

Introduction to Machine Learning: Discriminative Models

Farzad Kamalabadi

University of Illinois Urbana-Champaign

Heliophysics Summer School, 2026

Outline of this lecture

- Introduction to pattern recognition and classification
- Linear regression
- Logistic regression
- Support Vector Machines
- Multiclass regression and SVM
- Deep Learning

Reading Material

- K. Murphy; Machine Learning: A Probabilistic Perspective

History (1930s to 1950s):

- Neurology research showed brain as electrical network of neurons
- Cybernetics (Norbert Wiener) described control and stability in electrical networks
- Information theory (Claude Shannon) described digital signals
- Theory of computation (Alan Turing) described digital nature of computation
- Electronic brain/neural net (Walter Pitts and Warren McCulloch)

Marvin Minsky (1955):

A Proposal for the

DARTMOUTH SUMMER RESEARCH PROJECT ON ARTIFICIAL INTELLIGENCE

June 17 - Aug. 16

We propose that a 2 month, 10 man study of artificial intelligence be carried out during the summer of 1956 at Dartmouth College in Hanover, New Hampshire. The study is to proceed on the basis of the conjecture that every aspect of learning or any other feature of intelligence can in principle be so precisely described that a machine can be made to simulate it. An attempt will be made to find how to make machines use language, form abstractions and concepts, solve kinds of problems now reserved for humans, and improve themselves. We think that a significant advance can be made in one or more of these problems if a carefully selected group of scientists work on it together for a summer.

Mentioned aspects:

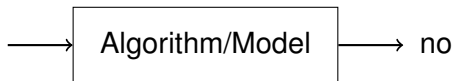
- Automatic computers
- Programming a computer with language
- Neuron nets
- Theory on complexity of calculation
- Self-improvement
- Abstractions
- Randomness and Creativity

Where pattern recognition/AI shows up ‘recently’:

- Chess play (Deep Blue, IBM, 1996, 1997)
- Jeopardy (Watson, IBM, 2011)
- Playing Atari games (DQN, Deepmind, 2015)
- Game of Go (AlphaGo, Deepmind, 2016)
- Autonomous Driving
- Generative Pre-trained Transformer (GPT), OpenAI, 2022; GPT5, 2025)
- DeepMind’s latest weather prediction model, GenCast, outperforms traditional physics-based models and even other AI models, in terms of accuracy and speed.

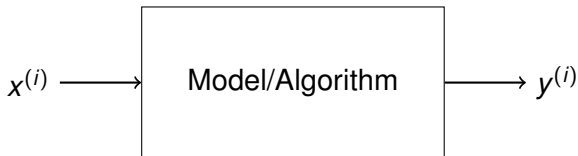
Example:

Let's get a computer to recognize whether there is a cat in the image.



Formulation of pattern recognition for this lecture (discriminative models):

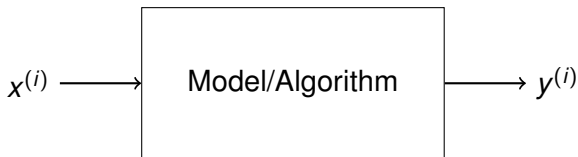
- input data/value/vector: $x^{(i)}$
- label/output: $y^{(i)}$



What do we call this process?

- Inference
- Prediction

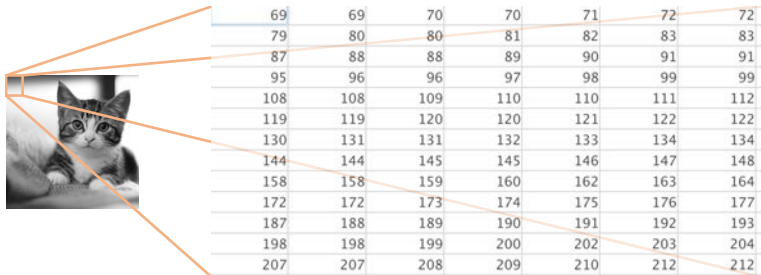
What is learning?



- Model/Algorithm depends on parameters $\mathbf{w}$
- Learning/Fitting of parameters $\mathbf{w}$
- Based on dataset $\mathcal{D} = \{(x^{(i)}, y^{(i)})\}_{i=1}^N$

How could we recognize a cat? What should the algorithm do?

How is an image represented in a computer?



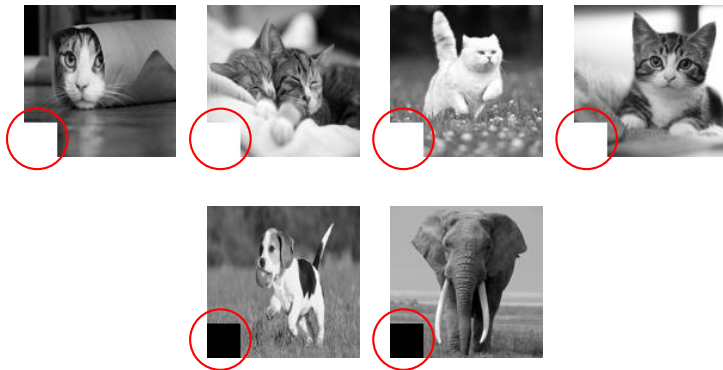
Given that we know how a computer ‘sees’ an image, how can it recognize cats?



Still very hard to describe what a cat looks like.

Let's look at tons of examples (Dataset).

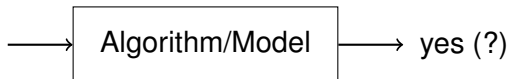
Dataset (Thousands of examples):



Instead of asking to recognize cats in general, let's recognize cats in this dataset

Algorithm: if bottom left corner is black (0) say no, otherwise say yes

Algorithm: if bottom left corner is black (0) say no, otherwise say yes



Works perfect on our dataset. :)
But does not generalize to other data. :(

Conclusion:

We designed a simple “classifier” that works on this dataset but doesn’t work on real data.

To our rescue:

Machine learning has found mechanisms to search for mappings which **generalize**.

Scope of this lecture:

In this lecture we talk about algorithms and models. A detailed treatment of **generalization** is beyond our scope.

Categorization of machine learning/pattern recognition algorithms according to

- Available annotated data (supervised vs. unsupervised)
- Complexity of model (linear vs. non-linear)
- Structure of output (independent vs. structured)
- Modeling of data ($x^{(i)}$) or label ($y^{(i)}$) (generative vs. discriminative)

Classification Framework: Formalism

Main Components

- an *input* (also called *observation* or *data point*), denoted by x
- a **discrete** *output* (also called *label* or *class*), denoted by y
- a *decision function* (also called *classification rule*), $y = f(x)$

The main problem in Pattern Recognition/Machine Learning is to construct the *decision function*.

Linear Regression

Notation:

- Vector: $\mathbf{x} = \begin{bmatrix} x_1 \\ \vdots \\ x_n \end{bmatrix} \in \mathbb{R}^n$
- Matrix: $\mathbf{X} = \begin{bmatrix} x_{1,1} & \cdots & x_{1,m} \\ \vdots & \ddots & \vdots \\ x_{n,1} & \cdots & x_{n,m} \end{bmatrix} \in \mathbb{R}^{n \times m}$
- Norm: $\|x^{(1)} - x^{(2)}\|_2^2 = \sum_{i=1}^n (x_i^{(1)} - x_i^{(2)})^2$ distance between two points in n dimensions
- Transpose: $\mathbf{X}^T = \begin{bmatrix} x_{1,1} & \cdots & x_{n,1} \\ \vdots & \ddots & \vdots \\ x_{1,m} & \cdots & x_{m,n} \end{bmatrix} \in \mathbb{R}^{m \times n}$
 $\mathbf{x}^T = \begin{bmatrix} x_1 & \cdots & x_n \end{bmatrix} \in \mathbb{R}^{1 \times n}$
- Matrix multiplication: $\mathbf{X}^T \mathbf{x}$ or $\mathbf{X} \mathbf{x}$?

Discrete Probability: $y \in \{1, \dots, 6\}$

- Discrete probability distribution: $p(Y = y) \in [0, 1]$ with $\sum_{y \in \{1, \dots, 6\}} p(Y = y) = 1$
- Abbreviation: $p(Y = y) = p(y) \in [0, 1]$
- Expectation: $\mathbb{E}_{p(y)}[f(y)] = \sum_{y \in \{1, \dots, 6\}} p(y)f(y)$

Continuous probability: $y \in \mathbb{R}$

- $p(Y = 1) = 0$
- Probability density function: $p(y) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left(\frac{-1}{2\sigma^2}(y - \mu)^2\right)$
- Mean: $\mathbb{E}_{p(y)}[y] = \int_{-\infty}^{\infty} yp(y)dy = \mu$
- Variance: $\mathbb{E}_{p(y)}[(y - \mu)^2] = \sigma^2$

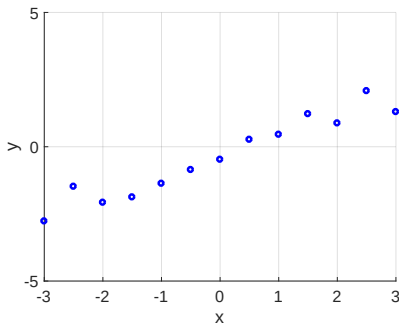
Multivariate continuous probability: $\mathbf{y} \in \mathbb{R}^n$ $\mu \in \mathbb{R}^n$

- n-dimensional density:
$$p(\mathbf{y}) = p(y_1, \dots, y_n) = \frac{1}{\sqrt{(2\pi)^n |\Sigma|}} \exp\left(\frac{-1}{2}(\mathbf{y} - \mu)^T \Sigma^{-1}(\mathbf{y} - \mu)\right)$$
- Covariance matrix: Σ

Multivariate calculus: $\mathbf{x} \in \mathbb{R}^n$, $\mathbf{w} \in \mathbb{R}^n$, $\mathbf{A} \in \mathbb{R}^{n \times n}$

- Multivariate function: $f(\mathbf{x}) = \mathbf{w}^T \mathbf{x}$
- Derivative: $\frac{\partial f}{\partial \mathbf{x}} = \mathbf{w}$ (e.g., Eq. (69))
- Multivariate function: $f(\mathbf{x}) = \mathbf{x}^T \mathbf{A} \mathbf{x}$
- Derivative: $\frac{\partial f}{\partial \mathbf{x}} = (\mathbf{A} + \mathbf{A}^T) \mathbf{x}$ (e.g., Eq. (78) or Eq. (81))

Linear Regression - The Problem:



Given outcomes $y^{(i)} \in \mathbb{R}$ for covariates $x^{(i)} \in \mathbb{R}$,

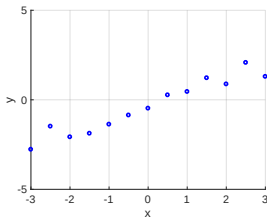
what is/are the underlying system/model/model parameters?

