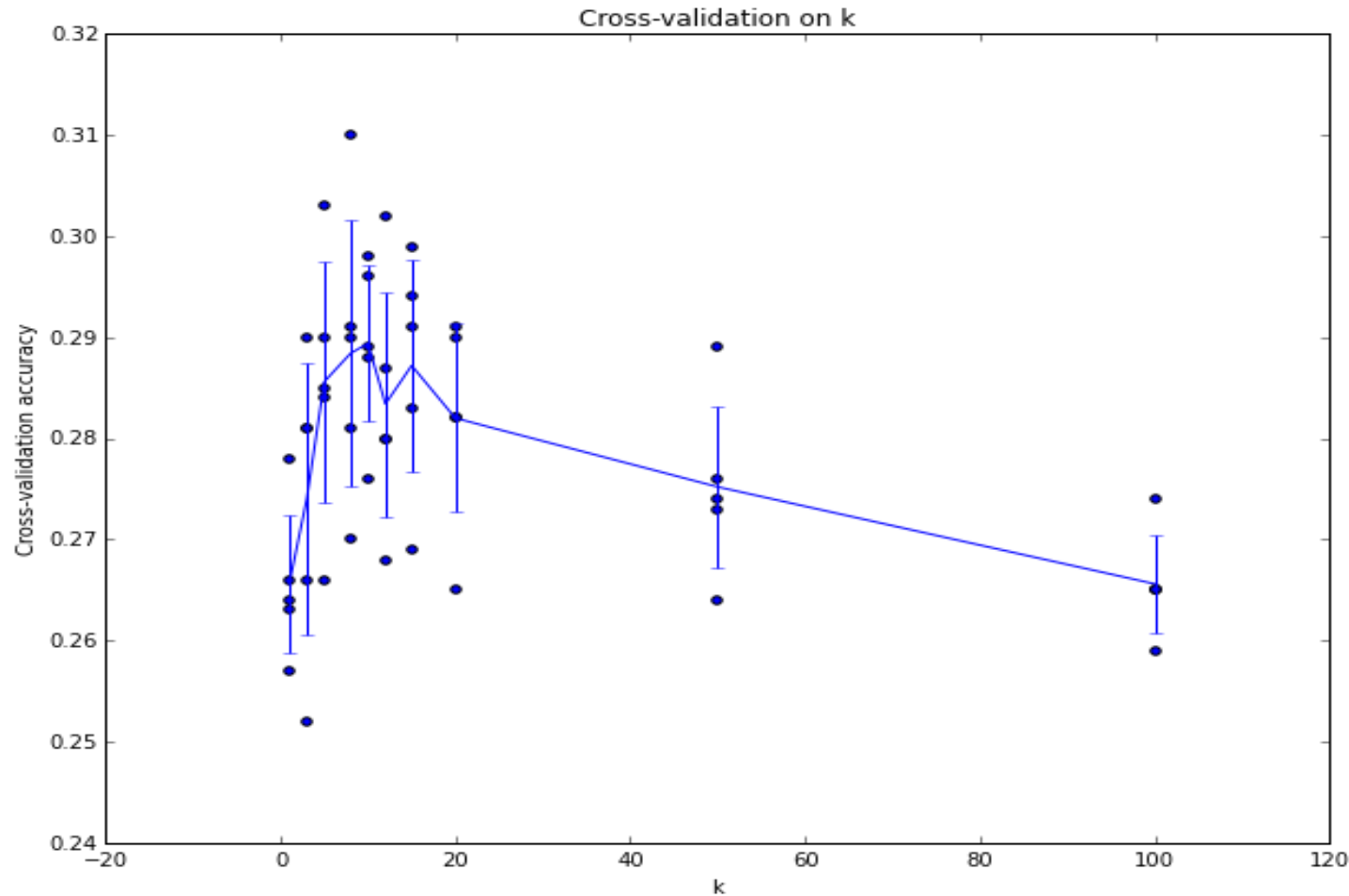
5-NN classifier



- Which classifier is more robust to *outliers*?



CHOICE OF K IN KNN CLASSIFIER



Credit: cs231n,
<http://cs231n.github.io/classification/>

Example of a 5-fold cross-validation run for the parameter k . Note that in this particular case, the cross-validation suggests that a value of about $k = 7$ works best on this particular CIFAR10 dataset (corresponding to the peak in the plot).



CLASSIFIERS: K-NEAREST NEIGHBOR

“Non-parametric” classifier: the entire training set is essentially the model parameters.



CLASSIFIERS: K-NEAREST NEIGHBOR

“Non-parametric” classifier: the entire training set is essentially the model parameters.

Pros:

- **Very fast at training** time
- **Flexible**: all it requires is a way to compute similarity or distances between pairs of features. Applies to many different kinds of features.
- Works with **any number of classes**.
- Works well in practice for large datasets (but see cons)



CLASSIFIERS: K-NEAREST NEIGHBOR

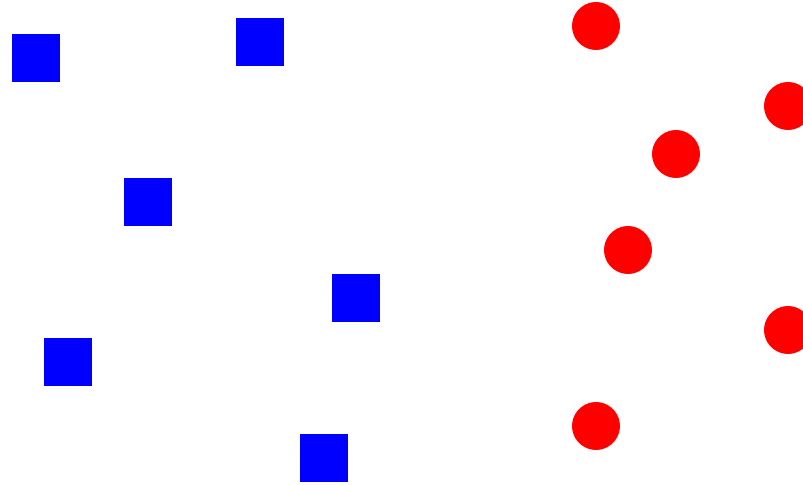
“Non-parametric” classifier: the entire training set is essentially the model parameters.

Cons:

- The classifier must ***remember*** all of the training data and store it for future comparisons with the test data.
- This is **space inefficient** because datasets may easily be gigabytes in size.
- **Slow at test time** (need to compute distances between test example and every training example)
- Optimum value of **K is not known**.



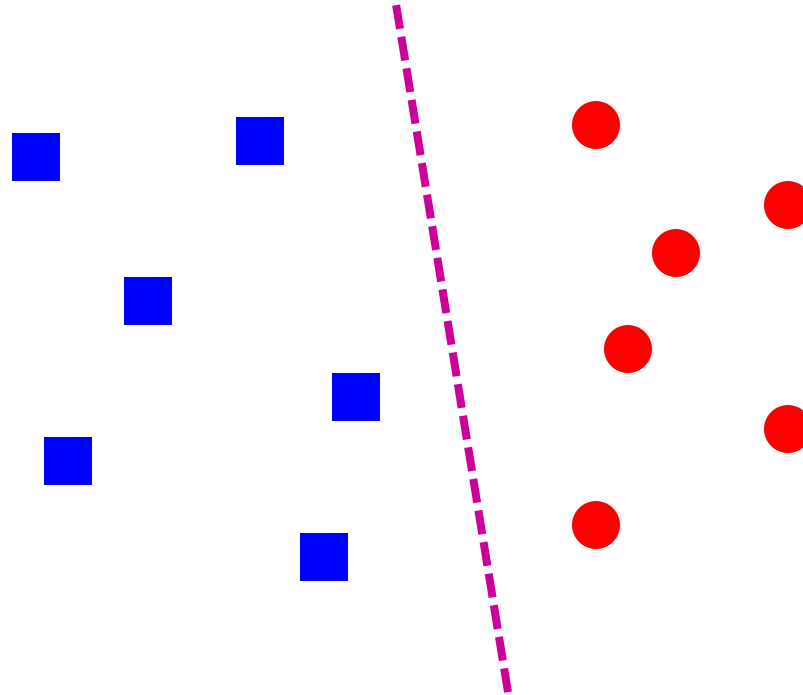
LINEAR CLASSIFIERS — 2 CLASS PROBLEM



- Find a *linear function* to separate the classes:



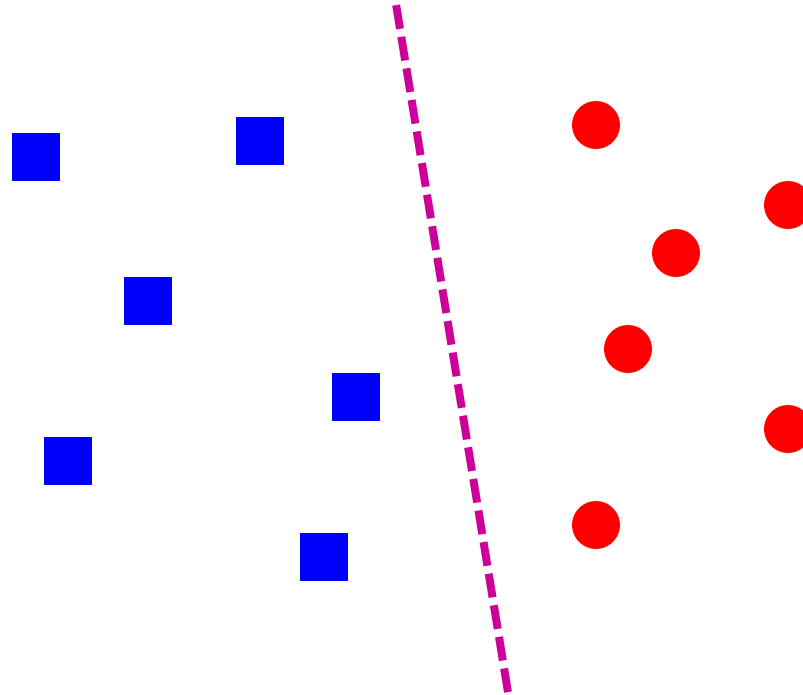
LINEAR CLASSIFIERS — 2 CLASS PROBLEM



- Find a *linear function* to separate the classes:



LINEAR CLASSIFIERS – 2 CLASS PROBLEM

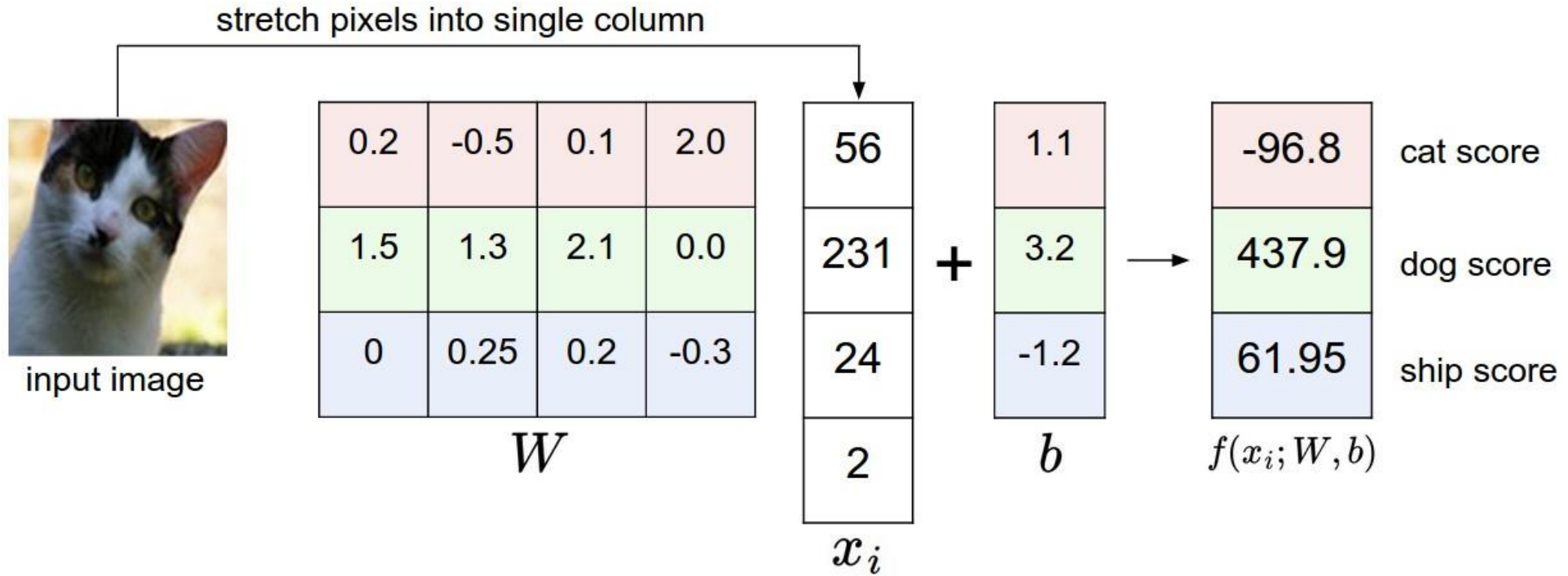


- Find a *linear function* to separate the classes:

