

Classification

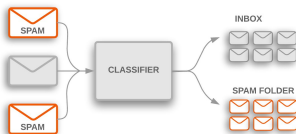
Ryan McDonald¹

AthNLP

September, 2026

¹Thanks to Andre Martins and Barbara Plank for use of materials.

Classifiers



"I love this movie.
I've seen it many times
and it's still awesome."



"This movie is bad.
I don't like it at all.
It's terrible."

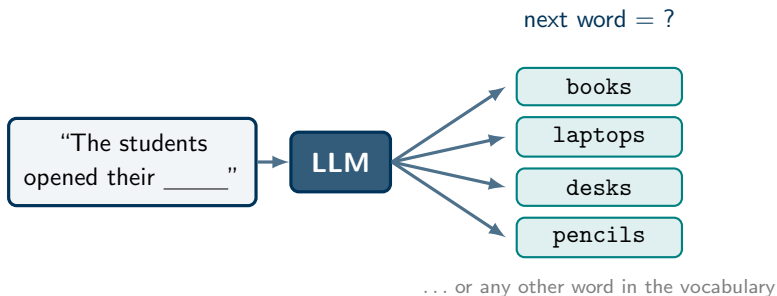


Set my alarm tomorrow for 10am -> **Alarm**

Quickest way to Boston -> **Navigation**

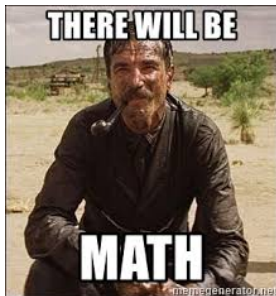
Why is there summer and winter -> **Answer seeking**

Classifiers



Predicting the next word is **classification** — the label set is the vocabulary.

Warning!



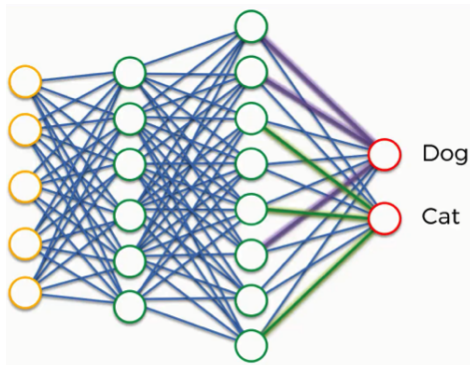
- Focus: machine learning fundamentals
 - Specific to language as input modality
 - Not specific applications
- If you miss a detail, don't worry
- Important to get broad concepts

This lecture is 1/2 about linear classifiers!

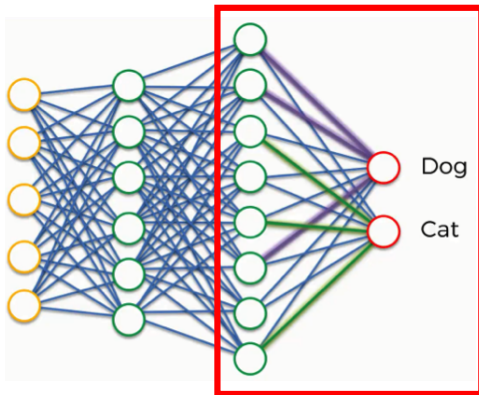
What's a linear classifier? It's 2026 and everybody uses neural networks.

- The underlying machine learning concepts are the same
- The theory (statistics and optimization) are much better understood
- Linear classifiers are **a component of neural networks**.

Linear Classifiers and Neural Networks



Linear Classifiers and Neural Networks

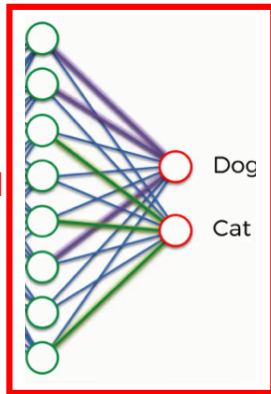


Linear Classifier

Linear Classifiers and Neural Networks



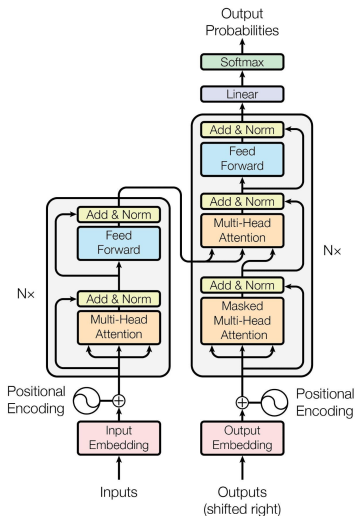
**Handcrafted
Features**



Linear Classifier

Linear Classifiers and Generative AI

- **Transformers:** 99% of LLMs/GenAI
 - ChatGPT; GPT
 - Claude
 - Gemini
 - Llama
 - DeepSeek
 - xAI/Grok
- Last layer = linear classifier
- Last layer predicts next word/token
- I.e., last layer is a classifier!
- Also: Many other linear layers!



Let's Start Simple

- Example 1 – sequence: ★ ◇ ○; label: ♣
- Example 2 – sequence: ★ ♥ △; label: ♣
- Example 3 – sequence: ★ △ ♠; label: □
- Example 4 – sequence: ◇ △ ○; label: □

Let's Start Simple

- Example 1 – sequence: ★ ◇ ○; label: ♣
- Example 2 – sequence: ★ ♥ △; label: ♣
- Example 3 – sequence: ★ △ ♠; label: □
- Example 4 – sequence: ◇ △ ○; label: □
- New sequence: ★ ◇ ○; label ?

Let's Start Simple

- Example 1 – sequence: ★ ◇ ○; label: ♣
- Example 2 – sequence: ★ ♥ △; label: ♣
- Example 3 – sequence: ★ △ ♠; label: □
- Example 4 – sequence: ◇ △ ○; label: □
- New sequence: ★ ◇ ○; label ♣
- New sequence: ★ ◇ ♥; label ?