Let's assume a linear model with parameters $w_1 \in \mathbb{R}$ and $w_2 \in \mathbb{R}$

$$y = w_1 \cdot x + w_2$$

Given a dataset of N pairs (x, y) :

$$\mathcal{D} = \{(x^{(i)}, y^{(i)})\}_{i=1}^N$$



How do we find the parameters w_1, w_2 ?

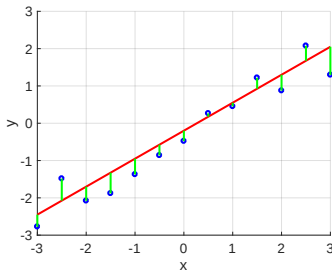
Assuming model

$$y = w_1 \cdot x + w_2$$

Find parameters w_1 , w_2 such that the squared error is small

$$\arg \min_{w_1, w_2} \frac{1}{2} \sum_{i=1}^N \left(y^{(i)} - w_1 \cdot x^{(i)} - w_2 \right)^2$$

What exactly is the error?



Program:

$$\arg \min_{w_1, w_2} \frac{1}{2} \sum_{i=1}^N \left(y^{(i)} - w_1 \cdot x^{(i)} - w_2 \right)^2$$

Vector notation:

$$\arg \min_{w_1, w_2} \frac{1}{2} \left\| \underbrace{\begin{bmatrix} y^{(1)} \\ \vdots \\ y^{(N)} \end{bmatrix}}_{\mathbf{Y} \in \mathbb{R}^N} - \underbrace{\begin{bmatrix} x^{(1)} & 1 \\ \vdots & \vdots \\ x^{(N)} & 1 \end{bmatrix}}_{\mathbf{X}^T \in \mathbb{R}^{N \times 2}} \cdot \underbrace{\begin{bmatrix} w_1 \\ w_2 \end{bmatrix}}_{\mathbf{w} \in \mathbb{R}^2} \right\|_2^2$$

Program:

$$\arg \min_{\mathbf{w}} \underbrace{\frac{1}{2} \|\mathbf{Y} - \mathbf{X}^\top \mathbf{w}\|_2^2}_{\text{cost/loss function}}$$

How to solve the program:

- Take derivative w.r.t. $\mathbf{w}$ of cost function
- Set derivative w.r.t. $\mathbf{w}$ to zero
- Solve for $\mathbf{w}$

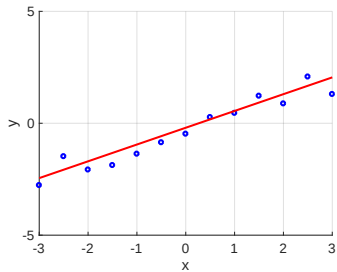
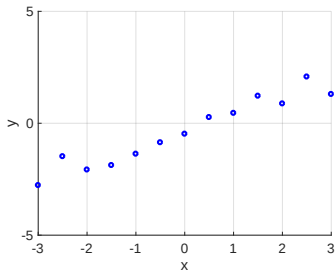
Derivative:

$$\mathbf{X}\mathbf{X}^\top \mathbf{w}^* - \mathbf{X}\mathbf{Y} = 0$$

Solution:

$$\mathbf{w}^* = (\mathbf{X}\mathbf{X}^\top)^{-1} \mathbf{X}\mathbf{Y}$$

Linear regression:



Extensions:

- Higher dimensional problems ($\mathbf{x}^{(i)} \in \mathbb{R}^d$)
- Regularization
- Higher order polynomials

Higher dimensional problems ($\mathbf{x}^{(i)} \in \mathbb{R}^d, y^{(i)} \in \mathbb{R}$)

Model:

$$y^{(i)} = w_0 + \sum_{k=1}^d \mathbf{x}_k^{(i)} w_k$$

Program:

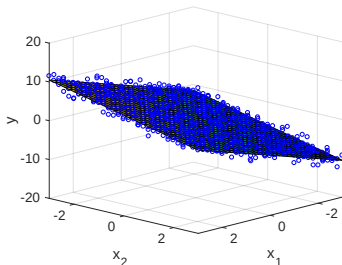
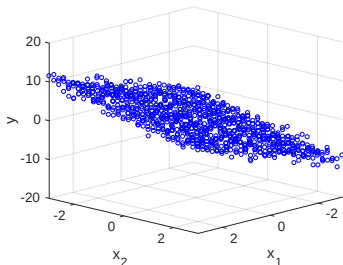
$$\arg \min_{\mathbf{w}} \frac{1}{2} \left\| \underbrace{\mathbf{Y}}_{\in \mathbb{R}^N} - \underbrace{\mathbf{X}^\top}_{\in \mathbb{R}^{N \times (d+1)}} \underbrace{\mathbf{w}}_{\in \mathbb{R}^{d+1}} \right\|_2^2$$

Solution: (obviously the same as before)

$$\mathbf{w}^* = \left(\mathbf{X} \mathbf{X}^\top \right)^{-1} \mathbf{X} \mathbf{Y}$$

Example:

$$\mathbf{w}^* = (\mathbf{X}\mathbf{X}^\top)^{-1} \mathbf{X}\mathbf{Y}$$



What if $N < d + 1$?

Regularization:

we want to make sure that the parameters are not too large

we want to make sure we can invert the matrix

Program:

$$\arg \min_{\mathbf{w}} \underbrace{\frac{1}{2} \|\mathbf{Y} - \mathbf{X}^\top \mathbf{w}\|_2^2 + \frac{C}{2} \|\mathbf{w}\|_2^2}_{\text{cost function}}$$

Solution:

$$\mathbf{w}^* = \left(\mathbf{X}\mathbf{X}^\top + C\mathbf{I} \right)^{-1} \mathbf{X}\mathbf{Y}$$

Higher order polynomials ($x^{(i)} \in \mathbb{R}, y^{(i)} \in \mathbb{R}$)

Model:

$$y^{(i)} = w_2 \cdot (x^{(i)})^2 + w_1 \cdot x^{(i)} + w_0$$

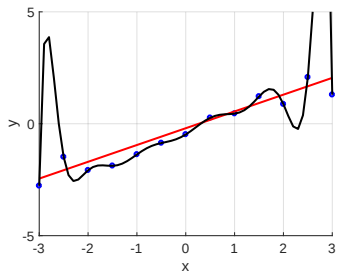
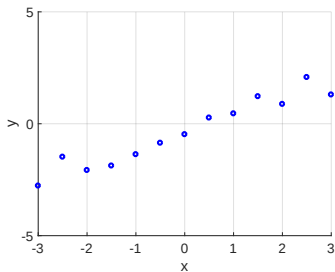
Program:

$$\arg \min_{w_0, w_1, w_2} \frac{1}{2} \left\| \underbrace{\begin{bmatrix} y^{(1)} \\ \vdots \\ y^{(N)} \end{bmatrix}}_{\mathbf{Y} \in \mathbb{R}^N} - \underbrace{\begin{bmatrix} (x^{(1)})^2 & x^{(1)} & 1 \\ \vdots & \vdots & \vdots \\ (x^{(N)})^2 & x^{(N)} & 1 \end{bmatrix}}_{\Phi^T \in \mathbb{R}^{N \times M}} \cdot \underbrace{\begin{bmatrix} w_2 \\ w_1 \\ w_0 \end{bmatrix}}_{\mathbf{w} \in \mathbb{R}^M} \right\|_2^2$$

Solution:

$$\mathbf{w}^* = (\Phi \Phi^T)^{-1} \Phi \mathbf{Y}$$

Example:



Which model is more reasonable?

Generalizing all aforementioned cases:

- $\mathbf{x}^{(i)}$ is some data (e.g., images)
- $\phi(\mathbf{x}^{(i)}) \in \mathbb{R}^M$ is a transformation into a feature vector

Model:

$$y^{(i)} = \phi(\mathbf{x}^{(i)})^\top \mathbf{w}$$

Program:

$$\arg \min_{\mathbf{w}} \frac{1}{2} \sum_{i=1}^N \left(y^{(i)} - \phi(\mathbf{x}^{(i)})^\top \mathbf{w} \right)^2$$

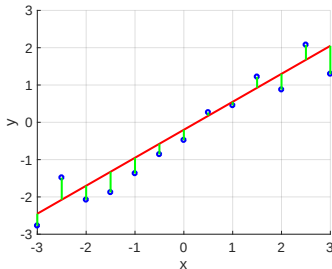
Solution:

$$\mathbf{w}^* = \left(\Phi \Phi^\top \right)^{-1} \Phi \mathbf{Y} \quad \text{where} \quad \Phi = \left[\phi(\mathbf{x}^{(1)}), \dots, \phi(\mathbf{x}^{(N)}) \right] \in \mathbb{R}^{M \times N}$$

Linear regression:

- So far: Error view

$$\left(y^{(i)} - \phi(x^{(i)})^\top \mathbf{w}\right)^2$$

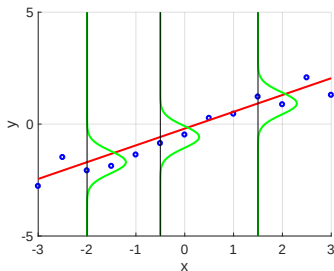


- Alternatively: Probabilistic view

A probabilistic interpretation of linear regression:

Model: Gaussian distribution

$$p(y^{(i)}|x^{(i)}) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{1}{2\sigma^2}(y^{(i)} - \mathbf{w}^\top \phi(x^{(i)}))^2\right)$$



How to find $\mathbf{w}$?