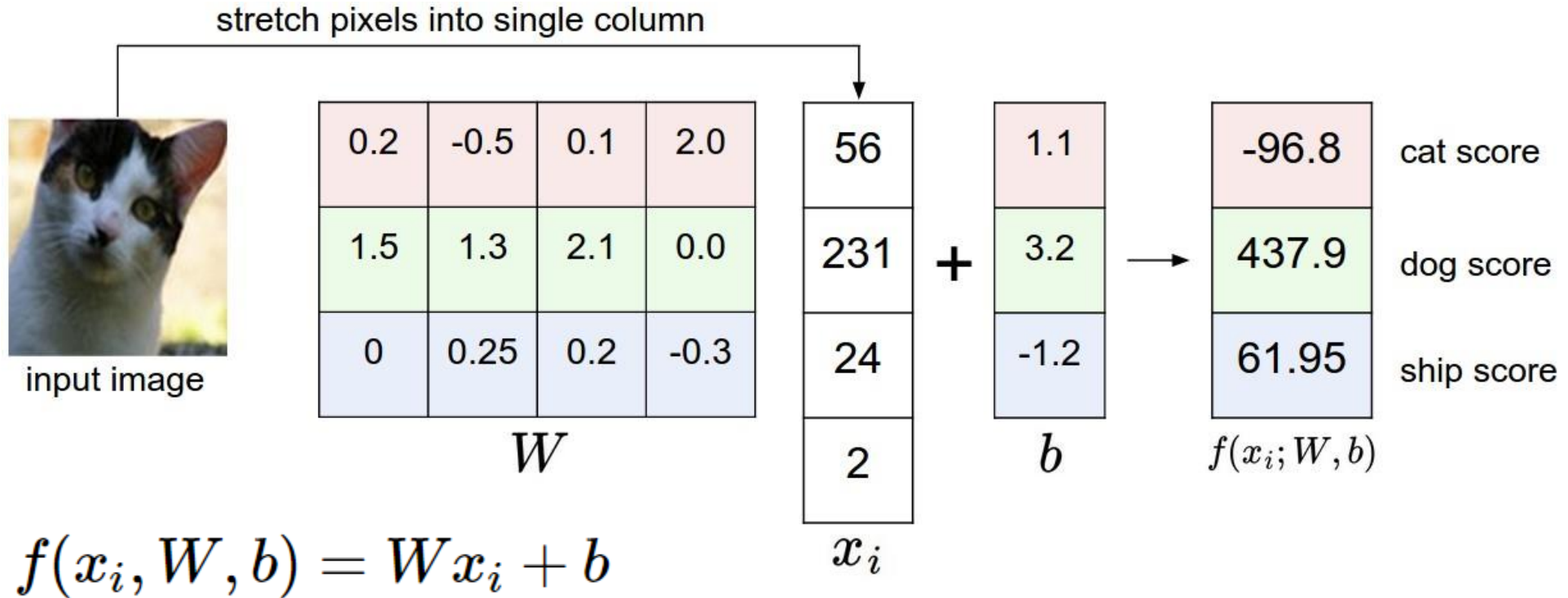
$$f(\mathbf{x}) = \text{sign}(\mathbf{w} \cdot \mathbf{x} + b)$$



LINEAR CLASSIFIERS – MORE THAN 2 CLASS



LINEAR CLASSIFIERS – MORE THAN 2 CLASS



LINEAR CLASSIFIERS – MORE THAN 2 CLASS

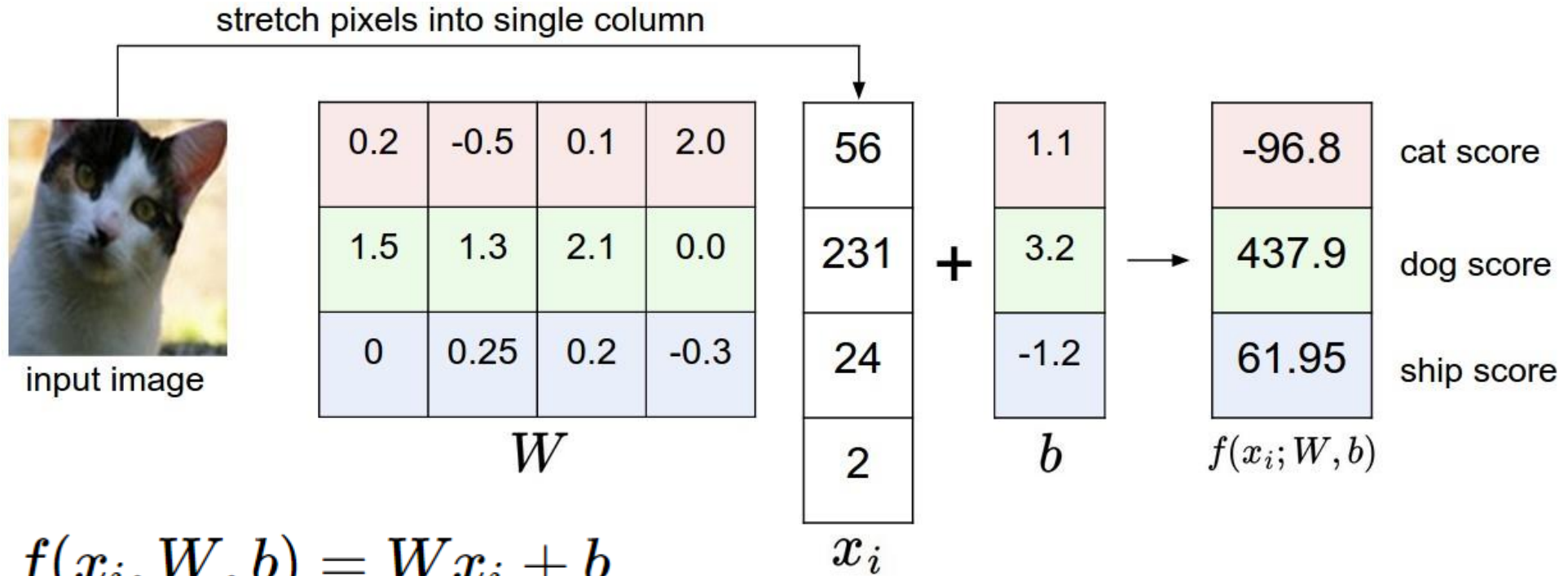


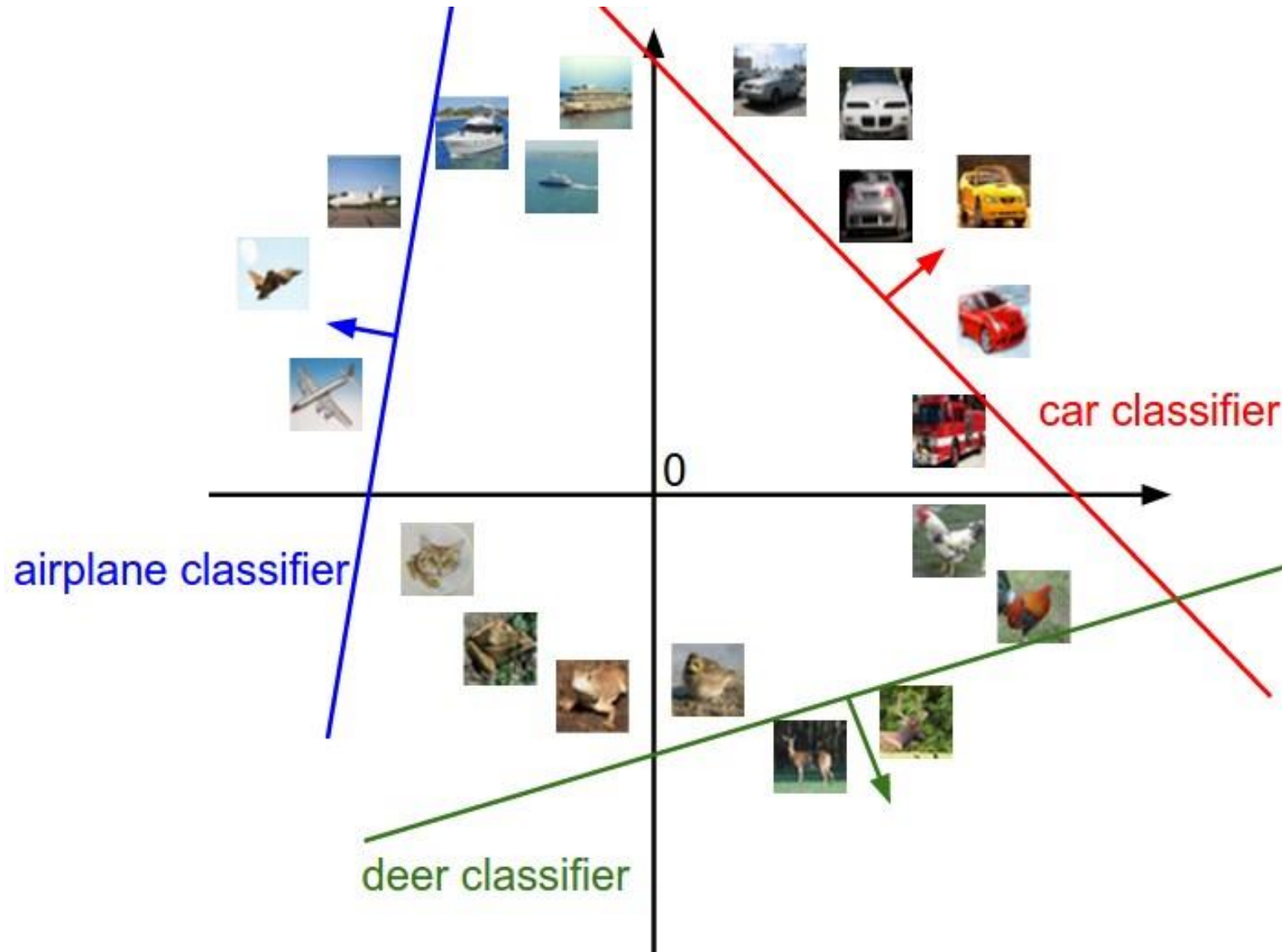
Image x_i has all of its pixels flattened out to a single column vector of shape $[D \times 1]$.

Matrix W (of size $[K \times D]$), and vector b (of size $[K \times 1]$) are the **parameters**.

K is the number of classes.



ANALOGY OF IMAGES AS HIGH-DIMENSIONAL POINTS



Source: cs231n, <http://cs231n.github.io/linear-classify/>



INTERPRETATION OF LINEAR CLASSIFIERS AS TEMPLATE MATCHING

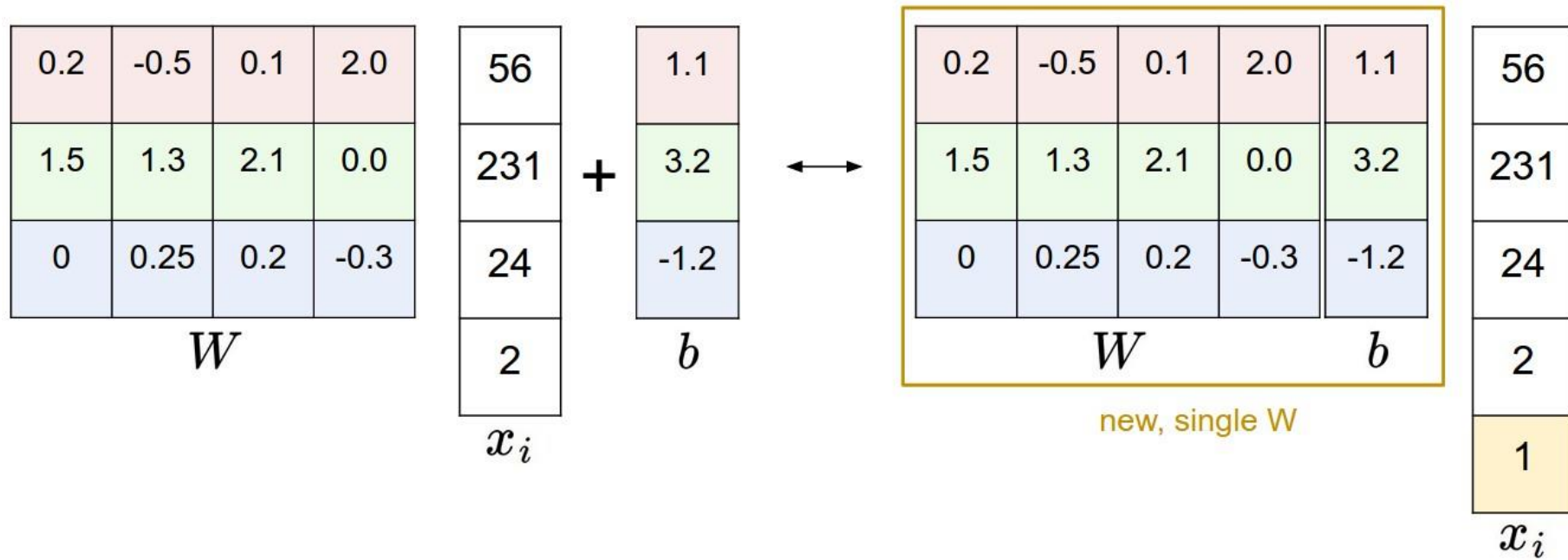


Example learned weights at the end of learning for CIFAR-10. Note that, for example, the ship template contains a lot of blue pixels as expected. This template will therefore give a high score once it is matched against images of ships on the ocean with an inner product.