Let's Start Simple

- Example 1 – sequence: ★ ◇ ○; label: ♣
- Example 2 – sequence: ★ ♥ △; label: ♣
- Example 3 – sequence: ★ △ ♠; label: □
- Example 4 – sequence: ◇ △ ○; label: □

- New sequence: ★ ◇ ○; label ♣
- New sequence: ★ ◇ ♥; label ♣
- New sequence: ★ △ ○; label ?

Let's Start Simple

- Example 1 – sequence: ★ ◇ ○; label: ♣
- Example 2 – sequence: ★ ♥ △; label: ♣
- Example 3 – sequence: ★ △ ♠; label: □
- Example 4 – sequence: ◇ △ ○; label: □

- New sequence: ★ ◇ ○; label ♣
- New sequence: ★ ◇ ♥; label ♣
- New sequence: ★ △ ○; label ?

Why can we do this?

Let's Start Simple: Machine Learning

- Example 1 – sequence: $\star \diamond \circ$; label: $\clubsuit$
- Example 2 – sequence: $\star \heartsuit \triangle$; label: $\clubsuit$
- Example 3 – sequence: $\star \triangle \spadesuit$; label: $\square$
- Example 4 – sequence: $\diamond \triangle \circ$; label: $\square$
- New sequence: $\star \diamond \heartsuit$; label $\clubsuit$

Label $\clubsuit$

Label $\square$

$$P(\clubsuit|\star) = \frac{\text{count}(\star \text{ and } \clubsuit)}{\text{count}(\star)} = \frac{2}{3} = 0.67 \text{ vs. } P(\square|\star) = \frac{\text{count}(\star \text{ and } \square)}{\text{count}(\star)} = \frac{1}{3} = 0.33$$

$$P(\clubsuit|\diamond) = \frac{\text{count}(\diamond \text{ and } \clubsuit)}{\text{count}(\diamond)} = \frac{1}{2} = 0.5 \text{ vs. } P(\square|\diamond) = \frac{\text{count}(\diamond \text{ and } \square)}{\text{count}(\diamond)} = \frac{1}{2} = 0.5$$

$$P(\clubsuit|\heartsuit) = \frac{\text{count}(\heartsuit \text{ and } \clubsuit)}{\text{count}(\heartsuit)} = \frac{1}{1} = 1.0 \text{ vs. } P(\square|\heartsuit) = \frac{\text{count}(\heartsuit \text{ and } \square)}{\text{count}(\heartsuit)} = \frac{0}{1} = 0.0$$

Let's Start Simple: Machine Learning

- Example 1 – sequence: $\star \diamond \circ$; label: $\clubsuit$
- Example 2 – sequence: $\star \heartsuit \triangle$; label: $\clubsuit$
- Example 3 – sequence: $\star \triangle \spadesuit$; label: $\square$
- Example 4 – sequence: $\diamond \triangle \circ$; label: $\square$
- New sequence: $\star \triangle \circ$; label ?

Label $\clubsuit$

Label $\square$

$$P(\clubsuit|\star) = \frac{\text{count}(\star \text{ and } \clubsuit)}{\text{count}(\star)} = \frac{2}{3} = 0.67 \text{ vs. } P(\square|\star) = \frac{\text{count}(\star \text{ and } \square)}{\text{count}(\star)} = \frac{1}{3} = 0.33$$

$$P(\clubsuit|\triangle) = \frac{\text{count}(\triangle \text{ and } \clubsuit)}{\text{count}(\triangle)} = \frac{1}{3} = 0.33 \text{ vs. } P(\square|\triangle) = \frac{\text{count}(\triangle \text{ and } \square)}{\text{count}(\triangle)} = \frac{2}{3} = 0.67$$

$$P(\clubsuit|\circ) = \frac{\text{count}(\circ \text{ and } \clubsuit)}{\text{count}(\circ)} = \frac{1}{2} = 0.5 \text{ vs. } P(\square|\circ) = \frac{\text{count}(\circ \text{ and } \square)}{\text{count}(\circ)} = \frac{1}{2} = 0.5$$

Let's Start Simple – Language Models

Same data, but now treat the label as just the **next symbol** (a “word”) and move it to the **rightmost token** of the sequence:

- Example 1 – sequence: ★ ◇ ○ ♣
- Example 2 – sequence: ★ ♥ △ ♣
- Example 3 – sequence: ★ △ ♠ □
- Example 4 – sequence: ◇ △ ○ □

“Predict the label” is now “**predict the next symbol**”:

Let's Start Simple – Language Models

Same data, but now treat the label as just the **next symbol** (a “word”) and move it to the **rightmost token** of the sequence:

- Example 1 – sequence: ★ ◇ ○ ♣
- Example 2 – sequence: ★ ♥ △ ♣
- Example 3 – sequence: ★ △ ♠ □
- Example 4 – sequence: ◇ △ ○ □

“Predict the label” is now “**predict the next symbol**”:

- New sequence: ★ ◇ ○ ?

Let's Start Simple – Language Models

Same data, but now treat the label as just the **next symbol** (a “word”) and move it to the **rightmost token** of the sequence:

- Example 1 – sequence: ★ ◇ ○ ♣
- Example 2 – sequence: ★ ♥ △ ♣
- Example 3 – sequence: ★ △ ♠ □
- Example 4 – sequence: ◇ △ ○ □

“Predict the label” is now “**predict the next symbol**”:

- New sequence: ★ ◇ ○ ♣
- New sequence: ★ ◇ ♥ ?

Let's Start Simple – Language Models

Same data, but now treat the label as just the **next symbol** (a “word”) and move it to the **rightmost token** of the sequence:

- Example 1 – sequence: ★ ◇ ○ ♣
- Example 2 – sequence: ★ ♥ △ ♣
- Example 3 – sequence: ★ △ ♠ □
- Example 4 – sequence: ◇ △ ○ □

“Predict the label” is now “**predict the next symbol**”:

- New sequence: ★ ◇ ○ ♣
- New sequence: ★ ◇ ♥ ♣
- New sequence: ★ △ ○ ?