Model:

$$p(y^{(i)}|x^{(i)}) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{1}{2\sigma^2}(y^{(i)} - \mathbf{w}^\top \phi(x^{(i)}))^2\right)$$

Maximize likelihood of given dataset $\mathcal{D} = \{(x^{(i)}, y^{(i)})\}$ assuming samples to be drawn independently from an identical distribution (i.i.d.).

$$p(\mathcal{D}) = \prod_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} p(y^{(i)}|x^{(i)}) = \frac{1}{(2\pi\sigma^2)^{\frac{N}{2}}} \exp\left(-\frac{1}{2\sigma^2} \|\mathbf{Y} - \Phi^\top \mathbf{w}\|_2^2\right)$$

$$\arg \max_{\mathbf{w}} p(\mathcal{D}) = \arg \min_{\mathbf{w}} \frac{1}{2} \|\mathbf{Y} - \Phi^\top \mathbf{w}\|_2^2$$

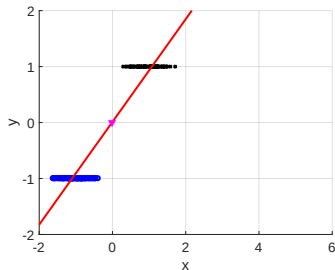
Linear regression for classification?

$$y^{(i)} \in \{-1, 1\}$$

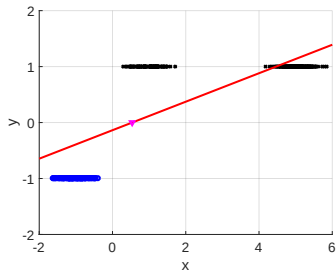
Model:

$$y = w_1 x + w_0$$

threshold at $y = 0$



perfect classification

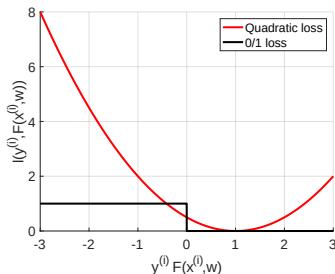


decision boundary shifted

Why is this?

Linear regression: Quadratic loss (recall $y^{(i)} \in \{-1, 1\}$)

$$\begin{aligned}\ell(y_i, \phi(x^{(i)})^\top \mathbf{w}) &= \frac{1}{2}(y^{(i)} - \phi(x^{(i)})^\top \mathbf{w})^2 \\ &\stackrel{(y^{(i)})^2=1}{=} \frac{1}{2}(1 - y^{(i)} \underbrace{\phi(x^{(i)})^\top \mathbf{w}}_{F(x^{(i)}, \mathbf{w})})^2 \\ &\quad \underbrace{\hspace{10em}}_{F(x^{(i)}, \mathbf{w}, y^{(i)})}\end{aligned}$$



We penalize samples that are ‘very easy to classify.’

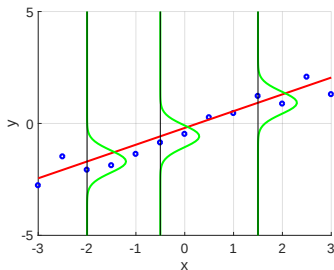
How to fix this?
Next topic...

Logistic Regression

A probabilistic interpretation of linear regression ($y^{(i)} \in \mathbb{R}$):

Model: Gaussian distribution

$$p(y^{(i)}|x^{(i)}) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{1}{2\sigma^2}(y^{(i)} - \mathbf{w}^\top \phi(x^{(i)}))^2\right)$$



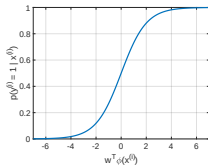
Logistic Regression:

Another probabilistic formulation for classification ($y^{(i)} \in \{-1, 1\}$):

Model:

$$p(y^{(i)} = 1|x^{(i)}) = \frac{1}{1 + \exp(-\mathbf{w}^T \phi(x^{(i)}))}$$

$$p(y^{(i)} = -1|x^{(i)}) = 1 - p(y^{(i)} = 1|x^{(i)}) = \frac{1}{1 + \exp(\mathbf{w}^T \phi(x^{(i)}))}$$



Taken together:

$$p(y^{(i)}|x^{(i)}) = \frac{1}{1 + \exp(-y^{(i)} \mathbf{w}^T \phi(x^{(i)}))}$$

What to do with this model?

$$p(y^{(i)}|x^{(i)}) = \frac{1}{1 + \exp(-y^{(i)} \mathbf{w}^T \phi(x^{(i)}))}$$

Recall that we are given a dataset $\mathcal{D} = \{(x^{(i)}, y^{(i)})\}$. How about we choose $\mathbf{w}$ which maximizes the likelihood/probability of this dataset?

Assumption:

Samples/Data points are i.i.d.

$$p(\mathcal{D}) = \prod_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} p(y^{(i)}|x^{(i)})$$

Choose $\mathbf{w}$ to maximize probability:

$$\max_{\mathbf{w}} \prod_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} p(y^{(i)}|x^{(i)})$$

Model:

$$p(y^{(i)}|x^{(i)}) = \frac{1}{1 + \exp(-y^{(i)} \mathbf{w}^T \phi(x^{(i)}))}$$

Task:

$$\arg \max_{\mathbf{w}} \prod_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} p(y^{(i)}|x^{(i)}) = \arg \min_{\mathbf{w}} \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} -\log p(y^{(i)}|x^{(i)})$$

Combined:

$$\min_{\mathbf{w}} \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \log \left(1 + \exp(-y^{(i)} \mathbf{w}^T \phi(x^{(i)})) \right)$$

Comparison

Linear regression

Program:

$$\min_{\mathbf{w}} \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \frac{1}{2} (1 - y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})})^2$$

$\underbrace{\hspace{10em}}_{F(x^{(i)}, \mathbf{w}, y^{(i)})}$

Logistic regression

Program:

$$\min_{\mathbf{w}} \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \log \left(1 + \exp(-y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})}) \right)$$

$\underbrace{\hspace{10em}}_{F(x^{(i)}, \mathbf{w}, y^{(i)})}$

Empirical risk minimization:

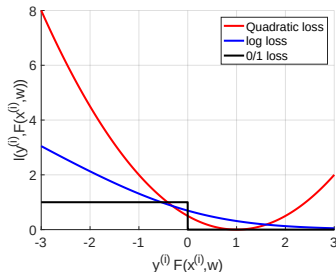
$$\min_{\mathbf{w}} \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \ell(y^{(i)}, F(x^{(i)}, \mathbf{w}))$$

Linear regression:

$$\min_{\mathbf{w}} \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \frac{1}{2} (1 - y^{(i)}) \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, w)}^2$$

Logistic regression:

$$\min_{\mathbf{w}} \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \log \left(1 + \exp(-y^{(i)}) \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, w)} \right)$$



How to optimize

$$\min_{\mathbf{w}} \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \log \left(1 + \exp(-y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})}) \right)$$

Can we set the gradient to zero and solve for $\mathbf{w}$?

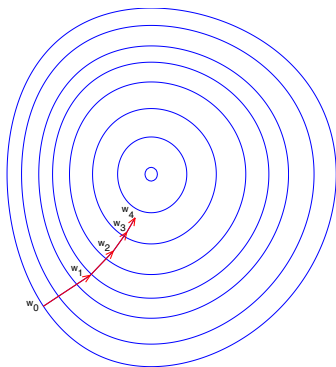
$$\sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \frac{-y^{(i)} \exp(-y^{(i)} \mathbf{w}^T \phi(x^{(i)}))}{1 + \exp(-y^{(i)} \mathbf{w}^T \phi(x^{(i)}))} \phi(x^{(i)}) = 0$$

No analytic solution for $\mathbf{w}$ in general

How to optimize

$$\min_{\mathbf{w}} \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \log \left(1 + \exp(-y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})}) \right)$$

Gradient descent: (walking down a mountain)



To solve

$$\min_{\mathbf{w}} f(\mathbf{w}) := \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \log \left(1 + \exp(-y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})}) \right)$$

we can use its gradient:

$$\nabla_{\mathbf{w}} f(\mathbf{w}) = \sum_{(x^{(i)}, y^{(i)}) \in \mathcal{D}} \frac{-y^{(i)} \exp(-y^{(i)} \mathbf{w}^T \phi(x^{(i)}))}{1 + \exp(-y^{(i)} \mathbf{w}^T \phi(x^{(i)}))} \phi(x^{(i)})$$

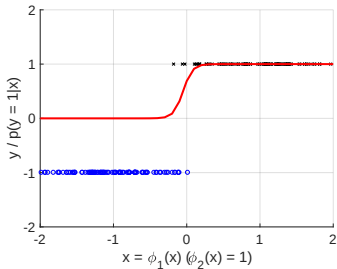
Simple algorithm: Initialize $t = 0$, $\mathbf{w}_t$, and stepsize α

- Compute gradient $\mathbf{g}_t = \nabla_{\mathbf{w}} f(\mathbf{w}_t)$
- Update parameters $\mathbf{w}_{t+1} \leftarrow \mathbf{w}_t - \alpha \mathbf{g}_t$
- Update $t \leftarrow t + 1$

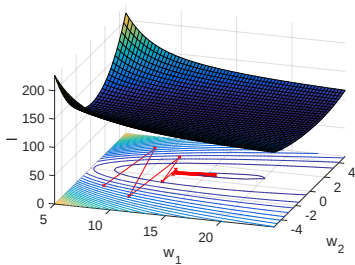
More complex algorithms may be ‘better.’