Source: cs231n, <http://cs231n.github.io/linear-classify/>



BIAS TRICK



LINEAR CLASSIFIERS

“Parametric” classifier: model defined by a small number of parameters (w, b)

Pros:

- Very fast at test time

Cons:

- Slow at training time: need to estimate the parameters
- Data may not be linearly separable



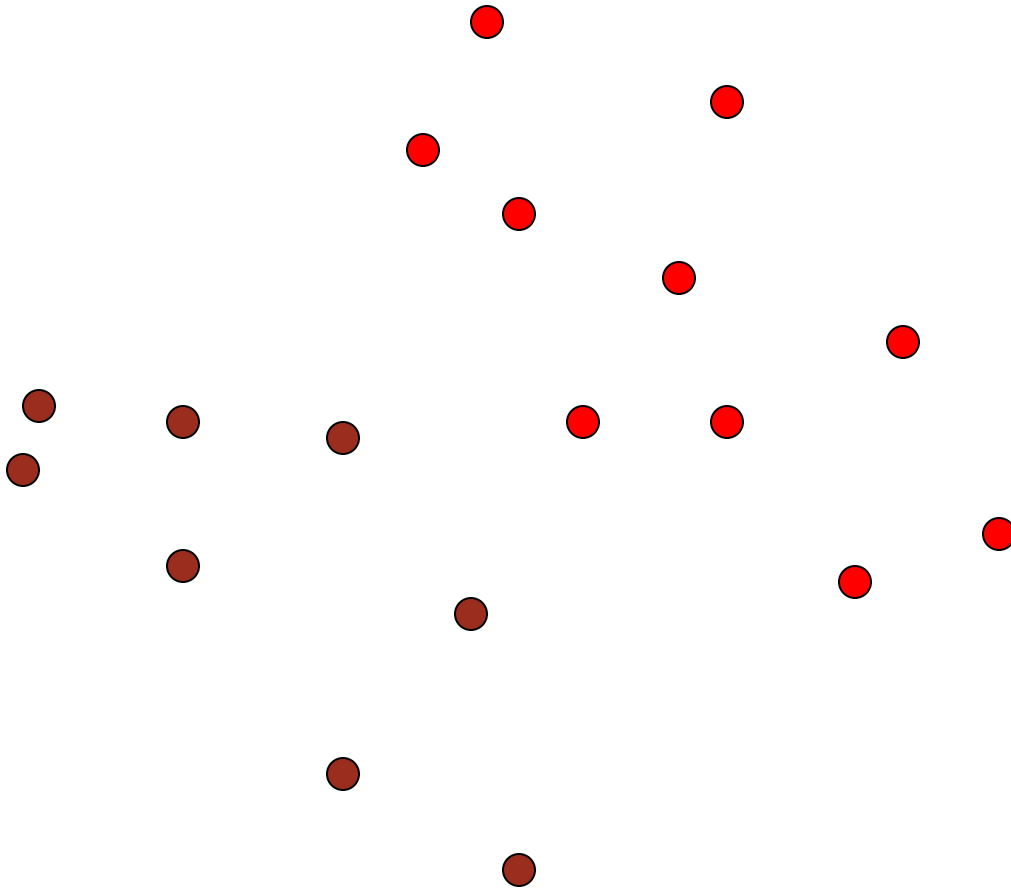
NEAREST NEIGHBOR VS. LINEAR CLASSIFIERS

- **NN pros:**
 - Simple to implement
 - Decision boundaries not necessarily linear
 - Works for any number of classes
 - *Nonparametric* method
- **NN cons:**
 - Need good distance function
 - Slow at test time, Memory in-efficient
- **Linear pros:**
 - Low-dimensional *parametric* representation
 - Very fast at test time
- **Linear cons:**
 - How to train the linear function?
 - What if data is not linearly separable?



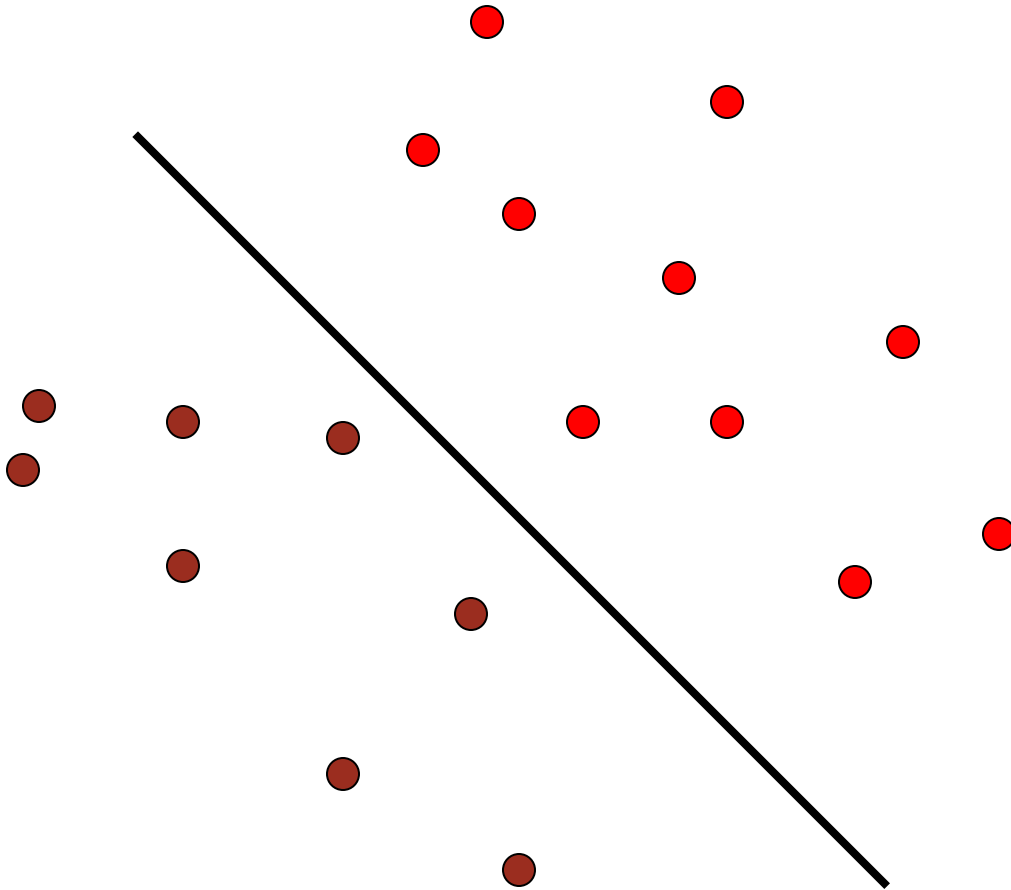
SUPPORT VECTOR MACHINES

- When the data is linearly separable, there may be more than one separator (hyperplane)



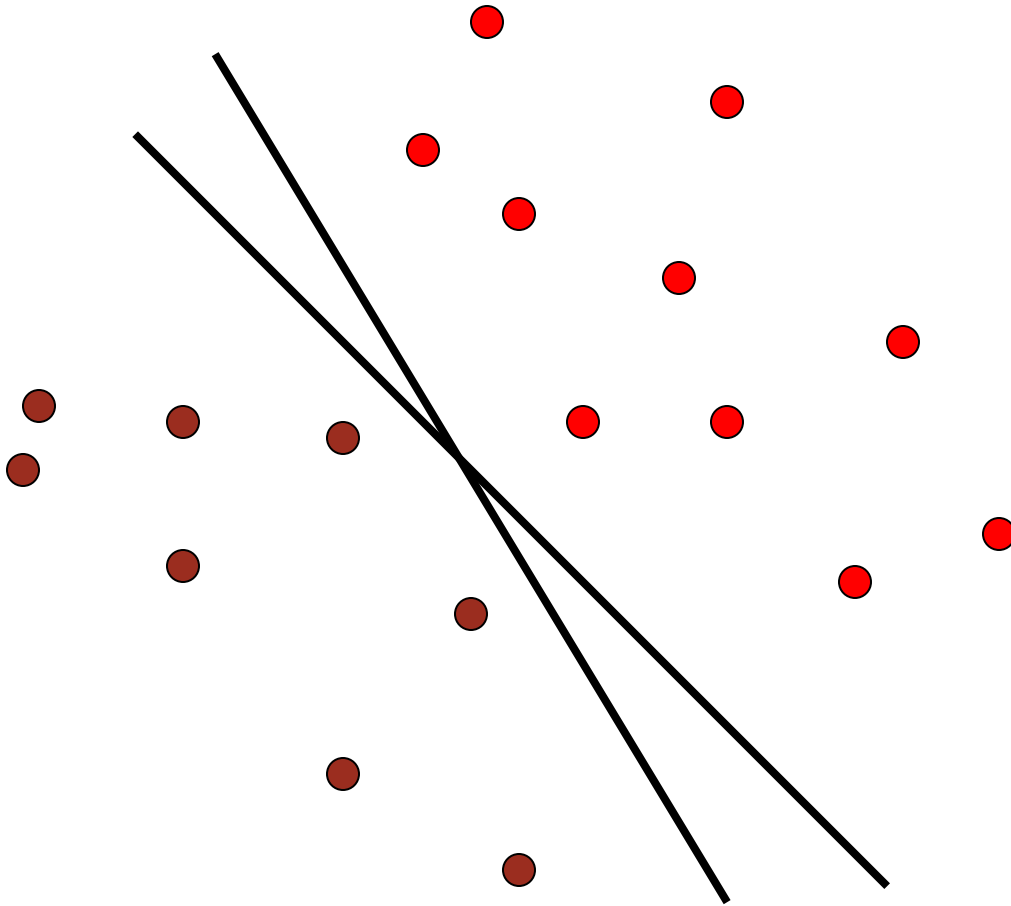
SUPPORT VECTOR MACHINES

- When the data is linearly separable, there may be more than one separator (hyperplane)



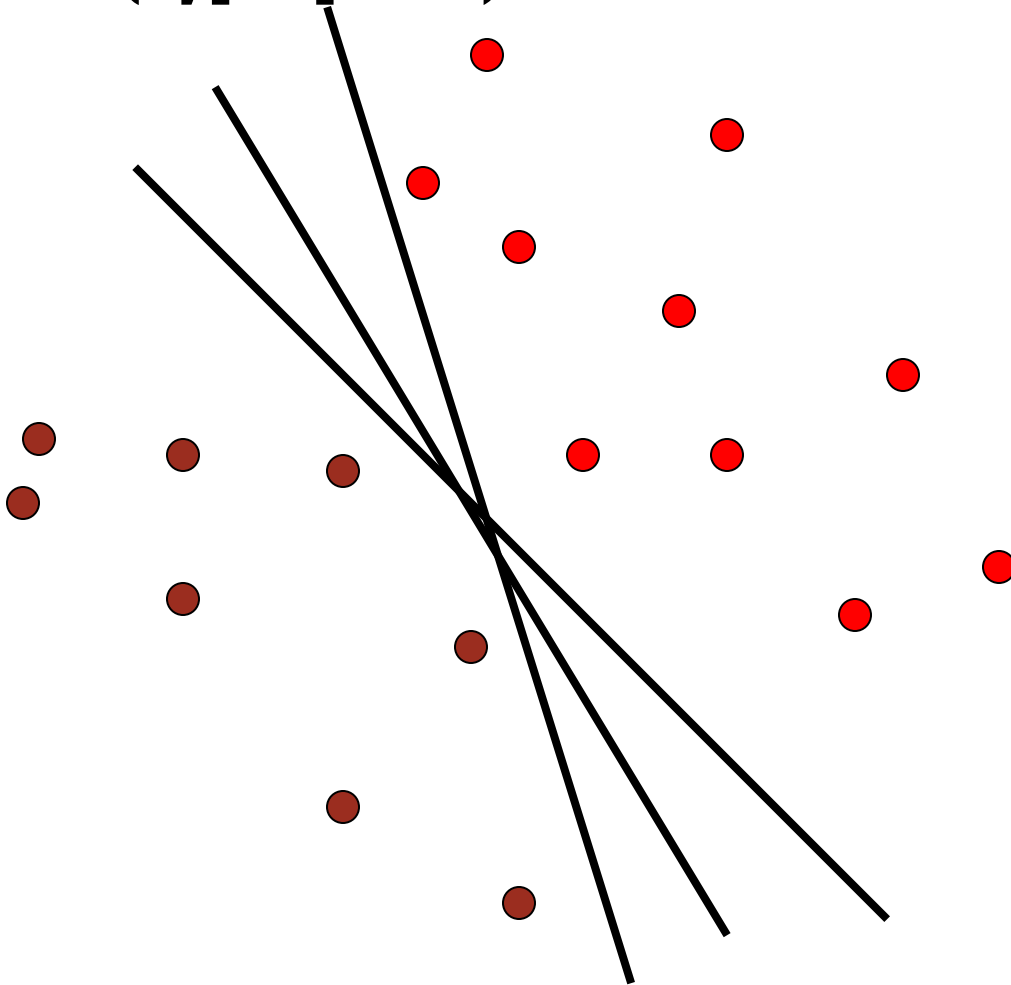
SUPPORT VECTOR MACHINES

- When the data is linearly separable, there may be more than one separator (hyperplane)