Let's Start Simple – Language Models

Same data, but now treat the label as just the **next symbol** (a “word”) and move it to the **rightmost token** of the sequence:

- Example 1 – sequence: ★ ◇ ○ ♣
- Example 2 – sequence: ★ ♥ △ ♣
- Example 3 – sequence: ★ △ ♠ □
- Example 4 – sequence: ◇ △ ○ □

“Predict the label” is now “**predict the next symbol**”:

- New sequence: ★ ◇ ○ ♣
- New sequence: ★ ◇ ♥ ♣
- New sequence: ★ △ ○ ?

Same three examples – it's the same problem!

Machine Learning

- 1 Define a model/distribution of interest
- 2 Make some assumptions if needed
- 3 Fit the model to the data

Outline

① Terminology, notation and feature representations

② Linear Classifiers

Perceptron

Logistic Regression

③ Regularization

④ Dense Representations

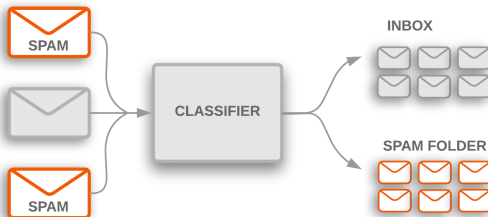
⑤ Similarity-based Learning

⑥ Neural Networks

Binary Classification: Spam Detection

Task: Identify if an incoming email/SMS/DM/etc. is spam or not.

This is a **binary classification problem**.



Multiclass Classification: Topic Labeling

Task: given a news article, determine its topic (politics, sports, etc.)

This is a **multi-class classification problem**.



sports
politics
technology
economy
weather
culture

Inputs and Outputs

- Input $x \in \mathcal{X}$
 - e.g., a news article, a sentence, an image, ...
- Output $y \in \mathcal{Y}$
 - e.g., spam/not spam, a topic, a translation, a next word
- Input/Output pair: $(x, y) \in \mathcal{X} \times \mathcal{Y}$
 - e.g., a **news article** together with a **topic**
 - e.g., a **email** together with a **spam/no spam label**
 - e.g., a **sentence prefix** together with a **the next word/token**

(Supervised) Machine Learning

- We are given a **labeled dataset** of input/output pairs:

$$\mathcal{D} = \{(\mathbf{x}_t, \mathbf{y}_t)\}_{t=1}^{|\mathcal{X}|} \subseteq \mathcal{X} \times \mathcal{Y}$$

- **Goal:** use it to learn a **classifier** $h : \mathcal{X} \rightarrow \mathcal{Y}$ that generalizes well to arbitrary inputs.
- At test time, given $\mathbf{x}_t \in \mathcal{X}$, we predict

$$\mathbf{y}' = h(\mathbf{x}_t).$$

- Hopefully, $\mathbf{y}' \approx \mathbf{y}_t$ most of the time.

Things can go by different names depending on what y is...

Regression

Deals with **continuous** output variables:

- **Regression:** $\mathcal{Y} = \mathbb{R}$
 - e.g., given a news article, how much time a user will spend reading it?
- **Multivariate regression:** $\mathcal{Y} = \mathbb{R}^K$, where $K > 1$
 - e.g., predict the X-Y coordinates in an image where the user will click

Classification

Deals with **discrete** output variables:

- **Binary classification:** $\mathcal{Y} = \{\pm 1\}$
 - e.g., spam detection, positive/negative sentiment
- **Multi-class classification:** $\mathcal{Y} = \{1, 2, \dots, K\}$
 - e.g., topic classification, positive/negative/neutral sentiment
- **Structured classification:** $\mathcal{Y}$ exponentially large and structured
 - e.g., machine translation, caption generation, image segmentation

What about GenerativeAI?

Our Setup

Let's assume a multi-class classification problem, with $|\mathcal{Y}|$ labels (classes).

Input x

What is x ?

x must be in a compute-friendly (numeric) form

x must be useful for the task

Feature Representations

Typical approach: define a feature map $\phi : \mathcal{X} \rightarrow \mathbb{R}^D$

- $\phi(x)$ is a high dimensional **feature vector** (always a **column** vector)

We can use feature vectors to encapsulate **Boolean**, **categorical**, and **continuous** features

- To start, we will focus on **sparse binary features**
- Categorical features can be reduced to a range of one-hot binary values
- We look at continuous (dense) features later

Feature Representations

Feature engineering is an important step in linear classifiers:

- Bag-of-words features for text, also lemmas, parts-of-speech, ...
- Embeddings (e.g., word2vec)
- SIFT features and wavelet representations in computer vision
- External database, APIs and knowledge resources

Later, we will talk about how we can learn these representations

Bag of Words Feature Representation

- x is a (noun) phrase

$$\phi_0(x) = \begin{cases} 1 & \text{if } x \text{ contains "George"} \\ 0 & \text{otherwise} \end{cases}$$

$$\phi_1(x) = \begin{cases} 1 & \text{if } x \text{ contains "Washington"} \\ 0 & \text{otherwise} \end{cases}$$

$$\phi_2(x) = \begin{cases} 1 & \text{if } x \text{ contains "Bridge"} \\ 0 & \text{otherwise} \end{cases}$$

$$\phi_3(x) = \begin{cases} 1 & \text{if } x \text{ contains "General"} \\ 0 & \text{otherwise} \end{cases}$$

$$\phi_4(x) = \begin{cases} 1 & \text{if } x \text{ contains an unknown word} \\ 0 & \text{otherwise} \end{cases}$$

- x =General George Washington $\rightarrow \phi(x) = [1 \ 1 \ 0 \ 1 \ 0]^\top$
- x =George Washington Bridge $\rightarrow \phi(x) = [1 \ 1 \ 1 \ 0 \ 0]^\top$
- x =George Washington University $\rightarrow \phi(x) = [1 \ 1 \ 0 \ 0 \ 1]^\top$
- x =George George George of the Jungle $\rightarrow \phi(x) = [1 \ 0 \ 0 \ 0 \ 1]^\top$

Feature Engineering and NLP Pipelines

Classical NLP pipelines consist of stacking together several linear classifiers
Each classifier's predictions are used to handcraft features for other classifiers

Example: Part-of-speech → Named Entities → Topic Classification

- **Part-of-speech**: nouns, determiners for Typed Named Entities
 - E.g., Google noun vs. Google verb
- **Typed Named Entities**: Categories for topic classification
 - E.g., Which George Washington? Person, University/Organization, Bridge/Location?