Example:

Data

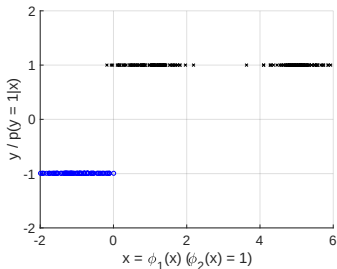


Loss

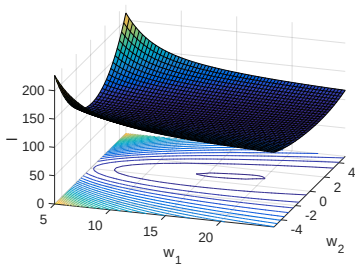


Example:

Data

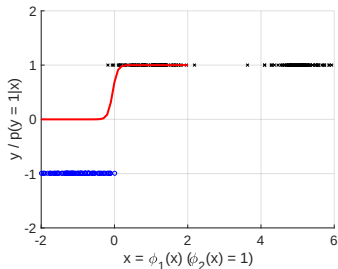


Loss

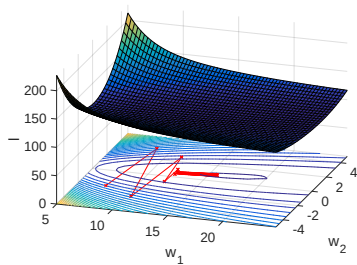


Example:

Data



Loss



Comparison:

Linear regression:

- Closed form solution
- Gaussian probability model
- Not too well suited for classification

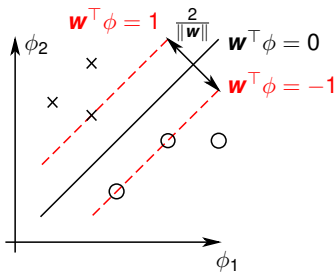
Logistic regression:

- Well suited for binary classification
- Logistic probability model
- No closed form solution

Support Vector Machines

Binary SVM

Intuitively:



Maximize margin $\frac{2}{\|\mathbf{w}\|}$:

$$\min_{\mathbf{w}} \frac{1}{2} \|\mathbf{w}\|_2^2 \quad \text{s.t.} \quad y^{(i)} \mathbf{w}^T \phi(x^{(i)}) \geq \underbrace{1}_{\text{Taskloss: } L} \quad \forall (\phi(x^{(i)}), y^{(i)}) \in \mathcal{D}$$

Note: any value $L \geq 0$ is okay.

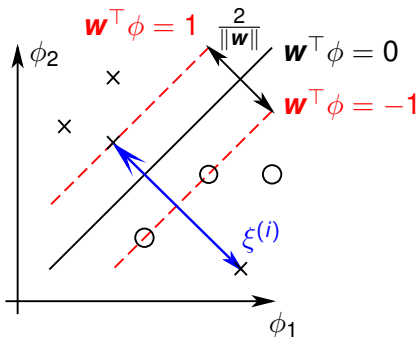
Issue: what if data not linearly separable?

Binary SVM

Introduce slack variables ξ :

$$\min_{\mathbf{w}, \xi^{(i)} \geq 0} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \xi^{(i)} \quad \text{s.t.} \quad y^{(i)} \mathbf{w}^\top \phi(x^{(i)}) \geq 1 - \xi^{(i)} \quad \forall i \in \mathcal{D}$$

Intuitively:



Binary SVM:

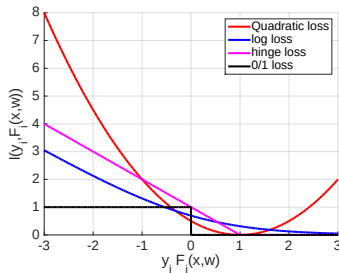
$$\min_{\mathbf{w}, \xi^{(i)} \geq 0} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \xi^{(i)} \quad \text{s.t.} \quad y^{(i)} \mathbf{w}^\top \phi(x^{(i)}) \geq 1 - \xi^{(i)} \quad \forall i \in \mathcal{D}$$

Equivalent:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \max\{0, 1 - y^{(i)} \mathbf{w}^\top \phi(x^{(i)})\}$$

Empirical risk minimization:

$$\min_{\mathbf{w}} R(\mathbf{w}) + \sum_{i \in \mathcal{D}} \ell(y^{(i)}, F(x^{(i)}, \mathbf{w}))$$



How to optimize the binary SVM objective (primal problem):

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \max\{0, 1 - y^{(i)} \mathbf{w}^\top \phi(\mathbf{x}^{(i)})\}$$

Approaches:

- Optimize the primal via gradient descent
- Optimize the corresponding dual problem

Recap:

- Linear regression:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \frac{1}{2} (1 - y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})})^2$$

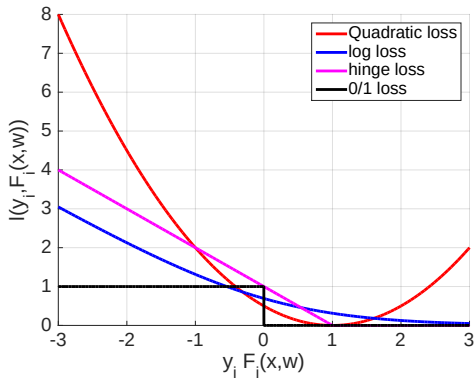
- Logistic regression:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \log \left(1 + \exp(-y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})}) \right)$$

- Binary SVM:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \max\{0, \underbrace{1}_{\text{taskloss}} - y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})}\}$$

Recap:



Combining log- and hinge-loss:

$$\begin{aligned}\lim_{\epsilon \rightarrow 0} \epsilon \log \left(1 + \exp \frac{-F}{\epsilon} \right) &= \\ &\stackrel{F \geq 0}{=} 0 \\ &\stackrel{F \leq 0}{=} \lim_{\epsilon \rightarrow 0} \frac{\log \left(1 + \exp \frac{-F}{\epsilon} \right)}{1/\epsilon} \\ &= \lim_{\epsilon \rightarrow 0} \frac{\frac{\exp \frac{-F}{\epsilon}}{1 + \exp \frac{-F}{\epsilon}} \cdot (F/\epsilon^2)}{-1/\epsilon^2} \\ &= \lim_{\epsilon \rightarrow 0} \frac{1}{1 + \exp \frac{F}{\epsilon}} \cdot (-F) \\ &= -F\end{aligned}$$

In summary:

$$\lim_{\epsilon \rightarrow 0} \epsilon \log \left(1 + \exp \frac{-F}{\epsilon} \right) = \max\{0, -F\}$$

SVM as 0-temperature limit of logistic regression

Recap:

- Linear regression:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \frac{1}{2} (1 - y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})})^2$$

- Logistic regression:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \log \left(1 + \exp(-y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})}) \right)$$

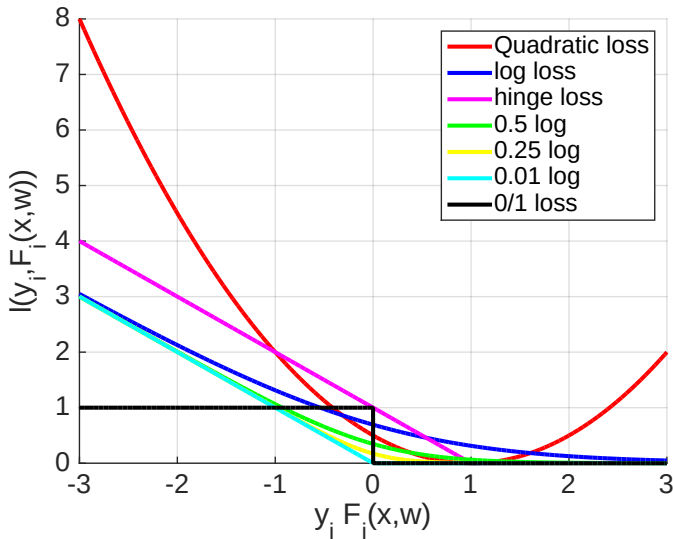
- Binary SVM:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \max\{0, \underbrace{1}_{\text{taskloss}} - y^{(i)} \underbrace{\mathbf{w}^T \phi(x^{(i)})}_{F(x^{(i)}, \mathbf{w})}\}$$

- General binary classification:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \epsilon \log \left(1 + \exp \left(\frac{L - y^{(i)} \mathbf{w}^T \phi(x^{(i)})}{\epsilon} \right) \right)$$

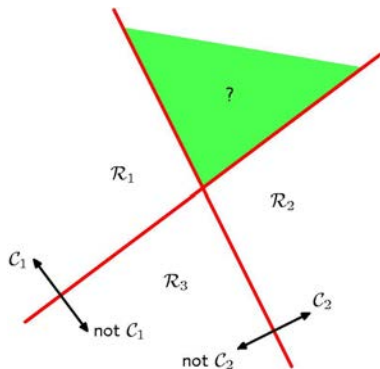
Different loss functions



Multiclass Classification and Kernel Methods

Multiclass classification: How to classify between K classes?

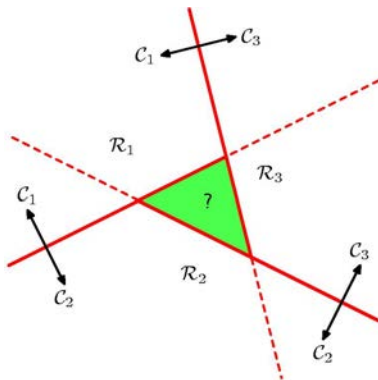
- **1 vs all** or **1 vs rest** classifier
- Use $K - 1$ classifiers, each solving a two class problem which separates a point in class k from point not in this class



- Issue: more than one good answer or no good answer

Multiclass classification: How to classify between K classes?