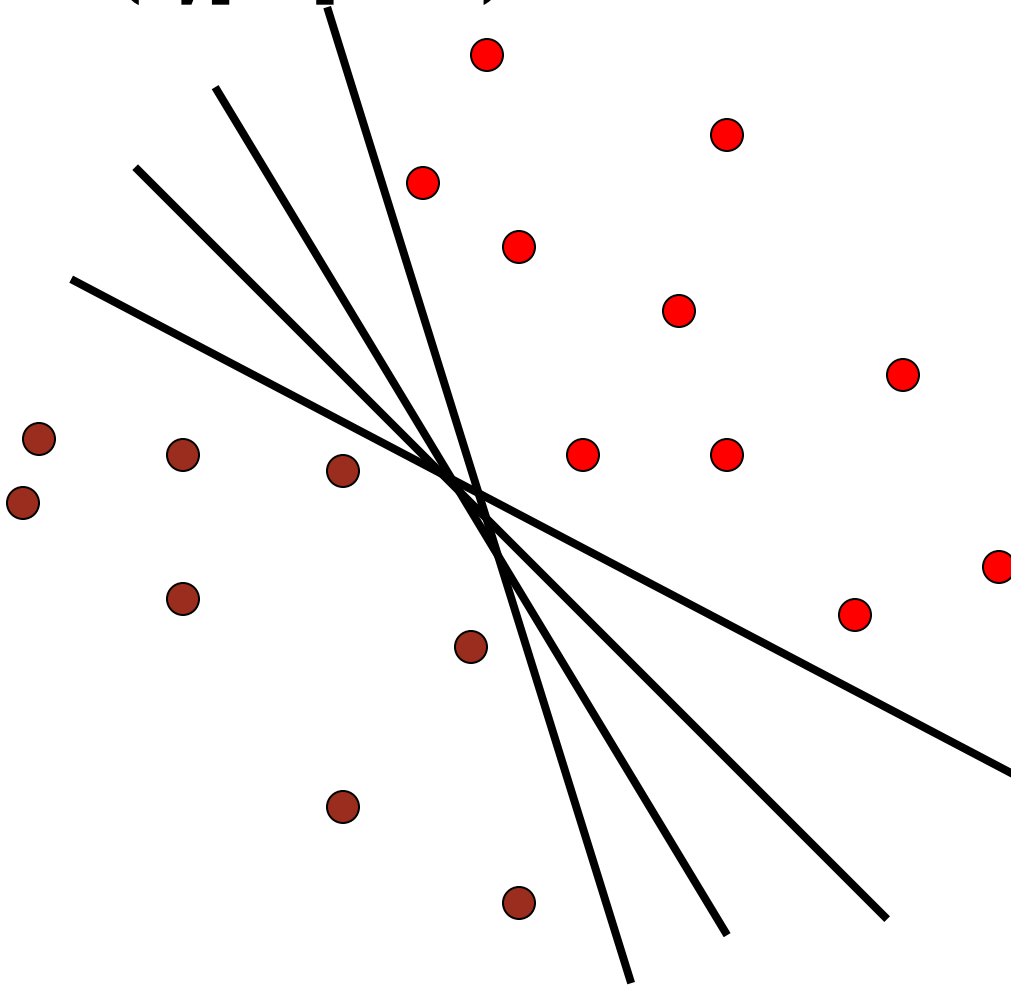
SUPPORT VECTOR MACHINES

- When the data is linearly separable, there may be more than one separator (hyperplane)



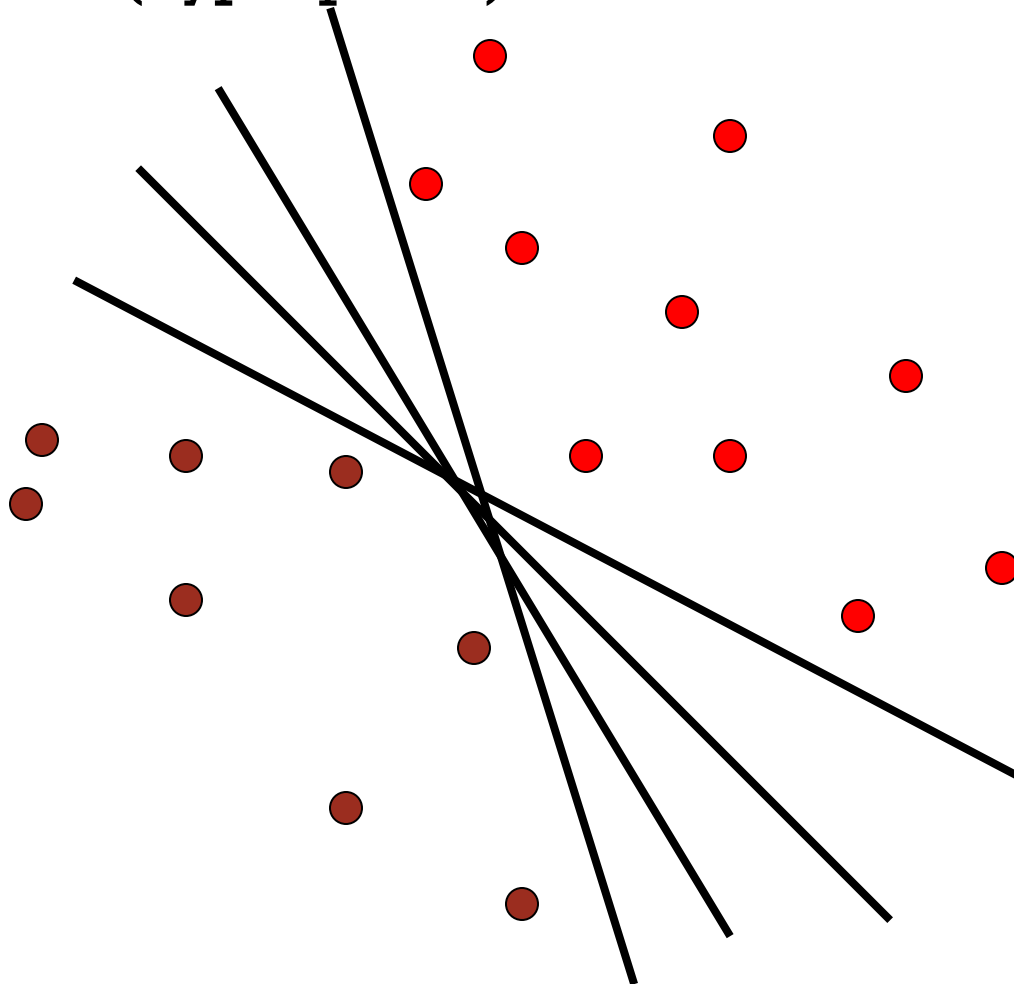
SUPPORT VECTOR MACHINES

- When the data is linearly separable, there may be more than one separator (hyperplane)



SUPPORT VECTOR MACHINES

- When the data is linearly separable, there may be more than one separator (hyperplane)

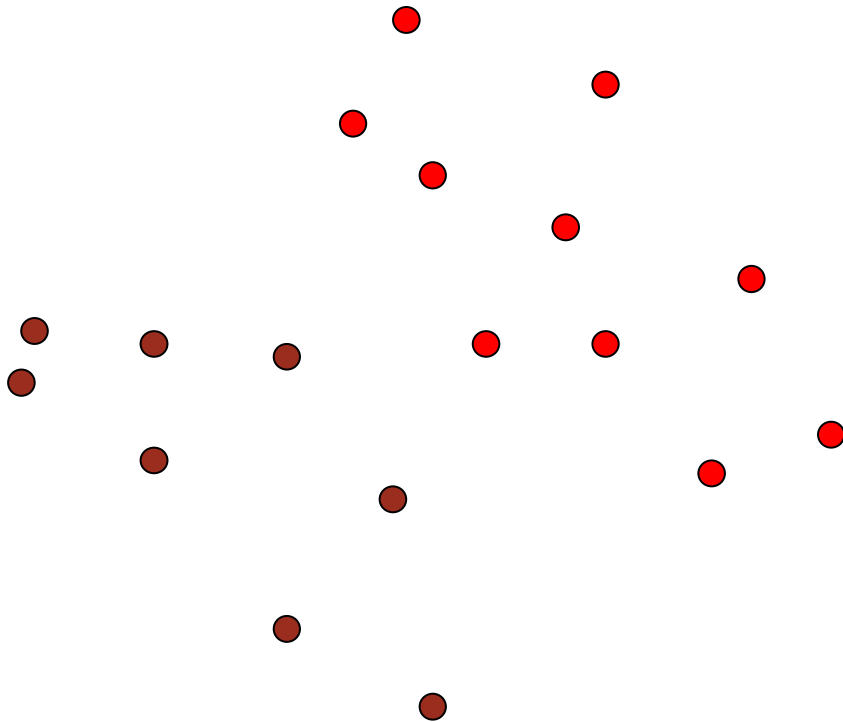


Which separator
is best?



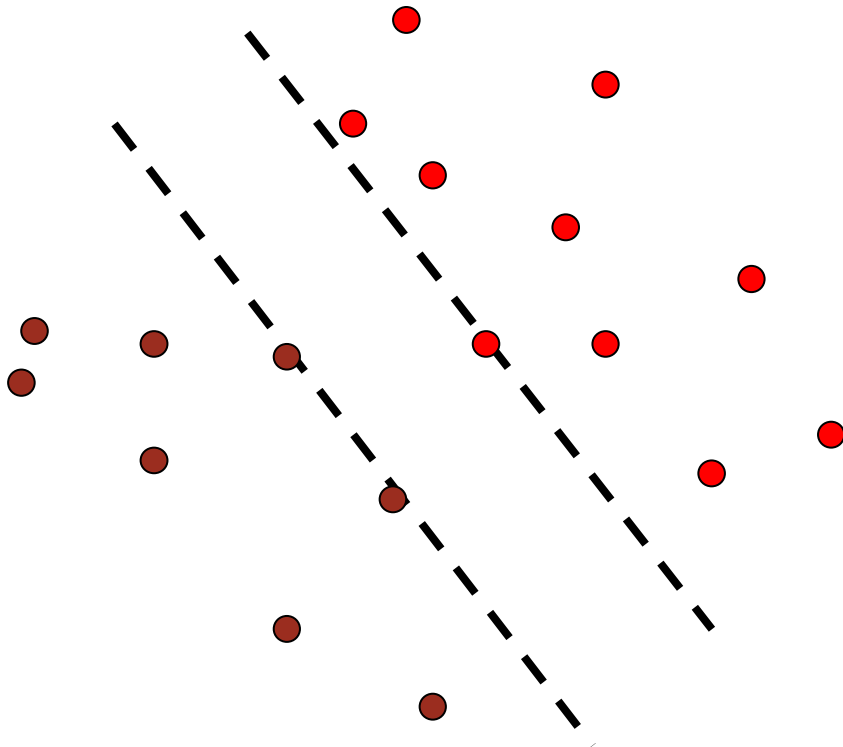
SUPPORT VECTOR MACHINES

- Find hyperplane that maximizes the *margin* between the positive and negative examples



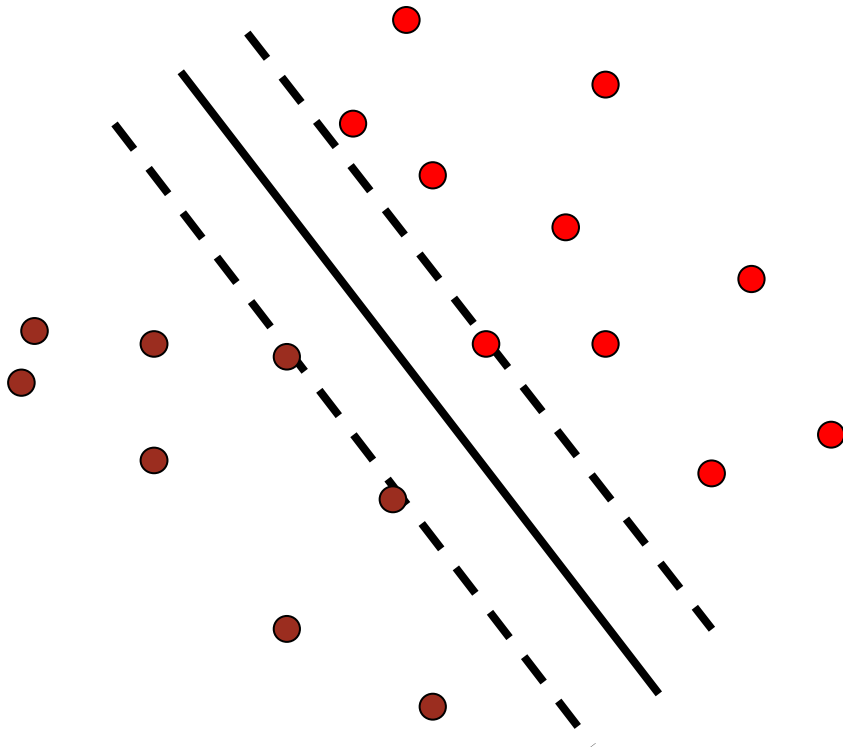
SUPPORT VECTOR MACHINES

- Find hyperplane that maximizes the *margin* between the positive and negative examples