Outline

① Terminology, notation and feature representations

② **Linear Classifiers**

Perceptron

Logistic Regression

③ Regularization

④ Dense Representations

⑤ Similarity-based Learning

⑥ Neural Networks

A Quick **Matrix** Primer (1/3): Vectors and Matrices

- A **vector** is a column of numbers; $\mathbf{u} \in \mathbb{R}^3$ means 3 rows, 1 column

$$\mathbf{u} = \begin{bmatrix} 1 \\ 2 \\ 3 \end{bmatrix} \quad (3 \times 1) \quad \mathbf{u}^\top = \begin{bmatrix} 1 & 2 & 3 \end{bmatrix} \quad (1 \times 3)$$

- The **transpose** $^\top$ turns a column into a row (and back)
- A **matrix** $\mathbf{A} \in \mathbb{R}^{2 \times 3}$ has 2 rows and 3 columns; $\mathbf{A}_{i,j}$ is row i , column j

$$\mathbf{A} = \begin{bmatrix} 1 & 2 & 3 \\ 4 & 5 & 6 \end{bmatrix} \quad \mathbf{A}_{2,3} = 6$$

- Addition** (same shape only) and **scaling** work entry by entry

$$\begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} + \begin{bmatrix} 5 & 6 \\ 7 & 8 \end{bmatrix} = \begin{bmatrix} 6 & 8 \\ 10 & 12 \end{bmatrix} \quad 2 \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} = \begin{bmatrix} 2 & 4 \\ 6 & 8 \end{bmatrix}$$

A Quick **Matrix** Primer (2/3): Inner and Outer Products

To multiply, the **inner dimensions must match**: $(m \times n)(n \times p) = (m \times p)$

- **Inner product** (row $\times$ column) gives a single **number**

$$\underbrace{\begin{bmatrix} 1 & 2 & 3 \end{bmatrix}}_{1 \times 3} \underbrace{\begin{bmatrix} 4 \\ 5 \\ 6 \end{bmatrix}}_{3 \times 1} = (1)(4) + (2)(5) + (3)(6) = 32$$

- **Outer product** (column $\times$ row) gives a **matrix**

$$\underbrace{\begin{bmatrix} 1 \\ 2 \end{bmatrix}}_{2 \times 1} \underbrace{\begin{bmatrix} 3 & 4 \end{bmatrix}}_{1 \times 2} = \begin{bmatrix} 3 & 4 \\ 6 & 8 \end{bmatrix} \quad (2 \times 2)$$

A Quick **Matrix** Primer (3/3): Matrix $\times$ Vector

Each entry of the result is the inner product of one **row** of the matrix with the vector:

$$\underbrace{\begin{bmatrix} 1 & 2 \\ 3 & 4 \\ 5 & 6 \end{bmatrix}}_{3 \times 2} \underbrace{\begin{bmatrix} 1 \\ 2 \end{bmatrix}}_{2 \times 1} = \begin{bmatrix} (1)(1) + (2)(2) \\ (3)(1) + (4)(2) \\ (5)(1) + (6)(2) \end{bmatrix} = \underbrace{\begin{bmatrix} 5 \\ 11 \\ 17 \end{bmatrix}}_{3 \times 1}$$

- The **inner** dimensions must match – here both are 2
 - each row of the matrix has 2 entries, and so does the vector
- The **outer** dimensions give the shape of the result
 - $(3 \times 2)(2 \times 1) = (3 \times 1)$
- **One row in, one number out**: a matrix with 3 rows gives 3 numbers

Our Setup

Reminder:

Multi-class classification problem, with $|\mathcal{Y}|$ labels (classes).

Input x , represented by a feature vector $\phi(x) \in \mathbb{R}^D$ (a **column** vector).

Model Parameter Notation – Weights

- Weight matrix $\mathbf{W} \in \mathbb{R}^{|\mathcal{Y}| \times D}$
- Row $\mathbf{y}$ of $\mathbf{W}$ is written $\mathbf{w}_y^\top$, with $\mathbf{w}_y \in \mathbb{R}^D$ the weight vector of label y

$$\mathbf{W} = \begin{bmatrix} \mathbf{w}_1^\top \\ \vdots \\ \mathbf{w}_{|\mathcal{Y}|}^\top \end{bmatrix}$$

- one row per label in $\mathcal{Y}$ and one column per dimension in $\phi(\mathbf{x}) \in \mathbb{R}^D$

Intuition: $\mathbf{W}_{y,i}$ (entry i of $\mathbf{w}_y$) says how much feature i counts as *evidence* for label y – positive weights push the score of y up, negative weights push it down.

Model Parameter Notation – Bias

- Bias vector $\mathbf{b} \in \mathbb{R}^{|\mathcal{Y}|}$: one bias per label, b_y is entry y of $\mathbf{b}$

$$\mathbf{b} = \begin{bmatrix} b_1 \\ \vdots \\ b_{|\mathcal{Y}|} \end{bmatrix}$$

Intuition: b_y is the score label y gets *before looking at any features* – it lets the model prefer labels that are simply more common.

Linear Classifiers: Scores and Predictions

- The **score** (a.k.a. **logit**) of a single label y :

$$z_y = \mathbf{w}_y^\top \phi(x) + b_y$$

$$\mathbf{w}_y^\top \phi(x) = \sum_{i=1}^D \mathbf{w}_{y,i} \phi_i(x) = \mathbf{w}_{y,1} \phi_1(x) + \mathbf{w}_{y,2} \phi_2(x) + \cdots + \mathbf{w}_{y,D} \phi_D(x)$$

- All labels at once – one matrix-vector product:

$$\underbrace{\mathbf{z}}_{|Y| \times 1} = \underbrace{\mathbf{W}}_{|Y| \times D} \underbrace{\phi(x)}_{D \times 1} + \underbrace{\mathbf{b}}_{|Y| \times 1}$$

- Prediction: pick the label with the highest score

$$y' = \operatorname{argmax}_{y \in Y} z_y = \operatorname{argmax} (\mathbf{W} \phi(x) + \mathbf{b})$$