- **1 vs 1** classifier
- Use $K(K - 1)/2$ two-way classifiers, one for each possible pair of classes
- Classify sample according to majority vote



- Issue: two-way preferences need not be transitive

Multiclass classification: How to classify between K classes?

- Use a multinomial distribution over $y \in \{0, 1, \dots, K - 1\}$
- Use K weight vectors $\mathbf{w}_{(y)}$
- How to parameterize the multinomial distribution?

$$p(y = k | \mathbf{x}^{(i)}) = \frac{\exp \mathbf{w}_{(k)}^\top \phi(\mathbf{x}^{(i)})}{\sum_{j \in \{0, 1, \dots, K-1\}} \exp \mathbf{w}_{(j)}^\top \phi(\mathbf{x}^{(i)})}$$

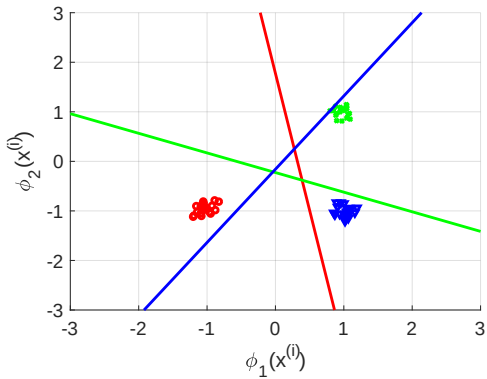
- We are using one parameter vector $\mathbf{w}_{(y)}$ per class
- Maximizing the likelihood as before

$$\arg \max_{\mathbf{w}} \prod_{(\mathbf{x}^{(i)}, y^{(i)}) \in \mathcal{D}} p(y = y^{(i)} | \mathbf{x}^{(i)}) = \arg \min_{\mathbf{w}} \sum_{(\mathbf{x}^{(i)}, y^{(i)}) \in \mathcal{D}} -\log p(y = y^{(i)} | \mathbf{x}^{(i)})$$

Questions:

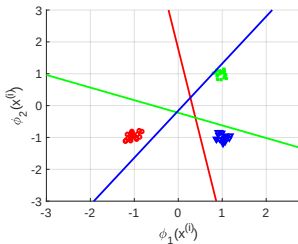
- Example to build intuition?
- How does this relate to our earlier binary setting?

Example: ($\phi_3(x^{(i)}) = 1$)

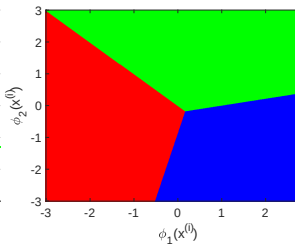


Example: $(\phi_3(x^{(i)}) = 1)$

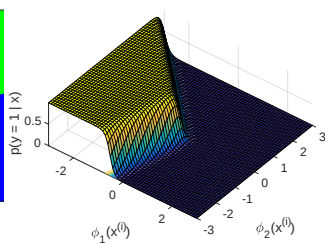
Setup



Decision Boundary

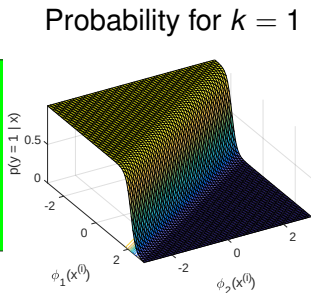
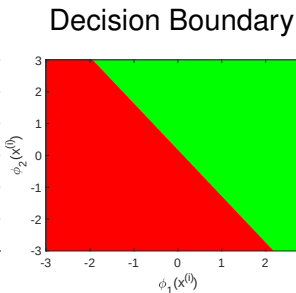
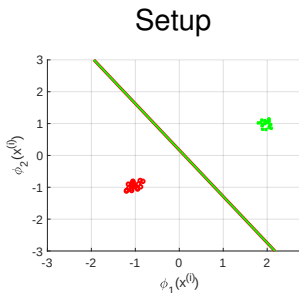


Probability for $k = 1$



Example: What happens when we have two classes?

- We get two hyperplanes (lines in 2D). How do they related to binary logistic regression?

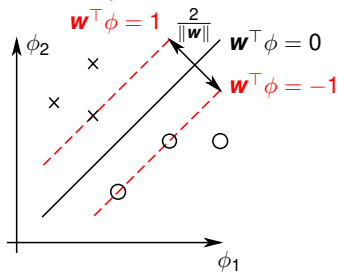


Lines coincide ($\mathbf{w}_{(1)} = -\mathbf{w}_{(2)}$)

How about multi-class SVM?

Recap: Binary SVM

Intuitively:



$$u_1 + c\mathbf{w} = u_2 \quad \mathbf{w}^\top u_2 = -1$$

$$\mathbf{w}^\top (u_1 + c\mathbf{w}) = -1$$

$$c = \frac{-2}{\|\mathbf{w}\|_2^2}$$

$$\|u_1 - u_2\|_2 = \|-c\mathbf{w}\|_2 = \frac{2}{\|\mathbf{w}\|_2^2} \|\mathbf{w}\| = \frac{2}{\|\mathbf{w}\|}$$

Maximize margin $\frac{2}{\|\mathbf{w}\|}$:

$$\min_{\mathbf{w}} \frac{1}{2} \|\mathbf{w}\|_2^2 \quad \text{s.t.} \quad y^{(i)} \mathbf{w}^\top \phi(x^{(i)}) \geq \underbrace{1}_{\text{Taskloss: } L} \quad \forall (x^{(i)}, y^{(i)}) \in \mathcal{D}$$

Note: any value $L \geq 0$ is okay.

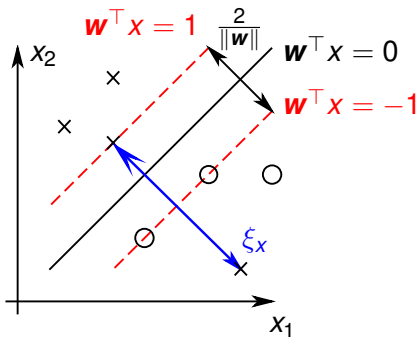
Issue: what if data not linearly separable?

Recap: Binary SVM

Introduce slack variables $\xi^{(i)}$:

$$\min_{\mathbf{w}, \xi^{(i)} \geq 0} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \xi^{(i)} \quad \text{s.t.} \quad y^{(i)} \mathbf{w}^\top \phi(\mathbf{x}^{(i)}) \geq 1 - \xi^{(i)} \quad \forall i \in \mathcal{D}$$

Intuitively:



How about multi-class SVM? Recall $y \in \{0, 1, \dots, K - 1\}$ and class vectors $\mathbf{w}_{(y)}$

What do we want?

$$\mathbf{w}_{y^{(i)}}^T \phi(\mathbf{x}^{(i)}) \geq \mathbf{w}_{\hat{y}}^T \phi(\mathbf{x}^{(i)}) \quad \forall i \in \mathcal{D}, \hat{y} \in \{0, 1, \dots, K - 1\}$$

$$\min_{\mathbf{w}, \xi^{(i)} \geq 0} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \xi^{(i)} \quad \text{s.t.} \quad \mathbf{w}_{y^{(i)}}^T \phi(\mathbf{x}^{(i)}) - \mathbf{w}_{\hat{y}}^T \phi(\mathbf{x}^{(i)}) \geq 1 - \xi^{(i)} \quad \forall i, \hat{y}$$

Multiclass SVM objective:

$$\min_{\mathbf{w}, \xi^{(i)} \geq 0} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \xi^{(i)} \quad \text{s.t.} \quad \mathbf{w}_{y^{(i)}}^T \phi(\mathbf{x}^{(i)}) - \mathbf{w}_{\hat{y}}^T \phi(\mathbf{x}^{(i)}) \geq 1 - \xi^{(i)} \quad \forall i, \hat{y}$$

Define mapping:

$$\mathbf{w} = \begin{bmatrix} \mathbf{w}_0 \\ \vdots \\ \mathbf{w}_{K-1} \end{bmatrix} \quad \psi(\mathbf{x}^{(i)}, \mathbf{y}^{(i)}) = \begin{bmatrix} \phi(\mathbf{x}^{(i)}) \delta(\mathbf{y}^{(i)} = 0) \\ \vdots \\ \phi(\mathbf{x}^{(i)}) \delta(\mathbf{y}^{(i)} = K - 1) \end{bmatrix}$$

Equivalent formulation:

$$\min_{\mathbf{w}, \xi^{(i)} \geq 0} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \xi^{(i)} \quad \text{s.t.} \quad \mathbf{w}^T (\psi(\mathbf{x}^{(i)}, \mathbf{y}^{(i)}) - \psi(\mathbf{x}^{(i)}, \hat{y})) \geq 1 - \xi^{(i)} \quad \forall i, \hat{y}$$

Multiclass SVM objective:

$$\min_{\mathbf{w}, \xi^{(i)} \geq 0} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \xi^{(i)} \quad \text{s.t.} \quad \mathbf{w}^T (\psi(\mathbf{x}^{(i)}, \mathbf{y}^{(i)}) - \psi(\mathbf{x}^{(i)}, \hat{\mathbf{y}})) \geq 1 - \xi^{(i)} \quad \forall i, \hat{\mathbf{y}}$$

Let's get rid of the slack variable:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \underbrace{\max\{0, \max_{\hat{\mathbf{y}}} (1 - \mathbf{w}^T (\psi(\mathbf{x}^{(i)}, \mathbf{y}^{(i)}) - \psi(\mathbf{x}^{(i)}, \hat{\mathbf{y}})))\}}_{\geq 0}$$

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \max_{\hat{\mathbf{y}}} (1 - \mathbf{w}^T (\psi(\mathbf{x}^{(i)}, \mathbf{y}^{(i)}) - \psi(\mathbf{x}^{(i)}, \hat{\mathbf{y}})))$$

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \underbrace{\max_{\hat{\mathbf{y}}} (1 + \mathbf{w}^T \psi(\mathbf{x}^{(i)}, \hat{\mathbf{y}})) - \mathbf{w}^T \psi(\mathbf{x}^{(i)}, \mathbf{y}^{(i)})}_{\text{Loss-augmented inference}}$$

How does multiclass SVM relate to multiclass logistic regression?