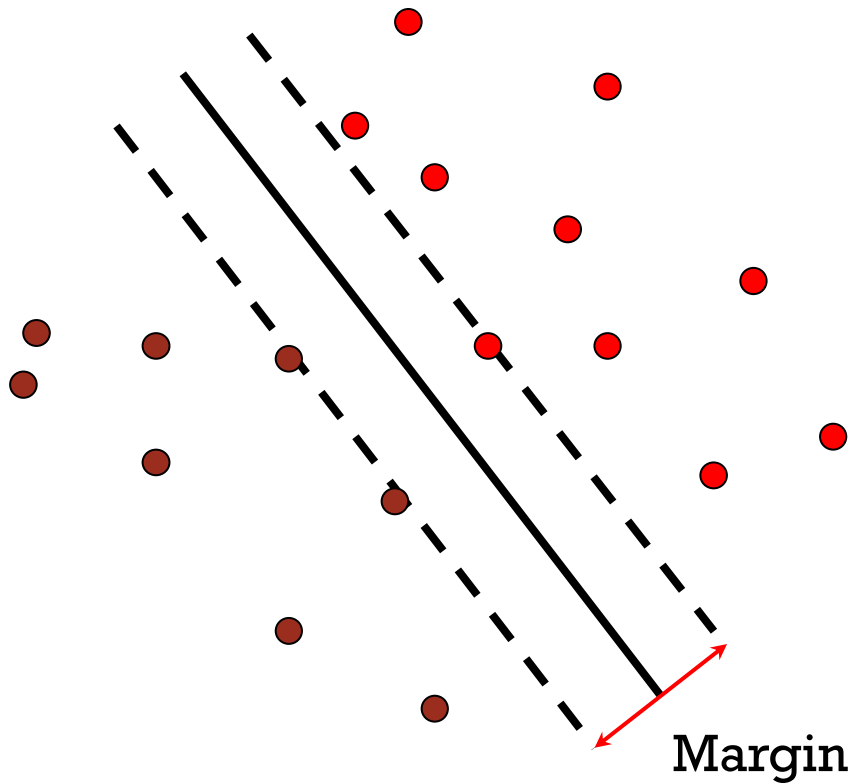
SUPPORT VECTOR MACHINES

- Find hyperplane that maximizes the *margin* between the positive and negative examples



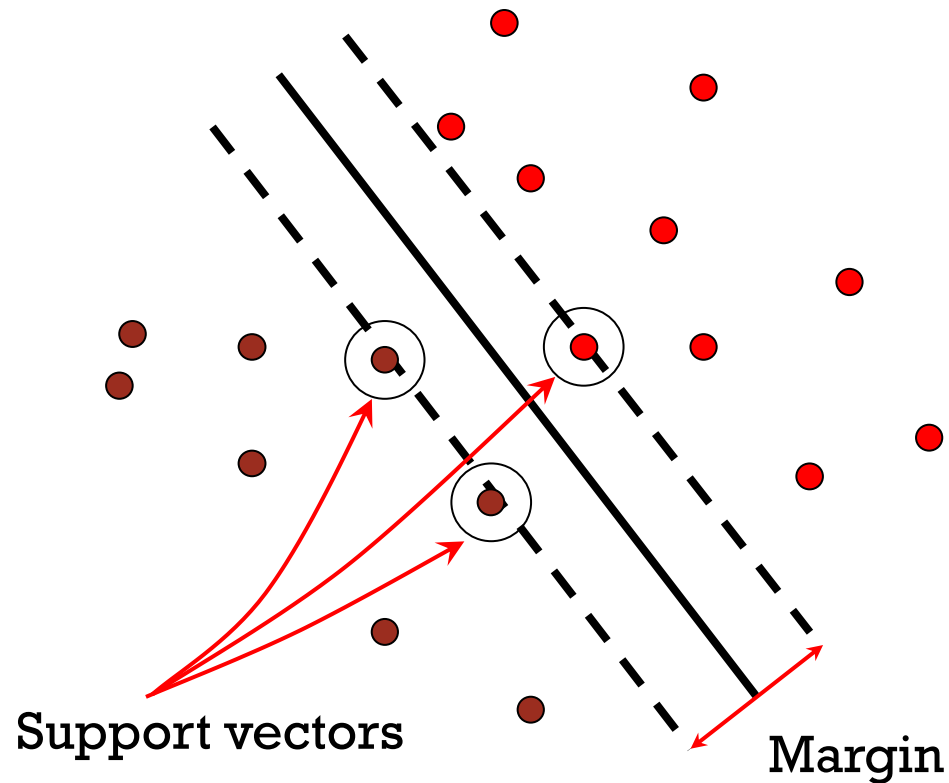
SUPPORT VECTOR MACHINES

- Find hyperplane that maximizes the *margin* between the positive and negative examples



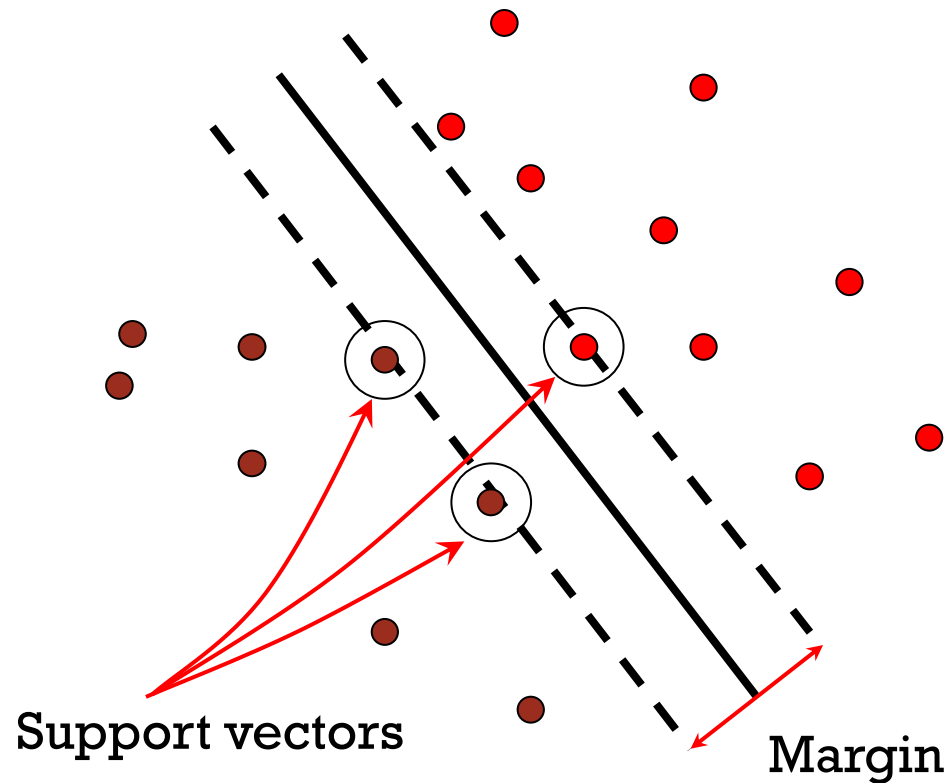
SUPPORT VECTOR MACHINES

- Find hyperplane that maximizes the *margin* between the positive and negative examples



SUPPORT VECTOR MACHINES

- Find hyperplane that maximizes the *margin* between the positive and negative examples



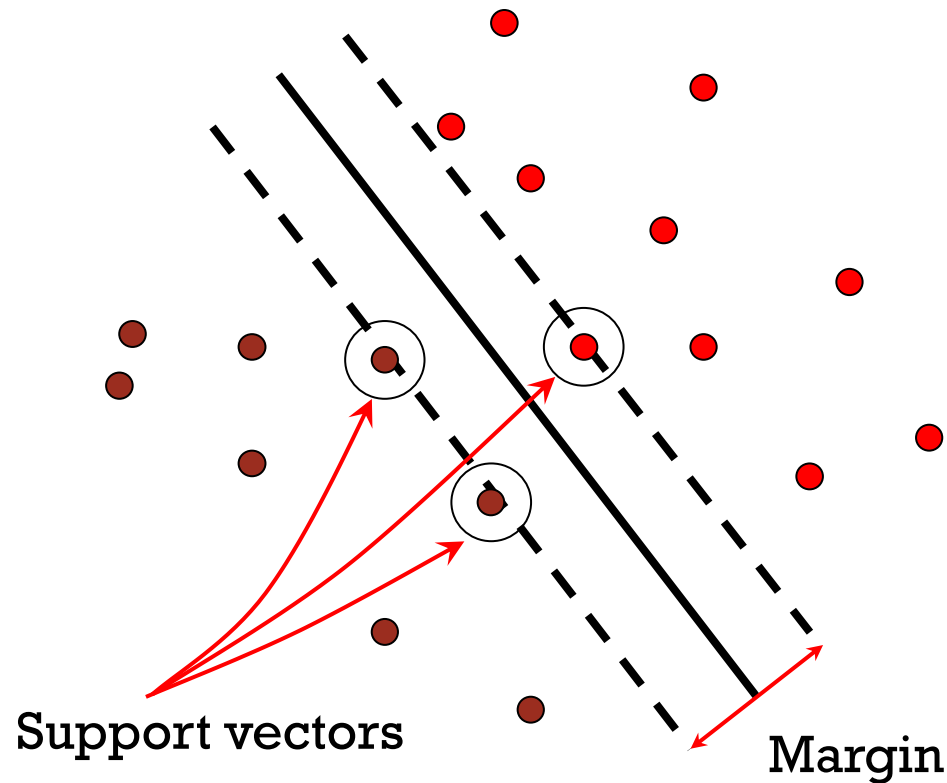
$$\mathbf{x}_i \text{ positive } (y_i = 1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \geq 1$$

$$\mathbf{x}_i \text{ negative } (y_i = -1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \leq -1$$



SUPPORT VECTOR MACHINES

- Find hyperplane that maximizes the *margin* between the positive and negative examples



$$\mathbf{x}_i \text{ positive } (y_i = 1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \geq 1$$

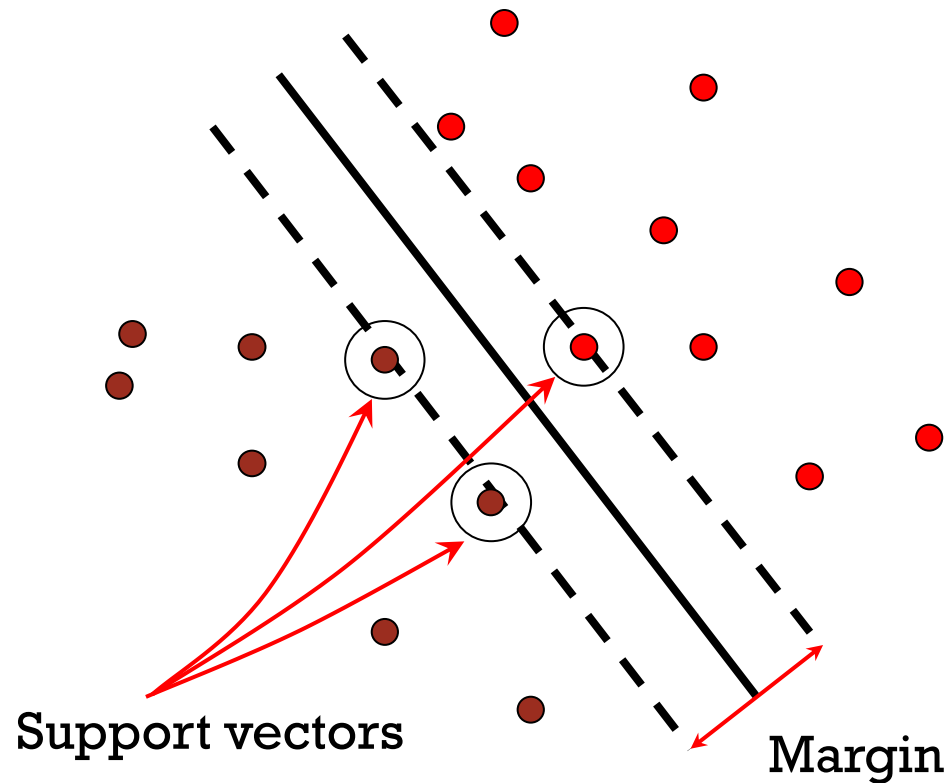
$$\mathbf{x}_i \text{ negative } (y_i = -1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \leq -1$$

$$\text{For support vectors, } \mathbf{x}_i \cdot \mathbf{w} + b = \pm 1$$



SUPPORT VECTOR MACHINES

- Find hyperplane that maximizes the *margin* between the positive and negative examples



$$\mathbf{x}_i \text{ positive } (y_i = 1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \geq 1$$

$$\mathbf{x}_i \text{ negative } (y_i = -1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \leq -1$$