Linear Classifiers as **Matrix Multiplication**

Let $|\mathcal{Y}| = 3$, $D = 3$, $\mathbf{w}_y = [w_1^y, w_2^y, w_3^y]^\top$, and write ϕ_i for $\phi_i(x)$

$$\begin{aligned}\mathbf{W}\phi(x) + \mathbf{b} &= \begin{bmatrix} w_1^1 & w_2^1 & w_3^1 \\ w_1^2 & w_2^2 & w_3^2 \\ w_1^3 & w_2^3 & w_3^3 \end{bmatrix} \begin{bmatrix} \phi_1 \\ \phi_2 \\ \phi_3 \end{bmatrix} + \begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix} \\ &= \begin{bmatrix} (w_1^1 \times \phi_1) + (w_2^1 \times \phi_2) + (w_3^1 \times \phi_3) + b_1 \\ (w_1^2 \times \phi_1) + (w_2^2 \times \phi_2) + (w_3^2 \times \phi_3) + b_2 \\ (w_1^3 \times \phi_1) + (w_2^3 \times \phi_2) + (w_3^3 \times \phi_3) + b_3 \end{bmatrix} \\ &= \begin{bmatrix} \mathbf{w}_1^\top \phi(x) + b_1 \\ \mathbf{w}_2^\top \phi(x) + b_2 \\ \mathbf{w}_3^\top \phi(x) + b_3 \end{bmatrix} = \begin{bmatrix} z_1 \\ z_2 \\ z_3 \end{bmatrix} = \mathbf{z}\end{aligned}$$

One row of **W** = one label = one score.

Linear Classifiers – Example

- $D = 5$, $\mathcal{Y} = \{\text{Location (loc)}, \text{Person (per)}\}$, so $\mathbf{W} \in \mathbb{R}^{2 \times 5}$
- $\mathbf{x} = \text{George Washington Bridge} \rightarrow \phi(\mathbf{x}) = [1, 1, 1, 0, 0]^\top$

$$\mathbf{z} = \mathbf{W}\phi(\mathbf{x}) + \mathbf{b}$$

$$= \begin{bmatrix} -0.6 & 2.4 & 4.0 & -2.1 & 0.1 \\ 0.3 & 1.2 & -5.4 & 3.8 & -0.09 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ 1 \\ 0 \\ 0 \end{bmatrix} + \begin{bmatrix} 0 \\ 0 \end{bmatrix}$$

$$= \begin{bmatrix} (-0.6 \times 1) + (2.4 \times 1) + (4.0 \times 1) + (-2.1 \times 0) + (0.1 \times 0) + 0 \\ (0.3 \times 1) + (1.2 \times 1) + (-5.4 \times 1) + (3.8 \times 0) + (-0.09 \times 0) + 0 \end{bmatrix}$$

$$= \begin{bmatrix} \mathbf{5.8} \\ -3.9 \end{bmatrix} \begin{array}{l} \text{loc} \\ \text{per} \end{array}$$

$$\mathbf{y}' = \operatorname{argmax}(\mathbf{W}\phi(\mathbf{x}) + \mathbf{b}) = \operatorname{argmax}\{5.8_{\text{loc}}, -3.9_{\text{per}}\} = \text{loc}$$

Model Parameters Θ

$$\Theta = (\mathbf{W}, b)$$

- Each possible Θ is an instance of a model
- Where does Θ come from?
- We will **learn** Θ from data
- Different ways of learning it give different linear classifiers
 - perceptron, logistic regression, SVMs, ...

Binary Linear Classifier

With **binary labels** ($\mathcal{Y} = \{\pm 1\}$), $\mathbf{W}$ has just two rows, $\mathbf{w}_{+1}^\top$ and $\mathbf{w}_{-1}^\top$. We often use a minimal parametrization:

$$y' = \arg \max_{y \in \{\pm 1\}} \mathbf{w}_y^\top \phi(x) + b_y$$

Binary Linear Classifier

With **binary labels** ($\mathcal{Y} = \{\pm 1\}$), $\mathbf{W}$ has just two rows, $\mathbf{w}_{+1}^\top$ and $\mathbf{w}_{-1}^\top$. We often use a minimal parametrization:

$$\begin{aligned} y' &= \arg \max_{y \in \{\pm 1\}} \mathbf{w}_y^\top \phi(\mathbf{x}) + b_y \\ &= \begin{cases} +1 & \text{if } \mathbf{w}_{+1}^\top \phi(\mathbf{x}) + b_{+1} > \mathbf{w}_{-1}^\top \phi(\mathbf{x}) + b_{-1} \\ -1 & \text{otherwise} \end{cases} \end{aligned}$$

Binary Linear Classifier

With **binary labels** ($\mathcal{Y} = \{\pm 1\}$), $\mathbf{W}$ has just two rows, $\mathbf{w}_{+1}^\top$ and $\mathbf{w}_{-1}^\top$. We often use a minimal parametrization:

$$\begin{aligned} y' &= \arg \max_{y \in \{\pm 1\}} \mathbf{w}_y^\top \phi(x) + b_y \\ &= \begin{cases} +1 & \text{if } \mathbf{w}_{+1}^\top \phi(x) + b_{+1} > \mathbf{w}_{-1}^\top \phi(x) + b_{-1} \\ -1 & \text{otherwise} \end{cases} \\ &= \begin{cases} +1 & \text{if } (\mathbf{w}_{+1} - \mathbf{w}_{-1})^\top \phi(x) + (b_{+1} - b_{-1}) > 0 \\ -1 & \text{otherwise} \end{cases} \end{aligned}$$