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \underbrace{\max_{\hat{y}} (1 + \mathbf{w}^T \psi(x^{(i)}, \hat{y})) - \mathbf{w}^T \psi(x^{(i)}, y^{(i)})}_{\text{Loss-augmented inference}}$$

Same temperature argument as in the last lecture.

Solution:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \epsilon \ln \sum_{\hat{y}} \exp \frac{L(y^{(i)}, \hat{y}) + \mathbf{w}^T \psi(x^{(i)}, \hat{y})}{\epsilon} - \mathbf{w}^T \psi(x^{(i)}, y^{(i)})$$

Our current framework:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \epsilon \ln \sum_{\hat{y}} \exp \frac{L(y^{(i)}, \hat{y}) + \mathbf{w}^T \psi(x^{(i)}, \hat{y})}{\epsilon} - \mathbf{w}^T \psi(x^{(i)}, y^{(i)})$$

What is a possible issue/limitation?

Linearity in the feature space $\psi(x^{(i)}, y^{(i)})$.

How to fix this?

Use Kernels

Dual of Logistic Regression:

$$\max_{0 \leq \lambda^{(i)} \leq 1} g(\lambda) := \frac{-1}{2C} \left\| \sum_i \lambda^{(i)} y^{(i)} \phi(x^{(i)}) \right\|_2^2 + \sum_i H(\lambda^{(i)})$$

Prediction with dual variables:

$$\mathbf{w}^\top \phi(x) = \frac{1}{C} \sum_i \lambda^{(i)} y^{(i)} \phi(x^{(i)})^\top \phi(x)$$

Dual of SVM:

$$\max_{0 \leq \alpha \leq 1} g(\alpha) := \frac{-1}{2C} \left\| \sum_i \alpha^{(i)} y^{(i)} \phi(x^{(i)}) \right\|_2^2 + \sum_i \alpha^{(i)}$$

Prediction with dual variables:

$$\mathbf{w}^\top \phi(x) = \frac{1}{C} \sum_i \alpha^{(i)} y^{(i)} \phi(x^{(i)})^\top \phi(x)$$

Observation:

Many inner products $\phi(x^{(i)})^\top \phi(x^{(j)})$ between different samples

Kernel trick:

Replace inner products with call to kernel $\kappa(x^{(i)}, x^{(j)}) \in \mathbb{R}$

Advantage:

No need to explicitly construct feature vector $\phi(x)$

But:

We need to ensure that κ is a valid kernel, i.e., it corresponds to an inner product in some feature space

Example:

$$\begin{aligned}\kappa(x, z) &= (x^\top z)^2 \\ &= x_1^2 z_1^2 + 2x_1 z_1 x_2 z_2 + x_2^2 z_2^2 \\ &= [x_1^2 \quad \sqrt{2}x_1 x_2 \quad x_2^2] [z_1^2 \quad \sqrt{2}z_1 z_2 \quad z_2^2]^\top \\ &= \phi(x)^\top \phi(z)\end{aligned}$$

How to test validity without explicitly constructing feature vectors:

- Gram matrix $(\mathbf{K})_{i,j} = \kappa(x^{(i)}, x^{(j)})$ is positive definite $\forall x^{(i)}, x^{(j)}$
- Decompose/Compose kernel into/from known kernels

Example kernels:

- Linear kernel:

$$\kappa(x^{(i)}, x^{(j)}) = x^{(i)\top} x^{(j)}$$

- Squared exponential (Gaussian) kernel:

$$\kappa(x^{(i)}, x^{(j)}) = \exp\left(-\frac{1}{2}(x^{(i)} - x^{(j)})^\top \Sigma^{-1}(x^{(i)} - x^{(j)})\right)$$

- Sigmoid kernel:

$$\kappa(x^{(i)}, x^{(j)}) = \tanh(a \cdot x^{(i)\top} x^{(j)} + b)$$

Techniques for constructing new kernels ($c > 0$):

- Positive scaling:

$$\kappa(x^{(i)}, x^{(j)}) = c\kappa_1(x^{(i)}, x^{(j)})$$

- Exponentiation:

$$\kappa(x^{(i)}, x^{(j)}) = \exp(\kappa_1(x^{(i)}, x^{(j)}))$$

- Addition:

$$\kappa(x^{(i)}, x^{(j)}) = \kappa_1(x^{(i)}, x^{(j)}) + \kappa_2(x^{(i)}, x^{(j)})$$

- Multiplication with function:

$$\kappa(x^{(i)}, x^{(j)}) = f(x^{(i)})\kappa_1(x^{(i)}, x^{(j)})f(x^{(j)})$$

- Multiplication:

$$\kappa(x^{(i)}, x^{(j)}) = \kappa_1(x^{(i)}, x^{(j)})\kappa_2(x^{(i)}, x^{(j)})$$

Deep Neural Networks

Goals of this lecture

- Getting to know deep nets
- Understanding forward and backward pass
- Learning about deep net components

Reading material

- I. Goodfellow et al.; Deep Learning; Chapters 6-9

Our current framework:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \epsilon \ln \sum_{\hat{y}} \exp \frac{L(y^{(i)}, \hat{y}) + \mathbf{w}^T \psi(x^{(i)}, \hat{y})}{\epsilon} - \mathbf{w}^T \psi(x^{(i)}, y^{(i)})$$

What is a possible issue/limitation?

Linearity in the feature space $\psi(x, y)$. Fix: use kernels. But still learning a model **linear** in the parameters $\mathbf{w}$

How to fix this?

Replace $\mathbf{w}^T \psi(x, y)$ with a general function $F(\mathbf{w}, x, y)$

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \epsilon \ln \sum_{\hat{y}} \exp \frac{L(y^{(i)}, \hat{y}) + F(\mathbf{w}, x^{(i)}, \hat{y})}{\epsilon} - F(\mathbf{w}, x^{(i)}, y^{(i)})$$

General framework:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \epsilon \ln \sum_{\hat{y}} \exp \frac{L(y^{(i)}, \hat{y}) + F(\mathbf{w}, x^{(i)}, \hat{y})}{\epsilon} - F(\mathbf{w}, x^{(i)}, y^{(i)})$$

How to get to

- **Logistic regression**
- **Binary SVM**
- **Multiclass regression**
- **Multiclass SVM**
- **Deep Learning**

Deep Learning:

What function $F(\mathbf{w}, x, y) \in \mathbb{R}$ to choose?

- Choose any differentiable composite function

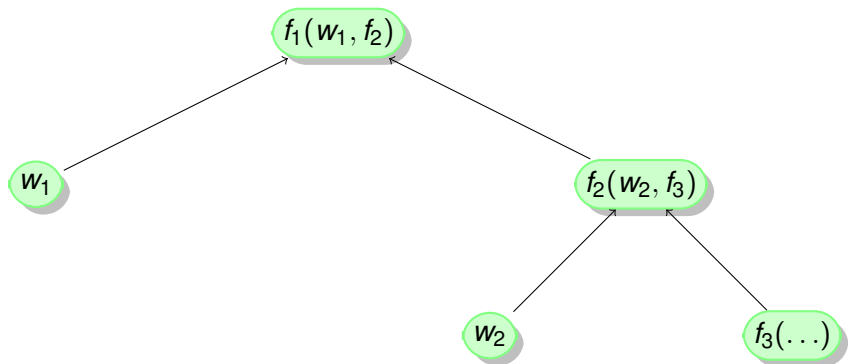
$$F(\mathbf{w}, x, y) = f_1(\mathbf{w}_1, y, f_2(\mathbf{w}_2, f_3(\dots f_n(\mathbf{w}_n, x) \dots)))$$

- More generally: functions can be represented by an acyclic graph (computation graph)

Example:

$$F(\mathbf{w}, x, y) = f_1(w_1, f_2(w_2, f_3(\dots)))$$

Nodes are weights, data, and functions:



Internal representation used by deep net packages.

Non-linear Functions

A neuron's activation is a non-linear function which takes $\xi \in \mathbb{R}$ as input and produces $\sigma(\xi) \in \mathbb{R}$ as output.