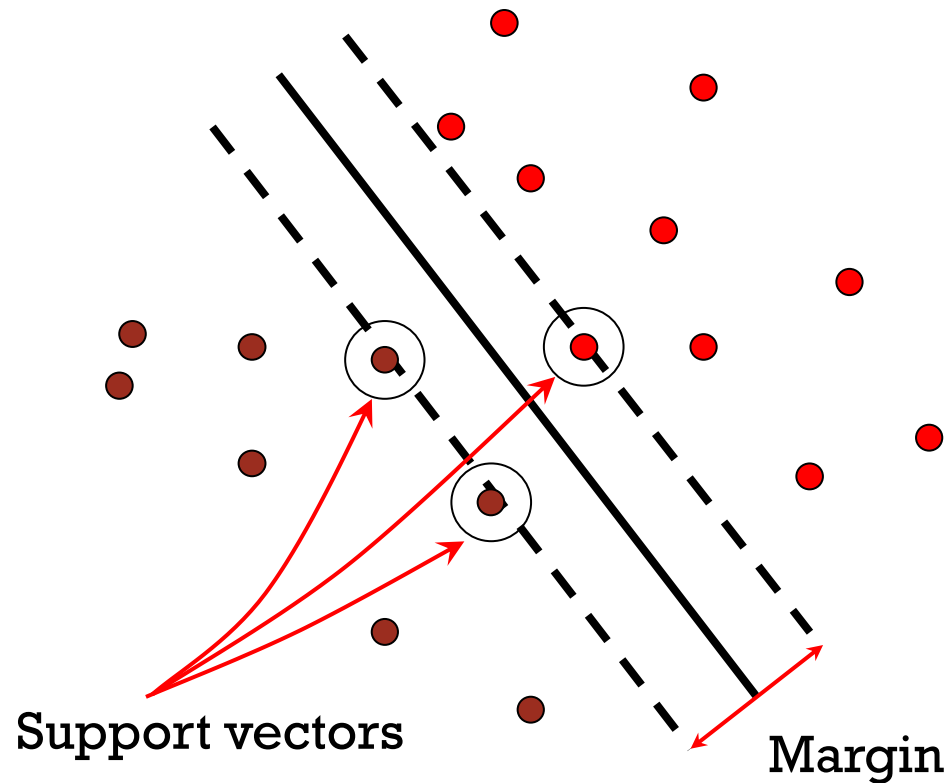
$$\text{For support vectors, } \mathbf{x}_i \cdot \mathbf{w} + b = \pm 1$$

$$\text{Distance between point and hyperplane: } \frac{|\mathbf{x}_i \cdot \mathbf{w} + b|}{\|\mathbf{w}\|}$$



SUPPORT VECTOR MACHINES

- Find hyperplane that maximizes the *margin* between the positive and negative examples



$$\mathbf{x}_i \text{ positive } (y_i = 1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \geq 1$$

$$\mathbf{x}_i \text{ negative } (y_i = -1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \leq -1$$

$$\text{For support vectors, } \mathbf{x}_i \cdot \mathbf{w} + b = \pm 1$$

$$\text{Distance between point and hyperplane: } \frac{|\mathbf{x}_i \cdot \mathbf{w} + b|}{\|\mathbf{w}\|}$$

$$\text{Therefore, the margin is } 2 / \|\mathbf{w}\|$$



FINDING THE MAXIMUM MARGIN HYPERPLANE

1. Maximize margin $2 / \|w\|$



FINDING THE MAXIMUM MARGIN HYPERPLANE

1. Maximize margin $2 / \|\mathbf{w}\|$
2. Correctly classify all training data:

$$\mathbf{x}_i \text{ positive } (y_i = 1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \geq 1$$

$$\mathbf{x}_i \text{ negative } (y_i = -1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \leq -1$$



FINDING THE MAXIMUM MARGIN HYPERPLANE

1. Maximize margin $2 / \|\mathbf{w}\|$
2. Correctly classify all training data:

$$\mathbf{x}_i \text{ positive } (y_i = 1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \geq 1$$

$$\mathbf{x}_i \text{ negative } (y_i = -1): \quad \mathbf{x}_i \cdot \mathbf{w} + b \leq -1$$

- *Quadratic optimization problem:*

$$\min_{\mathbf{w}, b} \frac{1}{2} \|\mathbf{w}\|^2 \quad \text{subject to} \quad y_i (\mathbf{w} \cdot \mathbf{x}_i + b) \geq 1$$



FINDING THE MAXIMUM MARGIN HYPERPLANE

- Solution: $\mathbf{w} = \sum_i \alpha_i y_i \mathbf{x}_i$

learned weight Support vector

- The weights α_i are non-zero only at support vectors.



FINDING THE MAXIMUM MARGIN HYPERPLANE

- Solution: $\mathbf{w} = \sum_i \alpha_i y_i \mathbf{x}_i$
 $b = y_i - \mathbf{w} \cdot \mathbf{x}_i$ (for any support vector)
 $\mathbf{w} \cdot \mathbf{x} + b = \sum_i \alpha_i y_i \mathbf{x}_i \cdot \mathbf{x} + b$



FINDING THE MAXIMUM MARGIN HYPERPLANE

- Solution: $\mathbf{w} = \sum_i \alpha_i y_i \mathbf{x}_i$
 $b = y_i - \mathbf{w} \cdot \mathbf{x}_i$ (for any support vector)

$$\mathbf{w} \cdot \mathbf{x} + b = \sum_i \alpha_i y_i \mathbf{x}_i \cdot \mathbf{x} + b$$

- Classification function:

$$f(x) = \text{sign}(\mathbf{w} \cdot \mathbf{x} + b) \quad \text{If } f(x) < 0, \text{ classify as negative,}$$

$$= \text{sign}\left(\sum_i \alpha_i y_i \mathbf{x}_i \cdot \mathbf{x} + b\right) \quad \text{if } f(x) > 0, \text{ classify as positive}$$

Dot product only!



SVM PARAMETER LEARNING

- Separable data: $\min_{\mathbf{w}, b} \frac{1}{2} \|\mathbf{w}\|^2$ subject to $y_i(\mathbf{w} \cdot \mathbf{x}_i + b) \geq 1$

Maximize
margin

Classify training data correctly



SVM PARAMETER LEARNING

- Separable data: $\min_{\mathbf{w}, b} \frac{1}{2} \|\mathbf{w}\|^2$ subject to $y_i(\mathbf{w} \cdot \mathbf{x}_i + b) \geq 1$

Maximize
margin

Classify training data correctly

- Non-separable data:

$$\min_{\mathbf{w}, b} \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^n \max(0, 1 - y_i(\mathbf{w} \cdot \mathbf{x}_i + b))$$

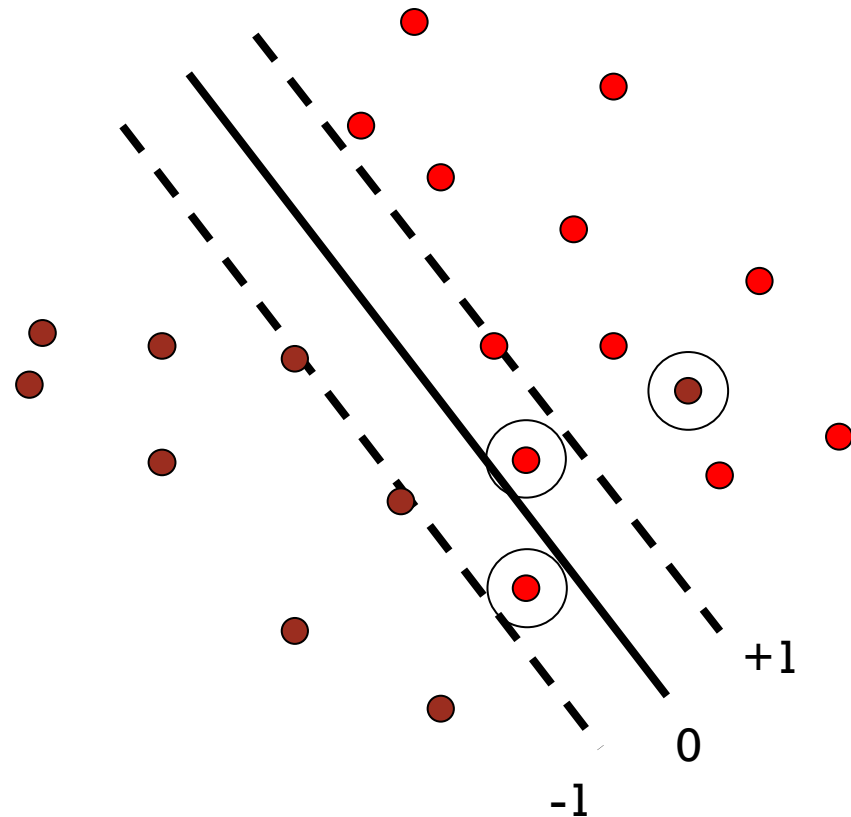
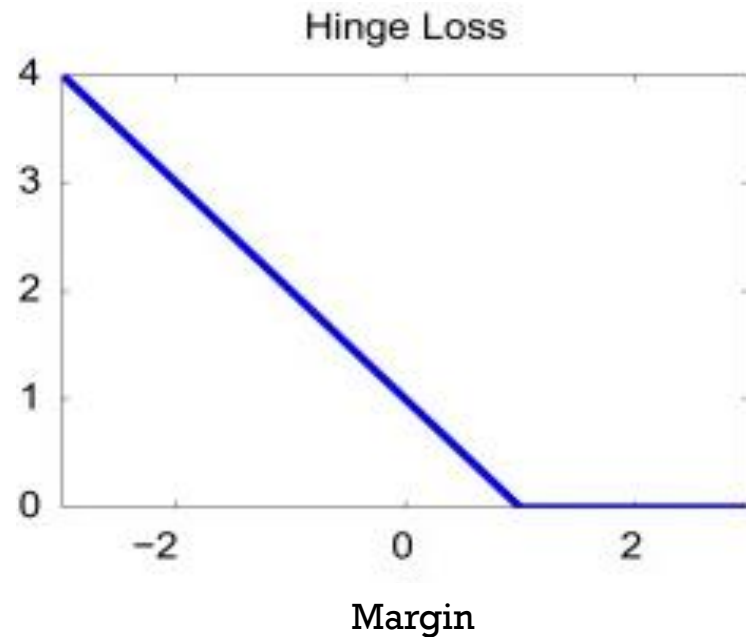
Maximize
margin

Minimize classification mistakes



SVM PARAMETER LEARNING

$$\min_{\mathbf{w}, b} \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^n \max(0, 1 - y_i (\mathbf{w} \cdot \mathbf{x}_i + b))$$

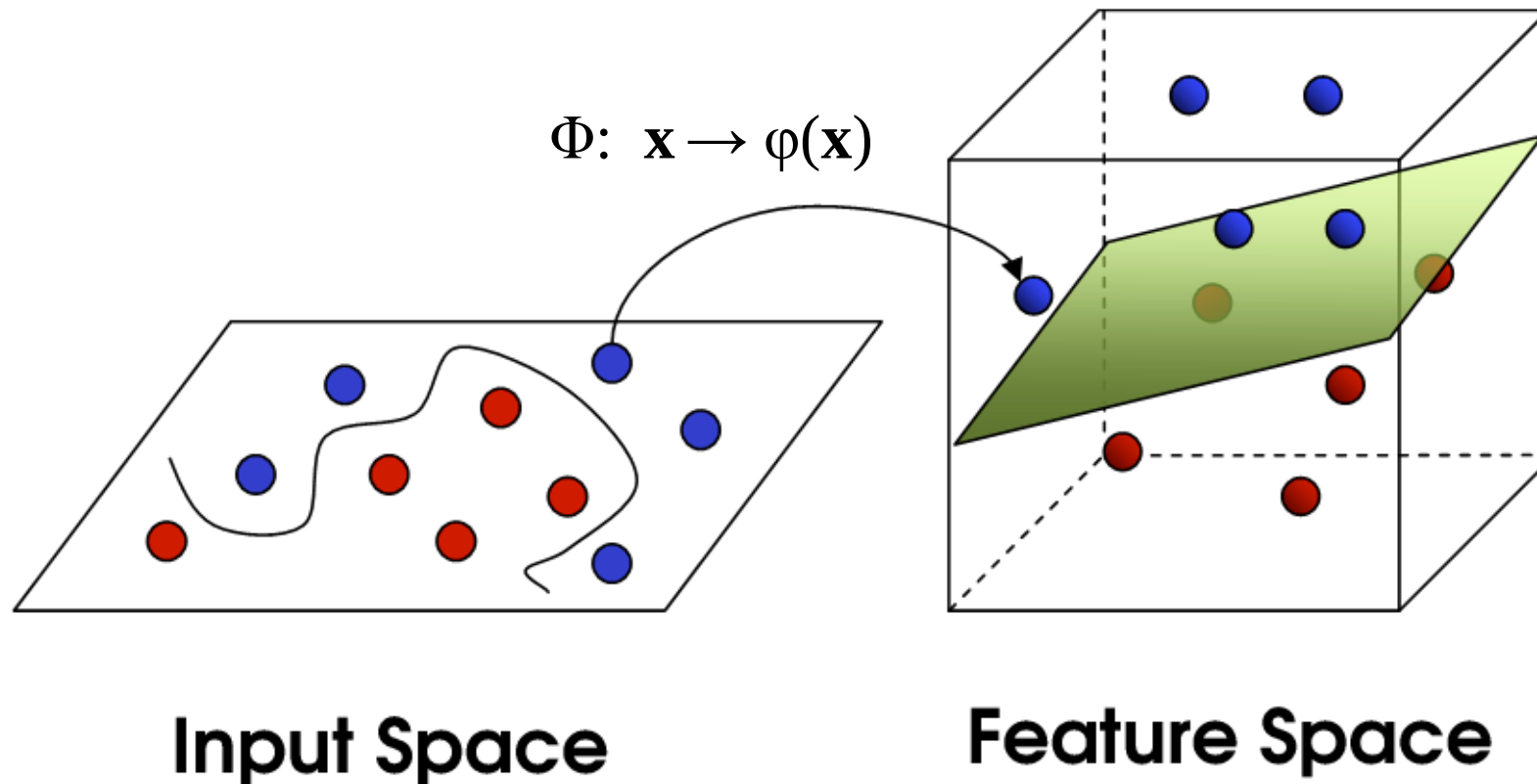


- Demo: <http://cs.stanford.edu/people/karpathy/svmjs/demo>



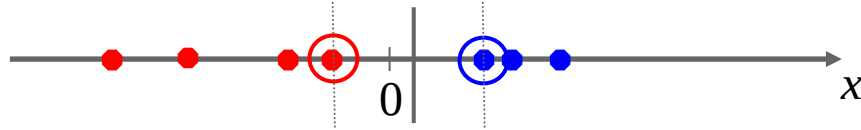
NONLINEAR SVMs

- General idea: the original input space can always be mapped to some higher-dimensional feature space where the training set is separable