Binary Linear Classifier

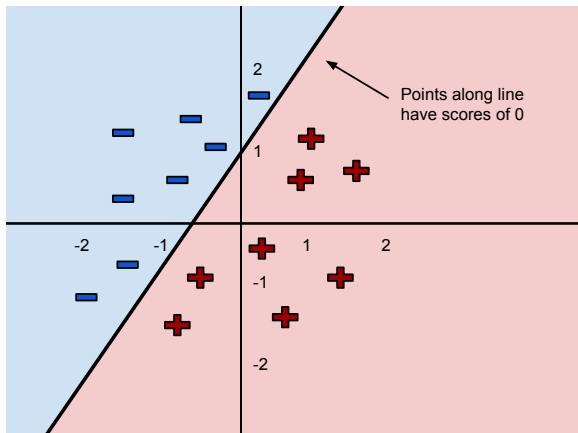
With **binary labels** ($\mathcal{Y} = \{\pm 1\}$), $\mathbf{W}$ has just two rows, $\mathbf{w}_{+1}^\top$ and $\mathbf{w}_{-1}^\top$. We often use a minimal parametrization:

$$\begin{aligned} y' &= \arg \max_{y \in \{\pm 1\}} \mathbf{w}_y^\top \phi(x) + b_y \\ &= \begin{cases} +1 & \text{if } \mathbf{w}_{+1}^\top \phi(x) + b_{+1} > \mathbf{w}_{-1}^\top \phi(x) + b_{-1} \\ -1 & \text{otherwise} \end{cases} \\ &= \begin{cases} +1 & \text{if } (\mathbf{w}_{+1} - \mathbf{w}_{-1})^\top \phi(x) + (b_{+1} - b_{-1}) > 0 \\ -1 & \text{otherwise} \end{cases} \\ &= \text{sign}(\underbrace{(\mathbf{w}_{+1} - \mathbf{w}_{-1})^\top}_{\mathbf{v}} \phi(x) + \underbrace{(b_{+1} - b_{-1})}_{c}). \end{aligned}$$

That is: only half of the parameters are needed.

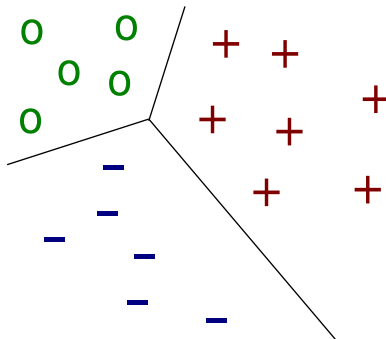
Binary Linear Classifier

Then $(\mathbf{v}, c)$ is an hyperplane that divides all points:



Multiclass Linear Classifier

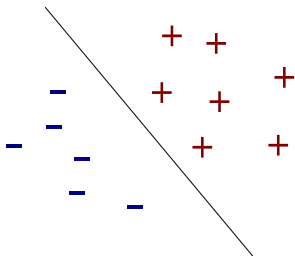
Defines regions of space.



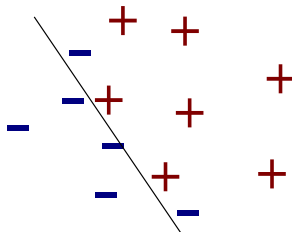
Linear Separability

- A set of points is **linearly separable** if there exists a $(\mathbf{W}, \mathbf{b})$ such that classification is perfect

Separable



Not Separable



Learning

- Machine Learning = finding the parameters $\Theta = (\mathbf{W}, \mathbf{b})$
- Using data! Specifically $\mathcal{D} = \{\mathbf{x}_t, \mathbf{y}_t\}_{t=1}$
- There are many algorithms for doing this

Outline

① Terminology, notation and feature representations

② **Linear Classifiers**

Perceptron

Logistic Regression

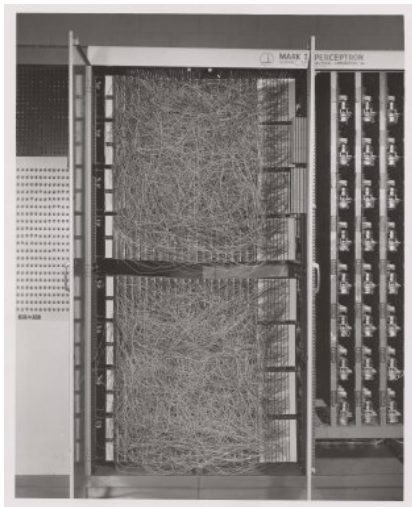
③ Regularization

④ Dense Representations

⑤ Similarity-based Learning

⑥ Neural Networks

Perceptron (Rosenblatt, 1958)



(Extracted from Wikipedia)

- Invented in 1957 at the Cornell Aeronautical Laboratory by Frank Rosenblatt
- Implemented in custom-built hardware as the “Mark 1 perceptron,” designed for image recognition
- 400 photocells, randomly connected to the “neurons.” Weights were encoded in potentiometers
- Weight updates during learning were performed by electric motors.

Perceptron in the News...

NEW NAVY DEVICE LEARNS BY DOING

Psychologist Shows Embryo
of Computer Designed to
Read and Grow Wiser

WASHINGTON, July 7 (UPI)—The Navy revealed the embryo of an electronic computer today that it expects will be able to walk, talk, see, write, reproduce itself and be conscious of its existence.

The embryo—the Weather Bureau's \$2,000,000 "704" computer—learned to differentiate between right and left after fifty attempts in the Navy's demonstration for newsmen.

The service said it would use this principle to build the first of its Perceptron thinking machines that will be able to read and write. It is expected to be finished in about a year at a cost of \$100,000.

Dr. Frank Rosenblatt, designer of the Perceptron, conducted the demonstration. He said the machine would be the first device to think as the human brain. As do human be-

ings, Perceptron will make mistakes at first, but will grow wiser as it gains experience, he said.

Dr. Rosenblatt, a research psychologist at the Cornell Aeronautical Laboratory, Buffalo, said Perceptrons might be fired to the planets as mechanical space explorers.

Without Human Controls

The Navy said the perceptron would be the first non-living mechanism "capable of receiving, recognizing and identifying its surroundings without any human training or control."

The "brain" is designed to remember images and information it has perceived itself. Ordinary computers remember only what is fed into them on punch cards or magnetic tape.

Later Perceptrons will be able to recognize people and call out their names and instantly translate speech in one language to speech or writing in another language, it was predicted.

Mr. Rosenblatt said in principle it would be possible to build brains that could reproduce themselves on an assembly line and which would be conscious of their existence.

1958 New York Times...

In today's demonstration, the "704" was fed two cards, one with squares marked on the left side and the other with squares on the right side.