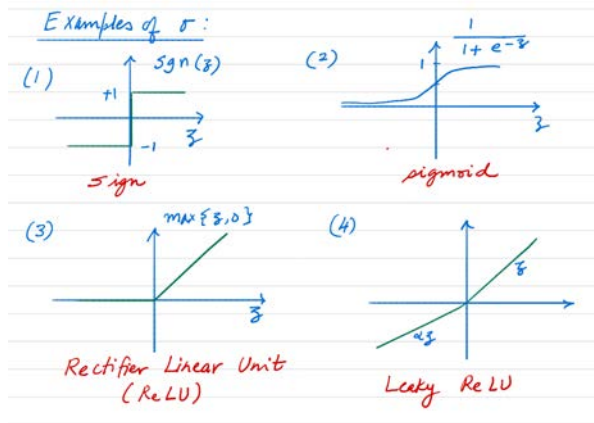


Figure: Examples of σ

Simplified Representation

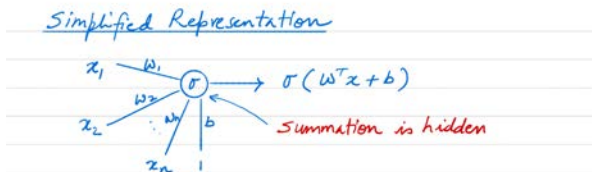
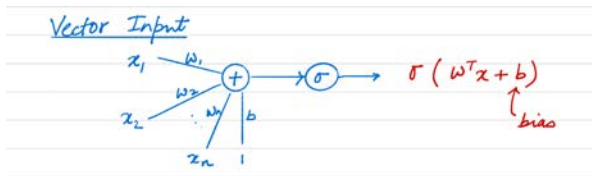
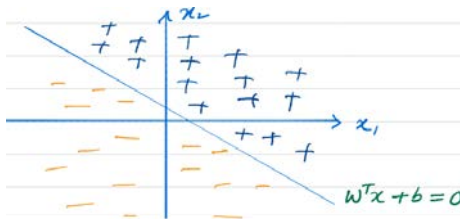


Figure: Enter Caption

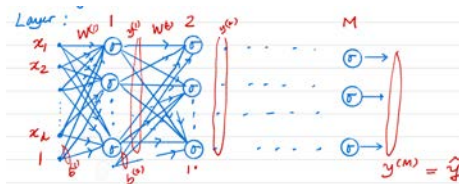
Bias and Boundary

Bias b is important, e.g., if σ is the sign function, we can output +1 if $\omega^T x \geq b$, and -1 if $\omega^T x < b$
 $\omega^T x = 0$ is in $\Re^n$, and a neuron simply divides the input space into two parts corresponding to +1 and -1



Multilayer Neural Networks allow us to divide high dimensional space in a less complicated way and approximate functions by mapping inputs and outputs of NN.

Multilayer Neural Network (Feed Forward)



- Number of neurons in each layer can be different.
- All weights on edge connecting layers $m-1$ and m are in matrix $W^{(m)}$, with $W_{ij}^{(m)}$ being the weight connecting output j of layer $m-1$ with neuron i of layer m
- Input to network is vector x ; output of layer m is vector $y^{(m)}$
 $y_i^{(1)} = \sigma(x_i^{(1)})$, with $x_i^{(1)} = \sum_j W_{ij}^{(1)} x_j + b_j^{(1)}$
- This can be written compactly as: $y^{(1)} = \sigma(x^{(1)})$, with
 $x^{(1)} = w^{(1)}x + b^{(1)}$, for σ is applied to each element of $x^{(1)}$

Multilayer Neural Network (Feed Forward)

- $y^{(1)} = \sigma(x^{(1)})$, with $x^{(1)} = w^{(1)}x + b^{(1)}$
- similarly, $y^{(2)} = \sigma(x^{(2)})$, with $x^{(2)} = w^{(2)}y^{(1)} + b^{(2)}$
-
- $y^{(M)} = \sigma(x^{(M)})$, with $x^{(M)} = w^{(M)}y^{(M-1)} + b^{(M)}$
- We want the output of NN approximates some desired functions well: $\hat{y} = y^{(M)} \approx f(x) = y$
- How is the function approximated? Use labeled training data, i.e. inputs for which we know the correct outputs, even though f is unknown.
- Minimize empirical loss on training data $J = \sum_{n=1}^N L(y[n], \hat{y}[n])$, $\hat{y} = y^{(M)}$

Multilayer Neural Network (Feed Forward)

- In case the desired output y is scalar, then L is often taken to be the square loss: $L(y, \hat{y}) = |y - \hat{y}|^2$ means each layer has only one neuron
- J is a function of W and b
- We wish to minimize J using a gradient descent procedure (assume the loss function is continuous and differentiable)
- Need to compute gradient of J with respect to all weights and biases
- We can compute the gradient for each term separately (partial derivative)
- we need: $\frac{\partial L}{\partial w_{ij}^l}$ and $\frac{\partial L}{\partial b_i^l}$

Back Propagation

- Exploit network structure and Chain Rule to compute gradient terms efficiently
- Layer M Recall $y^{(M)} = \sigma(x^{(M)})$, $x^{(M)} = w^{(M)}y^{(M-1)} + b^{(M)}$
- $\frac{\partial L}{\partial w_{ij}^{(M)}} = \frac{\partial L}{\partial y_i^{(M)}} * \frac{\partial y_i^{(M)}}{\partial w_{ij}^{(M)}} = \frac{\partial L}{\partial y_i^{(M)}} * \frac{\partial y_i^{(M)}}{\partial x_i^{(M)}} * \frac{\partial x_i^{(M)}}{\partial w_{ij}^{(M)}}$
- $\frac{\partial L}{\partial y_i^{(M)}}$ can easily be computed in practice for each data point ($x[n]$, $y[n]$)
- For example, if $y^{(M)}$ and y are scalar, and we are using square loss, then $\frac{\partial L}{\partial y_i^{(M)}}(y[n], y^{(M)}[n]) = 2(y^{(M)}[n] - y[n])$
- $\frac{\partial y_i^{(M)}}{\partial x_i^{(M)}} = \frac{d\sigma(x_i^{(M)})}{dx_i^{(M)}} = \sigma'(x_i^{(M)})$, assume σ is differentiable
- $x_i^{(M)} = \sum_k w_{(ik)}^{(M)} y_{(k)}^{(M-1)} + b_i^{(M)}$, $\frac{\partial x_i^{(M)}}{\partial w_{ij}^{(M)}} = y_j^{(M-1)}$
- Thus, $\frac{\partial L}{\partial w_{ij}^{(M)}} = \frac{\partial L}{\partial y_i^{(M)}} * \sigma'(x_i^{(M)}) * y_j^{(M-1)}$

Back Propagation

- Similarly, $\frac{\partial L}{\partial b_i^{(M)}} = \frac{\partial L}{\partial y_i^{(M)}} * \frac{\partial y_i^{(M)}}{\partial x_i^{(M)}} * \frac{\partial x_i^{(M)}}{\partial b_i^{(M)}} = \frac{\partial L}{\partial y_i^{(M)}} * \sigma'(x_i^{(M)})$
- Lay m, $1 \leq m \leq M$ Recall
 $y^{(m)} = \sigma(x^{(m)}), x^{(m)} = w^{(m)}y^{(m-1)} + b^{(m)}$
- $\frac{\partial L}{\partial w_{ij}^{(m)}} = \frac{\partial L}{\partial y_i^{(m)}} * \frac{\partial y_i^{(m)}}{\partial x_i^{(m)}} * \frac{\partial x_i^{(m)}}{\partial w_{ij}^{(m)}} = \frac{\partial L}{\partial y_i^{(m)}} * \sigma'(x_i^{(m)}) * y_j^{(m-1)}$
- $y^{(m+1)} = \sigma(x^{(m+1)}), x^{(m+1)} = w^{(m+1)}y^{(m)} + b^{(m+1)}$
- $\frac{\partial L}{\partial y_i^{(m)}} = \sum_k \frac{\partial L}{\partial y_k^{(m+1)}} * \frac{\partial y_k^{(m+1)}}{\partial x_k^{(m+1)}} * \frac{\partial x_k^{(m+1)}}{\partial y_i^{(m)}} = \sum_k \frac{\partial L}{\partial y_k^{(m+1)}} * \sigma'(x_k^{(m+1)}) * w_{(ki)}^{(m+1)}$
because m+1 is already computed in the previous step(layer)
- Similarly, $\frac{\partial L}{\partial b_i^{(m)}} = \frac{\partial L}{\partial y_i^{(m)}} * \frac{\partial y_i^{(m)}}{\partial x_i^{(m)}} * \frac{\partial x_i^{(m)}}{\partial b_i^{(m)}} = \frac{\partial L}{\partial y_i^{(m)}} * \sigma'(x_i^{(m)})$

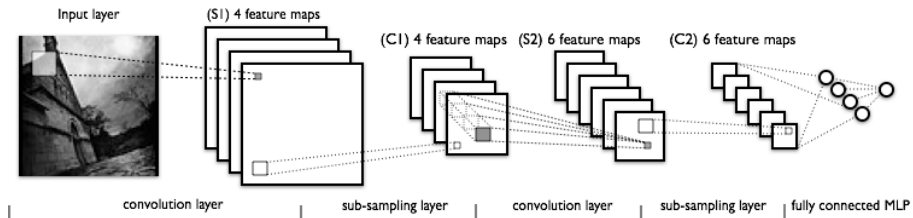
Back Propagation Summary

- Compute $\frac{\partial L}{\partial y_i^{(M)}}$ for all i
- For $m=M$ down to 1 $\frac{\partial L}{\partial w_{ij}^{(m)}} = \frac{\partial L}{\partial y_i^{(m)}} * \sigma'(x_i^{(m)}) * y_j^{(m-1)}$
- $\frac{\partial L}{\partial b_i^{(m)}} = \frac{\partial L}{\partial y_i^{(m)}} * \sigma'(x_i^{(m)})$
- where $\frac{\partial L}{\partial y_i^{(m)}} = \sum_k \frac{\partial L}{\partial y_k^{(m+1)}} * \sigma'(x_k^{(m+1)}) * w_{(ki)}^{(m+1)}$

What are the individual functions/layers f_1, f_2 etc.?

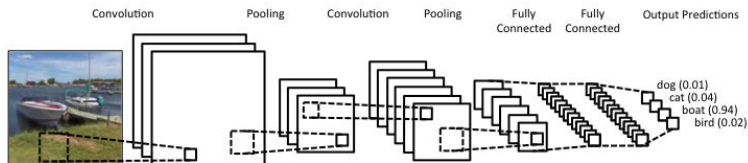
- Fully connected layers
- Convolutions
- Rectified linear units (ReLU): $\max\{0, x\}$
- Maximum-/Average pooling
- Soft-max layer
- Dropout

Example function architecture: LeNet



Decreasing spatial resolution and the increasing number of channels

Another deep net:



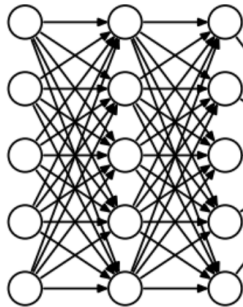
Those nets are structurally simple in that a layer's output is used as input for the next layer. This is not required.