NONLINEAR SVMs

- Linearly separable dataset in 1D:

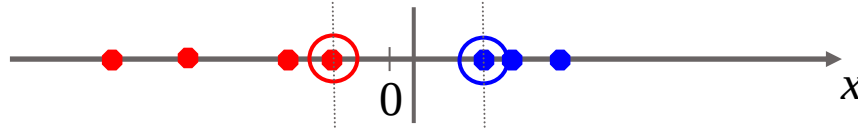


- Non-linearly separable dataset in 1D:

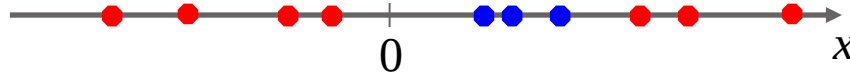


NONLINEAR SVMs

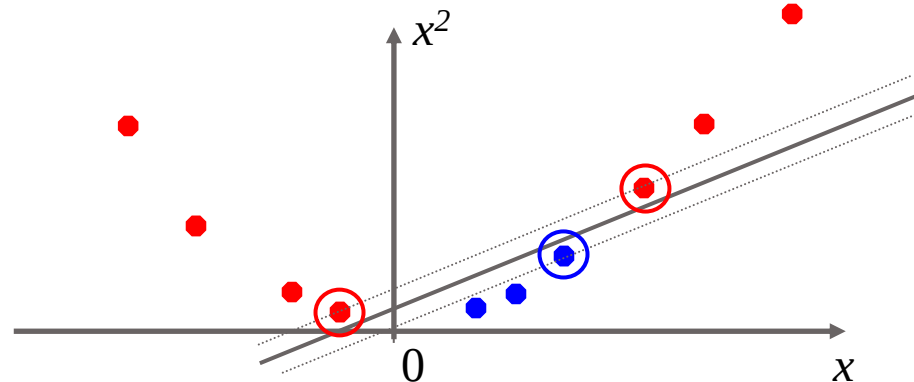
- Linearly separable dataset in 1D:



- Non-linearly separable dataset in 1D:



- We can map the data to a *higher-dimensional space*:



THE KERNEL TRICK

- General idea: the original input space can always be mapped to some higher-dimensional feature space where the training set is separable
- **The kernel trick:** instead of explicitly computing the lifting transformation $\phi(\mathbf{x})$, define a kernel function K such that

$$K(\mathbf{x}, \mathbf{y}) = \phi(\mathbf{x}) \cdot \phi(\mathbf{y})$$



THE KERNEL TRICK

- Linear SVM decision function:

$$\mathbf{w} \cdot \mathbf{x} + b = \sum_i \alpha_i y_i \mathbf{x}_i \cdot \mathbf{x} + b$$

learned
weight

Support
vector



THE KERNEL TRICK

- Linear SVM decision function:

$$\mathbf{w} \cdot \mathbf{x} + b = \sum_i \alpha_i y_i \mathbf{x}_i \cdot \mathbf{x} + b$$

- Kernel SVM decision function:

$$\sum_i \alpha_i y_i \varphi(\mathbf{x}_i) \cdot \varphi(\mathbf{x}) + b = \sum_i \alpha_i y_i K(\mathbf{x}_i, \mathbf{x}) + b$$

- This gives a nonlinear decision boundary in the original feature space



EXAMPLE

2-dimensional vectors $\mathbf{x}=[x_1 \ x_2]$;

let $K(\mathbf{x}_i, \mathbf{x}_j) = (1 + \mathbf{x}_i^T \mathbf{x}_j)^2$



EXAMPLE

2-dimensional vectors $\mathbf{x}=[x_1 \ x_2]$;

let $K(\mathbf{x}_i, \mathbf{x}_j) = (1 + \mathbf{x}_i^T \mathbf{x}_j)^2$

Need to show that $K(\mathbf{x}_i, \mathbf{x}_j) = \phi(\mathbf{x}_i)^T \phi(\mathbf{x}_j)$:



EXAMPLE

2-dimensional vectors $\mathbf{x}=[x_1 \ x_2]$;

let $K(\mathbf{x}_i, \mathbf{x}_j) = (1 + \mathbf{x}_i^T \mathbf{x}_j)^2$

Need to show that $K(\mathbf{x}_i, \mathbf{x}_j) = \phi(\mathbf{x}_i)^T \phi(\mathbf{x}_j)$:

$$\begin{aligned} K(\mathbf{x}_i, \mathbf{x}_j) &= (1 + \mathbf{x}_i^T \mathbf{x}_j)^2, \\ &= 1 + x_{i1}^2 x_{j1}^2 + 2 x_{i1} x_{j1} x_{i2} x_{j2} + x_{i2}^2 x_{j2}^2 + 2 x_{i1} x_{j1} + 2 x_{i2} x_{j2} \end{aligned}$$



EXAMPLE

2-dimensional vectors $\mathbf{x}=[x_1 \ x_2]$;

let $K(\mathbf{x}_i, \mathbf{x}_j) = (1 + \mathbf{x}_i^T \mathbf{x}_j)^2$

Need to show that $K(\mathbf{x}_i, \mathbf{x}_j) = \boldsymbol{\phi}(\mathbf{x}_i)^T \boldsymbol{\phi}(\mathbf{x}_j)$:

$$\begin{aligned} K(\mathbf{x}_i, \mathbf{x}_j) &= (1 + \mathbf{x}_i^T \mathbf{x}_j)^2, \\ &= 1 + x_{i1}^2 x_{j1}^2 + 2 x_{i1} x_{j1} x_{i2} x_{j2} + x_{i2}^2 x_{j2}^2 + 2 x_{i1} x_{j1} + 2 x_{i2} x_{j2} \\ &= [1 \ x_{i1}^2 \ \sqrt{2} x_{i1} x_{i2} \ x_{i2}^2 \ \sqrt{2} x_{i1} \ \sqrt{2} x_{i2}]^T \\ &\quad [1 \ x_{j1}^2 \ \sqrt{2} x_{j1} x_{j2} \ x_{j2}^2 \ \sqrt{2} x_{j1} \ \sqrt{2} x_{j2}] \end{aligned}$$



EXAMPLE

2-dimensional vectors $\mathbf{x}=[x_1 \ x_2]$;

let $K(\mathbf{x}_i, \mathbf{x}_j) = (1 + \mathbf{x}_i^T \mathbf{x}_j)^2$

Need to show that $K(\mathbf{x}_i, \mathbf{x}_j) = \boldsymbol{\varphi}(\mathbf{x}_i)^T \boldsymbol{\varphi}(\mathbf{x}_j)$:

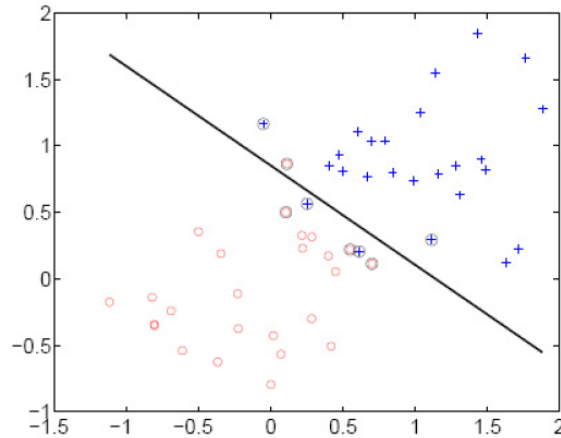
$$\begin{aligned} K(\mathbf{x}_i, \mathbf{x}_j) &= (1 + \mathbf{x}_i^T \mathbf{x}_j)^2, \\ &= 1 + x_{i1}^2 x_{j1}^2 + 2 x_{i1} x_{j1} x_{i2} x_{j2} + x_{i2}^2 x_{j2}^2 + 2 x_{i1} x_{j1} + 2 x_{i2} x_{j2} \\ &= [1 \ x_{i1}^2 \ \sqrt{2} x_{i1} x_{i2} \ x_{i2}^2 \ \sqrt{2} x_{i1} \ \sqrt{2} x_{i2}]^T \\ &\quad [1 \ x_{j1}^2 \ \sqrt{2} x_{j1} x_{j2} \ x_{j2}^2 \ \sqrt{2} x_{j1} \ \sqrt{2} x_{j2}] \\ &= \boldsymbol{\varphi}(\mathbf{x}_i)^T \boldsymbol{\varphi}(\mathbf{x}_j), \end{aligned}$$

where $\boldsymbol{\varphi}(\mathbf{x}) = [1 \ x_1^2 \ \sqrt{2} x_1 x_2 \ x_2^2 \ \sqrt{2} x_1 \ \sqrt{2} x_2]$

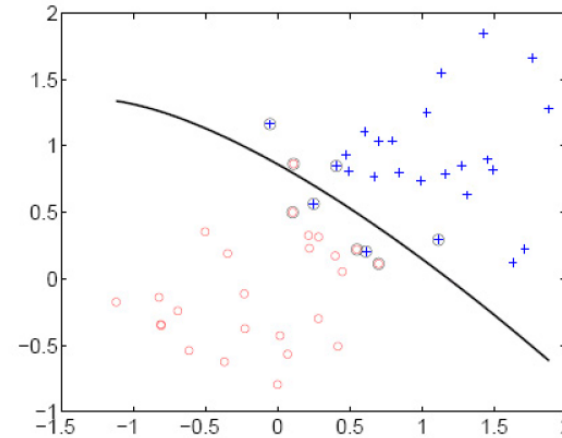


POLYNOMIAL KERNEL

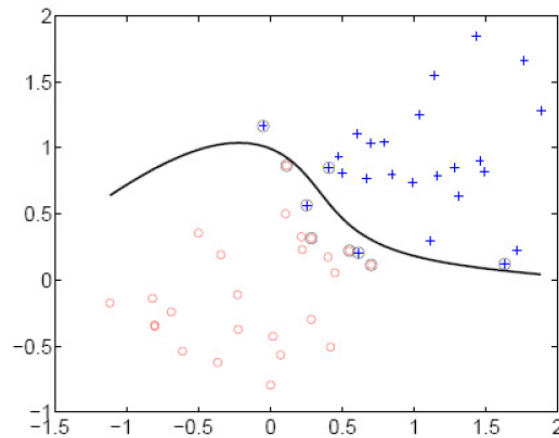
$$K(\mathbf{x}, \mathbf{y}) = (c + \mathbf{x} \cdot \mathbf{y})^d$$