Learns by Doing

In the first fifty trials, the machine made no distinction between them. It then started registering a "Q" for the left squares and "O" for the right squares.

Dr. Rosenblatt said he could explain why the machine learned only in highly technical terms. But he said the computer had undergone a "self-induced change in the wiring diagram."

The first Perceptron will have about 1,000 electronic "association cells" receiving electrical impulses from an eye-like scanning device with 400 photo-cells. The human brain has 10,000,000,000 responsive cells, including 100,000,000 connections with the eyes.

Perceptron Algorithm

- **Online** algorithm: process one data point (x_t, y_t) at time
 - Take x_t ; apply the current model to make a prediction for it

$$y' = \operatorname{argmax}_{y \in \mathcal{Y}} z_y = \operatorname{argmax} (\mathbf{W}\phi(x_t) + \mathbf{b})$$

- If prediction is **correct**, proceed
- **Else**, correct model: add the feature vector to the **row of the correct output** & subtract it from the **row of the predicted (wrong) output**

$$w_{y_t} = w_{y_t} + \phi(x_t) \quad w_{y'} = w_{y'} - \phi(x_t)$$

- Same for the corresponding biases

$$b_{y_t} = b_{y_t} + 1 \quad b_{y'} = b_{y'} - 1$$

Perceptron Update as a **Matrix** Update

- Let $\mathbf{e}_y \in \mathbb{R}^{|\mathcal{Y}|}$ be the **one-hot** vector of label y : a 1 in position y , zeros elsewhere
- Add $\phi(\mathbf{x}_t)$ to row $\mathbf{y}_t$, subtract it from row $\mathbf{y}'$ is exactly

$$\mathbf{W} = \mathbf{W} + (\mathbf{e}_{y_t} - \mathbf{e}_{y'}) \phi(\mathbf{x}_t)^\top, \quad \mathbf{b} = \mathbf{b} + (\mathbf{e}_{y_t} - \mathbf{e}_{y'})$$

- $(\mathbf{e}_{y_t} - \mathbf{e}_{y'}) \phi(\mathbf{x}_t)^\top$ is an **outer product**: an $|\mathcal{Y}| \times D$ matrix
 - row $\mathbf{y}_t$ is $+\phi(\mathbf{x}_t)^\top$, row $\mathbf{y}'$ is $-\phi(\mathbf{x}_t)^\top$, all other rows are $\mathbf{0}$

$$(\mathbf{e}_{y_t} - \mathbf{e}_{y'}) \phi(\mathbf{x}_t)^\top = \begin{bmatrix} \mathbf{0}^\top \\ +\phi(\mathbf{x}_t)^\top \\ \mathbf{0}^\top \\ -\phi(\mathbf{x}_t)^\top \\ \mathbf{0}^\top \end{bmatrix} \begin{matrix} \text{row } \mathbf{y}_t \\ \text{row } \mathbf{y}' \end{matrix}$$

Remember this shape – every update we see has it.

Perceptron Update – A Small Example

$|\mathcal{Y}| = 3$ labels, $D = 2$ features. Current parameters and a training example:

$$\mathbf{W} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 0 & 0 \end{bmatrix} \begin{array}{l} \text{label 1} \\ \text{label 2} \\ \text{label 3} \end{array} \quad \mathbf{b} = \begin{bmatrix} 0 \\ 0 \\ 0 \end{bmatrix} \quad \phi(\mathbf{x}_t) = \begin{bmatrix} 2 \\ 1 \end{bmatrix}, \quad \mathbf{y}_t = 3$$

Predict:

$$\mathbf{z} = \mathbf{W}\phi(\mathbf{x}_t) + \mathbf{b} = \begin{bmatrix} 2 \\ 1 \\ 0 \end{bmatrix} \implies \mathbf{y}' = 1 \neq \mathbf{y}_t = 3 \quad \text{mistake!}$$

Update: add $\phi(\mathbf{x}_t)$ to row 3, subtract it from row 1

$$\left(\underbrace{\begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}}_{\mathbf{e}_3} - \underbrace{\begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix}}_{\mathbf{e}_1} \right) \underbrace{\begin{bmatrix} 2 & 1 \end{bmatrix}}_{\phi(\mathbf{x}_t)^\top} = \begin{bmatrix} -1 \\ 0 \\ 1 \end{bmatrix} \begin{bmatrix} 2 & 1 \end{bmatrix} = \begin{bmatrix} -2 & -1 \\ 0 & 0 \\ 2 & 1 \end{bmatrix}$$

Perceptron Update – A Small Example

Apply it to the weights and the biases:

$$\mathbf{W} \leftarrow \begin{bmatrix} 1 & 0 \\ 0 & 1 \\ 0 & 0 \end{bmatrix} + \begin{bmatrix} -2 & -1 \\ 0 & 0 \\ 2 & 1 \end{bmatrix} = \begin{bmatrix} -1 & -1 \\ 0 & 1 \\ 2 & 1 \end{bmatrix} \quad \mathbf{b} \leftarrow \begin{bmatrix} -1 \\ 0 \\ 1 \end{bmatrix}$$

Only the rows of the **true** and the **predicted** label changed. Score the same example again:

$$\mathbf{z} = \begin{bmatrix} -1 & -1 \\ 0 & 1 \\ 2 & 1 \end{bmatrix} \begin{bmatrix} 2 \\ 1 \end{bmatrix} + \begin{bmatrix} -1 \\ 0 \\ 1 \end{bmatrix} = \begin{bmatrix} -4 \\ 1 \\ 6 \end{bmatrix} \implies y' = 3 \text{ correct!}$$

The update pushes the true label's score up and the wrong one's down.