Fully connected layer:

$$\mathbf{W}\mathbf{x} + \mathbf{b}$$

Trainable parameters $\mathbf{w}$:

- Matrix
- Bias



What's an issue with fully connected layers?

Issue with fully connected layers:

- Suppose the input is an image of size 256×256
- Let the output of this layer have identical size
- How many weights are necessary?

$$2^{32} = 4,294,967,296$$

- How to share weights?

Convolutions:

120	190	140	150	200
17	21	30	8	27
89	123	150	73	56
10	178	140	150	18
190	14	76	69	87

 $\times$

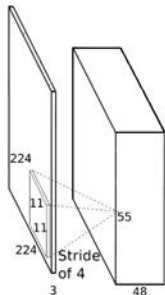
1/9	1/9	1/9
1/9	1/9	1/9
1/9	1/9	1/9

 $=$

	98	98	93	
	84	97	72	
	108	108	91	

Trainable parameters $\mathbf{w}$:

- Filters (width, height, depth, number)
- Bias

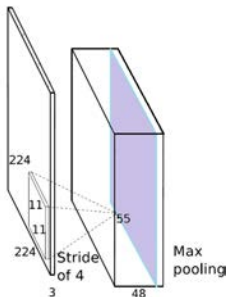


Maximum-/Average pooling

Maximum or average over a spatial region

Trainable parameters $\mathbf{w}$:

- None



Soft-max layer:

$$x \longrightarrow \frac{\exp x_i}{\sum_j \exp x_j}$$

Trainable parameters $\mathbf{w}$:

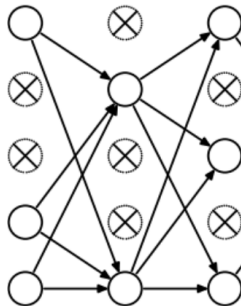
- None

Dropout layer:

Randomly set activations to zero

Trainable parameters $\mathbf{w}$:

- None



Recap: Our earlier framework:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \left(\epsilon \ln \sum_{\hat{y}} \exp \frac{L(y^{(i)}, \hat{y}) + \mathbf{w}^T \psi(x^{(i)}, \hat{y})}{\epsilon} - \mathbf{w}^T \psi(x^{(i)}, y^{(i)}) \right)$$

What is a possible issue/limitation?

Linearity in the feature space $\psi(x, y)$. Fix: use kernels. But still learning a model **linear** in the parameters $\mathbf{w}$

How to fix this?

Replace $\mathbf{w}^T \psi(x, y)$ with a general function $F(\mathbf{w}, x, y) \in \mathbb{R}$

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \left(\epsilon \ln \sum_{\hat{y}} \exp \frac{L(y^{(i)}, \hat{y}) + F(\mathbf{w}, x^{(i)}, \hat{y})}{\epsilon} - F(\mathbf{w}, x^{(i)}, y^{(i)}) \right)$$

Deep net training:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \left(\ln \sum_{\hat{y}} \exp F(\mathbf{w}, x^{(i)}, \hat{y}) - F(\mathbf{w}, x^{(i)}, y^{(i)}) \right)$$

Often also referred to as minimizing the regularized cross entropy:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 - \sum_{i \in \mathcal{D}} \sum_{\hat{y}} p_{\text{GT}}^{(i)}(\hat{y}) \ln p(\hat{y} | x^{(i)}) \quad \text{with} \quad \begin{cases} p_{\text{GT}}^{(i)}(\hat{y}) = \delta(\hat{y} = y^{(i)}) \\ p(\hat{y} | x) \propto \exp F(\mathbf{w}, x, \hat{y}) \end{cases}$$

What is C ? Weight decay (aka regularization constant)

$$\min_{\mathbf{w}} \underbrace{\frac{C}{2} \|\mathbf{w}\|_2^2}_{\text{weight decay}} - \underbrace{\sum_{i \in \mathcal{D}} \sum_{\hat{y}} p_{\text{GT}}^{(i)}(\hat{y}) \ln p(\hat{y} | x^{(i)})}_{\ell(\text{gt}, F)}$$

Program:

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \left(\ln \sum_{\hat{y}} \exp F(\mathbf{w}, x^{(i)}, \hat{y}) - F(\mathbf{w}, x^{(i)}, y^{(i)}) \right)$$

How to optimize this?

Stochastic gradient descent with momentum: What was this again?

Gradient of

$$\min_{\mathbf{w}} \frac{C}{2} \|\mathbf{w}\|_2^2 + \sum_{i \in \mathcal{D}} \left(\ln \sum_{\hat{y}} \exp F(\mathbf{w}, x^{(i)}, \hat{y}) - F(\mathbf{w}, x^{(i)}, y^{(i)}) \right)$$

is?

$$C\mathbf{w} + \sum_{i \in \mathcal{D}} \sum_{\hat{y}} \left(p(\hat{y}|x^{(i)}) - \delta(\hat{y} = y^{(i)}) \right) \frac{\partial F(\mathbf{w}, x^{(i)}, \hat{y})}{\partial \mathbf{w}}$$

How to compute this numerically:

- $p(\hat{y}|x) = \frac{\exp F(\mathbf{w}, x, \hat{y})}{\sum_{\tilde{y}} \exp F(\mathbf{w}, x, \tilde{y})}$ via soft-max which takes logits F as input
- $\frac{\partial F(\mathbf{w}, x, \hat{y})}{\partial \mathbf{w}}$ via backpropagation

Backpropagation example:

$$F(\mathbf{w}, x, y) = f_1(w_1, y, f_2(w_2, f_3(w_3, x))) \text{ with activations } \begin{cases} x_2 = f_3(w_3, x) \\ x_1 = f_2(w_2, x_2) \end{cases}$$

What is $\frac{\partial F(\mathbf{w}, x, y)}{\partial w_3}$?

$$\frac{\partial f_1}{\partial x_1} \cdot \frac{\partial x_1}{\partial x_2} \cdot \frac{\partial x_2}{\partial w_3} = \underbrace{\frac{\partial f_1}{\partial f_2} \cdot \frac{\partial f_2}{\partial f_3}}_{\text{chain rule}} \cdot \frac{\partial f_3}{\partial w_3}$$

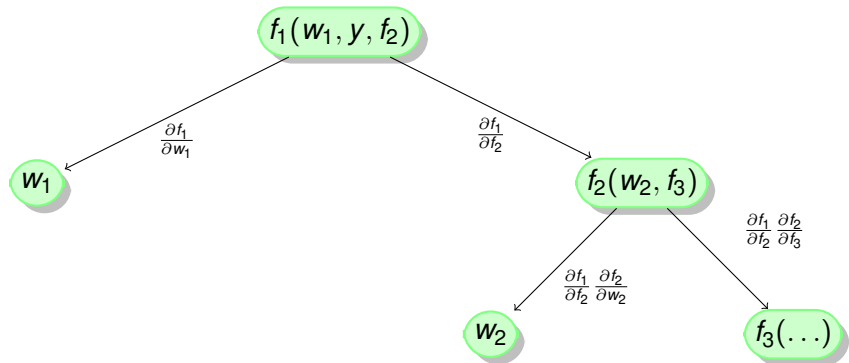
What is $\frac{\partial F(\mathbf{w}, x, y)}{\partial w_2}$?

$$\frac{\partial f_1}{\partial x_1} \cdot \frac{\partial x_1}{\partial w_2} = \underbrace{\frac{\partial f_1}{\partial f_2} \cdot \frac{\partial f_2}{\partial w_2}}_{\text{chain rule}}$$

Generally: To avoid repeated computation, backpropagation on an acyclic graph. Nodes in this graph are weights, data, and functions.

Composite function represented as acyclic graph

$$F(\mathbf{w}, x, y) = f_1(w_1, y, f_2(w_2, f_3(\dots)))$$



Repeated use of chain rule for efficient computation of all gradients

Remark:

Since $F(\mathbf{w}, x, y)$ is no longer constrained in any form, the loss function is generally no longer convex.

Implications:

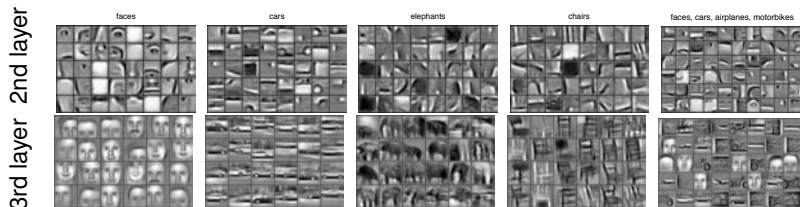
- We are no longer guaranteed to find the global optimum
- Initialization of $\mathbf{w}$ matters

Remark:

A deep net with a single fully connected layer is equivalent to logistic regression

Advantages of deep nets compared to usage of hand-crafted features:

Deep nets automatically learn feature space transformations (hierarchical abstractions of data) such that data is easily separable at the output



Disadvantage of deep nets compared to usage of features:

Deep nets are computationally demanding (GPUs) and require significant amounts of training data

Why this recent popularity:

- Sufficient computational resources
- Sufficient data
- Sufficient algorithmic advances
- Sufficient evidence that it works

This combination lead to significant performance improvements on many datasets

Algorithmic advances:

- Rectified linear unit ($\max\{0, x\}$) activation as opposed to sigmoid
 - ▶ Fixed the vanishing gradient problem for lower layers close to the input
- Dropout
 - ▶ Decorrelates different units, i.e., they learn different features
- Good initialization heuristics
 - ▶ Less prone to getting stuck in bad local optima
- Batch-Normalization during training
 - ▶ Normalizes data when training really deep nets
 - ▶ Normalize by subtracting mean and dividing by standard deviation