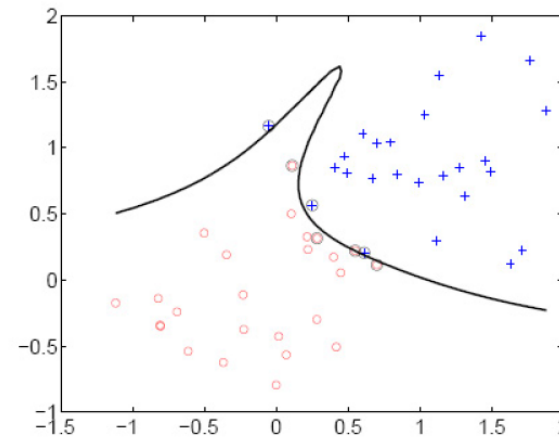
linear



2nd order polynomial



4th order polynomial



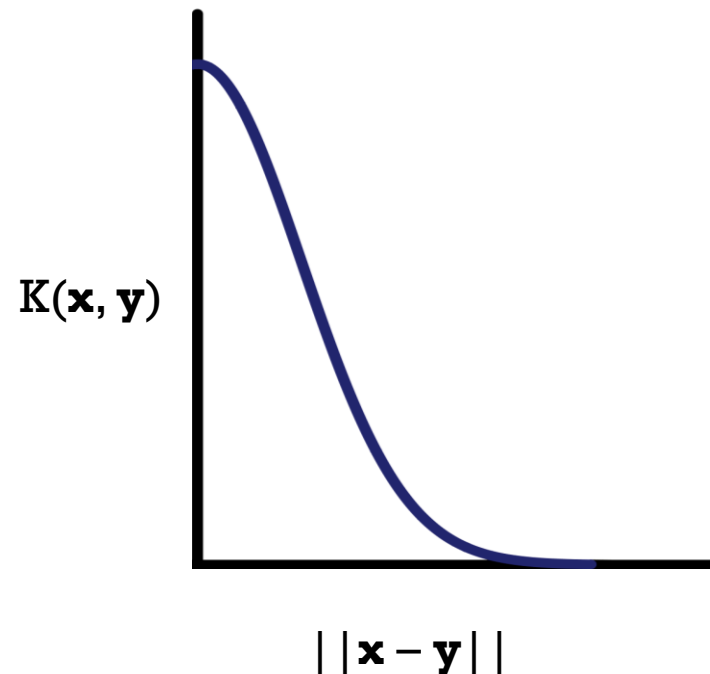
8th order polynomial



GAUSSIAN KERNEL

- Also known as the radial basis function (RBF) kernel:

$$K(\mathbf{x}, \mathbf{y}) = \exp\left(-\frac{1}{\sigma^2} \|\mathbf{x} - \mathbf{y}\|^2\right)$$



KERNELS FOR HISTOGRAMS

- Histogram intersection:

$$K(h_1, h_2) = \sum_{i=1}^N \min(h_1(i), h_2(i))$$

- Square root (Bhattacharyya kernel):

$$K(h_1, h_2) = \sum_{i=1}^N \sqrt{h_1(i)h_2(i)}$$



SVMS: PROS AND CONS

- Pros

- Kernel-based framework is **very powerful**, flexible
- Training is convex optimization, **globally optimal solution** can be found
- SVMs work very well in practice, even with **very small training** sample sizes

- Cons

- No “direct” **multi-class SVM**, must combine two-class SVMs (e.g., with one-vs-others)
- **Computation, memory** (esp. for nonlinear SVMs)



MULTICLASS SUPPORT VECTOR MACHINE LOSS

- i^{th} example: image x_i and the label y_i
- Score for the j^{th} class: $s_j = f(x_i, W)_j$



MULTICLASS SUPPORT VECTOR MACHINE LOSS

- i^{th} example: image x_i and the label y_i
- Score for the j^{th} class: $s_j = f(x_i, W)_j$
- The Multiclass SVM loss for the i^{th} example is then formalized as follows:

$$L_i = \sum_{j \neq y_i} \max(0, s_j - s_{y_i} + \Delta)$$

Hinge Loss

Margin



MULTICLASS SUPPORT VECTOR MACHINE LOSS

- i^{th} example: image x_i and the label y_i
- Score for the j^{th} class: $s_j = f(x_i, W)_j$
- The Multiclass SVM loss for the i^{th} example is then formalized as follows:

$$L_i = \sum_{j \neq y_i} \max(0, s_j - s_{y_i} + \Delta)$$

Hinge Loss

Margin

Problem: W is not necessarily unique



MULTICLASS SUPPORT VECTOR MACHINE LOSS

- Regularization Penalty:

$$R(W) = \sum_k \sum_l W_{k,l}^2$$



MULTICLASS SUPPORT VECTOR MACHINE LOSS

- Regularization Penalty:

$$R(W) = \sum_k \sum_l W_{k,l}^2$$

- Hinge Loss:

$$L = \underbrace{\frac{1}{N} \sum_i L_i}_{\text{data loss}} + \underbrace{\lambda R(W)}_{\text{regularization loss}}$$



MULTICLASS SUPPORT VECTOR MACHINE LOSS

- Regularization Penalty:

$$R(W) = \sum_k \sum_l W_{k,l}^2$$

- Hinge Loss:

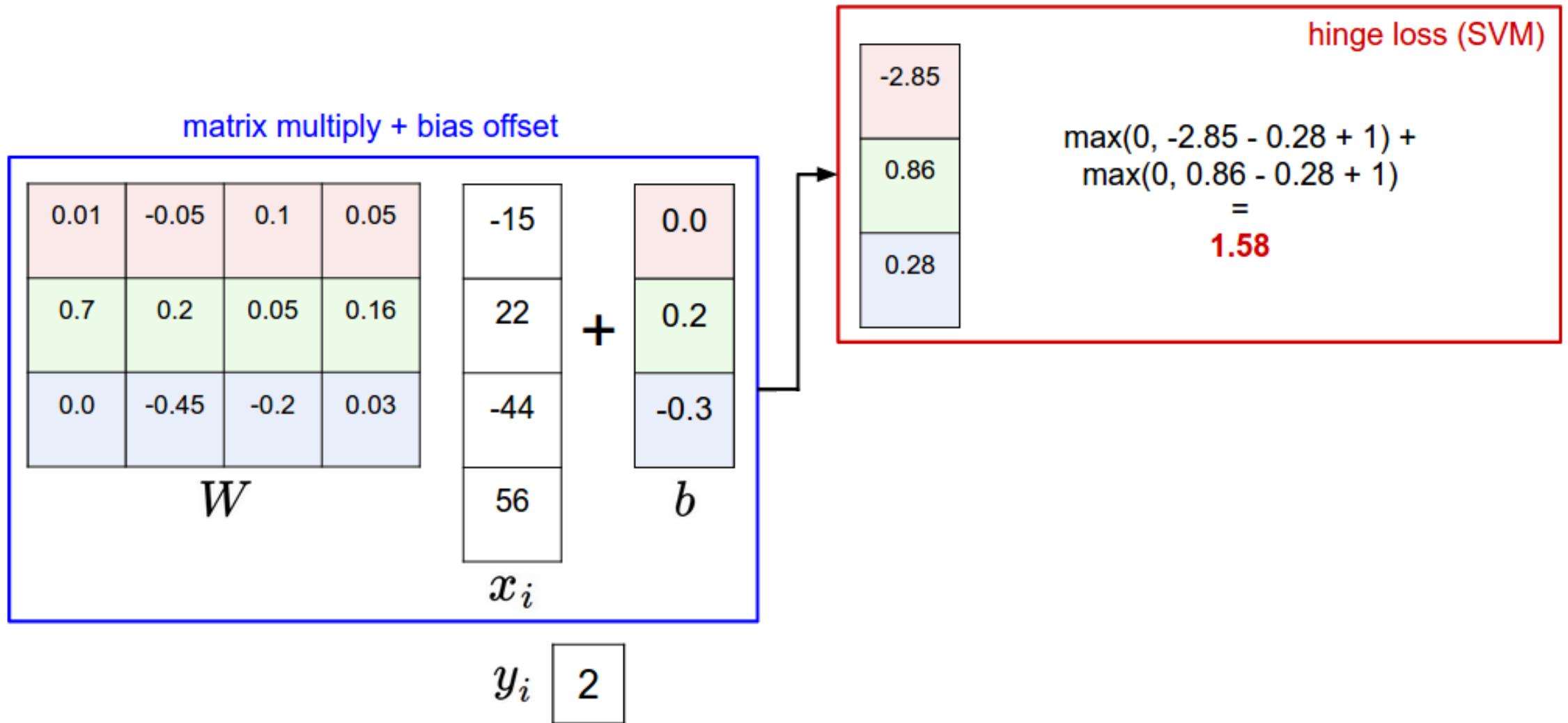
$$L = \underbrace{\frac{1}{N} \sum_i L_i}_{\text{data loss}} + \underbrace{\lambda R(W)}_{\text{regularization loss}}$$



$$L = \frac{1}{N} \sum_i \sum_{j \neq y_i} [\max(0, f(x_i; W)_j - f(x_i; W)_{y_i} + \Delta)] + \lambda \sum_k \sum_l W_{k,l}^2$$



HINGE LOSS



SOFTMAX CLASSIFIER

- Interprets the class scores as the unnormalized log probabilities for each class and replace the *hinge loss* with a **cross-entropy loss** that has the form:

$$L_i = -\log \left(\frac{e^{s_{y_i}}}{\sum_j e^{s_j}} \right) = -s_{y_i} + \log \sum_j e^{s_j}$$



SOFTMAX CLASSIFIER

- Interprets the class scores as the unnormalized log probabilities for each class and replace the *hinge loss* with a **cross-entropy loss** that has the form:

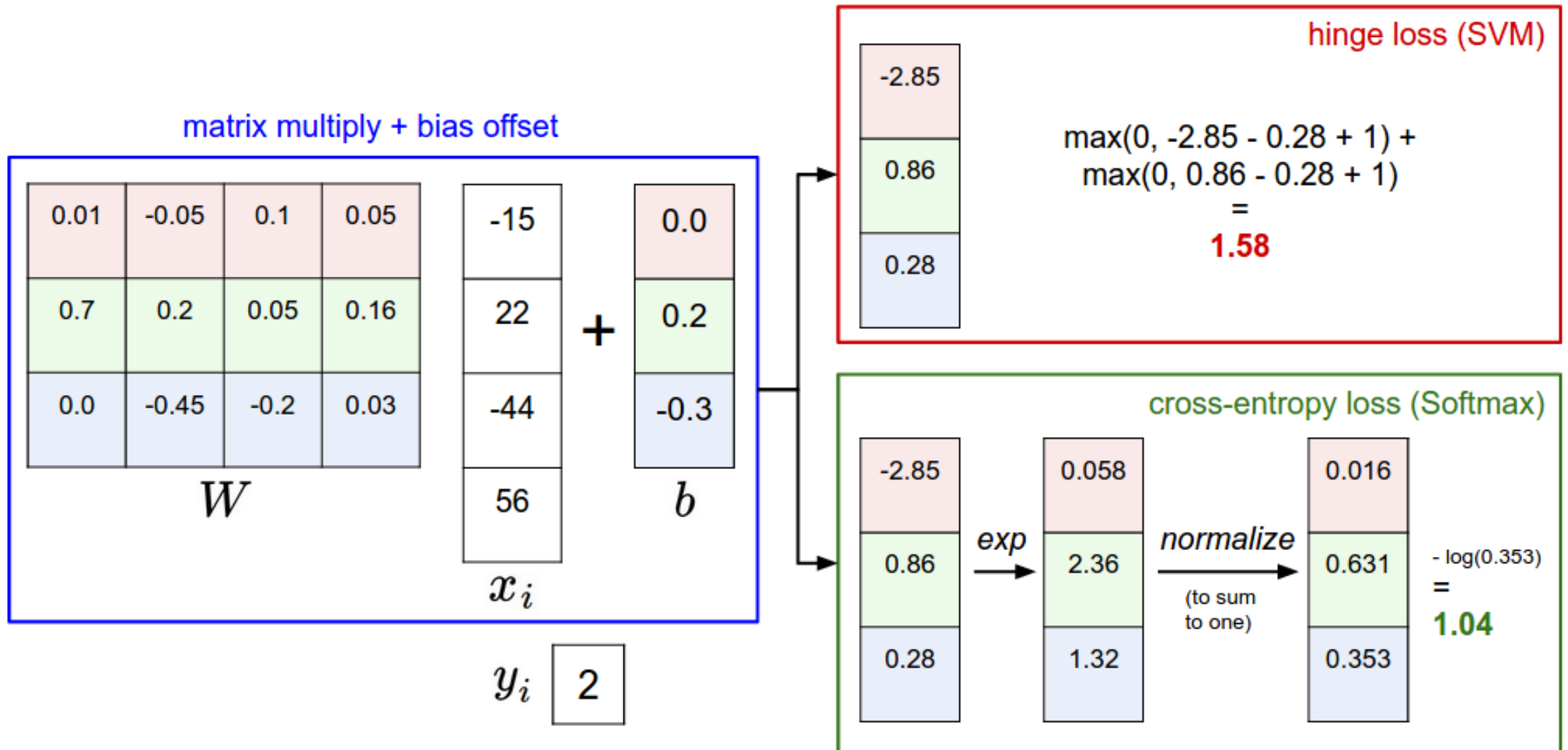
$$L_i = -\log \left(\frac{e^{s_{y_i}}}{\sum_j e^{s_j}} \right) = -s_{y_i} + \log \sum_j e^{s_j}$$

- Softmax Loss:

$$L = \underbrace{\frac{1}{N} \sum_i L_i}_{\text{data loss}} + \underbrace{\lambda R(W)}_{\text{regularization loss}}$$



HINGE VS CROSS-ENTROPY LOSS



ACKNOWLEDGEMENT

Thanks to the following courses and corresponding researchers for making their teaching/research material online

- Convolutional Neural Networks for Visual Recognition, Stanford University
- Deep Learning, Stanford University
- Introduction to Deep Learning, University of Illinois at Urbana-Champaign
- Introduction to Deep Learning, Carnegie Mellon University
- Natural Language Processing with Deep Learning, Stanford University
- And Many More Publicly Available Resources



Questions?