Perceptron Algorithm

```
input: labeled data  $\mathcal{D}$   
initialize  $\mathbf{W} = \mathbf{0}$  and  $\mathbf{b} = 0$   
repeat  
  observe example  $(\mathbf{x}_t, \mathbf{y}_t) \in \mathcal{D}$   
  predict  $\mathbf{y}' = \operatorname{argmax}(\mathbf{W}\phi(\mathbf{x}_t) + \mathbf{b})$   
  if  $\mathbf{y}' \neq \mathbf{y}_t$  then  
    update  $\mathbf{W} = \mathbf{W} + (\mathbf{e}_{\mathbf{y}_t} - \mathbf{e}_{\mathbf{y}'}) \phi(\mathbf{x}_t)^\top$   
    update  $\mathbf{b} = \mathbf{b} + (\mathbf{e}_{\mathbf{y}_t} - \mathbf{e}_{\mathbf{y}'})$   
  end if  
until stopping criterion2  
output: model parameters  $\Theta = (\mathbf{W}, \mathbf{b})$ 
```

²E.g., max iterations; zero/ ϵ errors

Perceptron's Mistake Bound

A couple definitions:

- the training data is **linearly separable** with margin $\gamma > 0$ iff there is a matrix $\mathbf{U}$ whose rows $\mathbf{u}_y^\top$ have $\|\mathbf{u}_y\| = 1$ such that

$$\mathbf{u}_{y_t}^\top \phi(\mathbf{x}_t) \geq \mathbf{u}_{y'}^\top \phi(\mathbf{x}_t) + \gamma, \quad \forall t, \forall y' \neq y_t.$$

- radius** of the data: $R = \max_t \|\phi(\mathbf{x}_t)\|$.

Perceptron's Mistake Bound

A couple definitions:

- the training data is **linearly separable** with margin $\gamma > 0$ iff there is a matrix $\mathbf{U}$ whose rows $\mathbf{u}_y^\top$ have $\|\mathbf{u}_y\| = 1$ such that

$$\mathbf{u}_{y_t}^\top \phi(\mathbf{x}_t) \geq \mathbf{u}_{y'}^\top \phi(\mathbf{x}_t) + \gamma, \quad \forall t, \forall y' \neq y_t.$$

- radius** of the data: $R = \max_t \|\phi(\mathbf{x}_t)\|$.

Then we have the following bound of the **number of mistakes**:

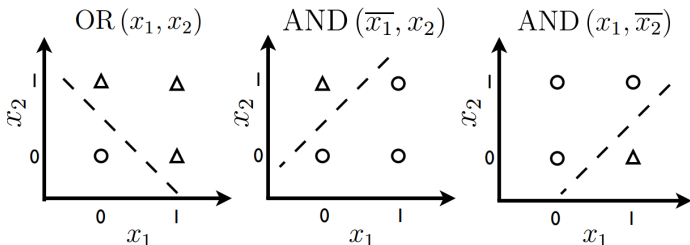
Theorem (Novikoff, 1962)

The perceptron algorithm is guaranteed to find a separating hyperplane after at most $2 \frac{R^2}{\gamma^2}$ mistakes.

Proof: <https://proceedings.mlr.press/v97/beygelzimer19a/beygelzimer19a-suppl.pdf>

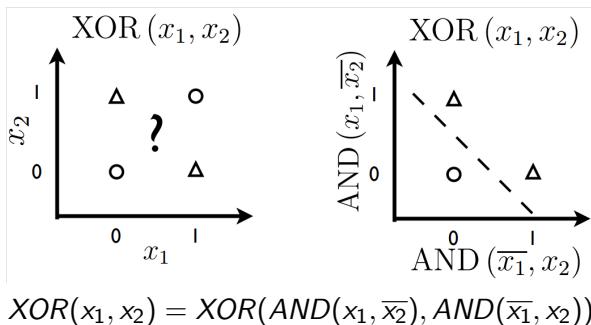
What a Simple Perceptron Can and Can't Do

- Remember: the decision boundary is linear (**linear classifier**)
- It **can** solve linearly separable problems (OR, AND)



What a Simple Perceptron Can and Can't Do

- ... but it **can't** solve **non-linearly separable** problems such as simple XOR (unless input is transformed into a better representation):



- Result attributed to Minsky and Papert (1969) but was known before.

Outline

① Terminology, notation and feature representations

② **Linear Classifiers**

Perceptron

Logistic Regression

③ Regularization

④ Dense Representations

⑤ Similarity-based Learning

⑥ Neural Networks

Logistic Regression

Turn the logits $\mathbf{z} = \mathbf{W}\phi(\mathbf{x}) + \mathbf{b}$ into a conditional probability:

$$P(\mathbf{y}|\mathbf{x}) = \frac{\exp(z_{\mathbf{y}})}{Z_{\mathbf{x}}}, \quad Z_{\mathbf{x}} = \sum_{\mathbf{y}' \in \mathcal{Y}} \exp(z_{\mathbf{y}'}), \quad \forall z_{\mathbf{y}} \in \mathbf{z}$$

This is called a **softmax** transformation

In vector form $\mathbf{p} = \mathbf{softmax}(\mathbf{z}) = [\frac{\exp(z_1)}{Z_{\mathbf{x}}}, \dots, \frac{\exp(z_{|\mathcal{Y}|})}{Z_{\mathbf{x}}}]^{\top}$

I.e., $\mathbf{p} = [p_1, \dots, p_{|\mathcal{Y}|}]^{\top}$, where $p_{\mathbf{y}} = P(\mathbf{y}|\mathbf{x}) = \frac{\exp(z_{\mathbf{y}})}{Z_{\mathbf{x}}}$

Probability distribution: $\sum_{\mathbf{y}} P(\mathbf{y}|\mathbf{x}) = 1$ and $P(\mathbf{y}|\mathbf{x}) \geq 0, \forall \mathbf{y}$

Logistic Regression

$$P_{\Theta}(\mathbf{y}|\mathbf{x}) = \frac{\exp(z_{\mathbf{y}})}{Z_{\mathbf{x}}}, \quad \mathbf{z} = \mathbf{W}\phi(\mathbf{x}) + \mathbf{b}$$

- The parameters are the matrix and the bias: $\Theta = (\mathbf{W}, \mathbf{b})$

How do we learn Θ ?

- Set Θ to minimize the negative **conditional log-likelihood**:

$$\begin{aligned} \Theta &= \arg \min_{\Theta} -\log \left(\prod_{t=1} P_{\Theta}(\mathbf{y}_t | \mathbf{x}_t) \right) = \arg \min_{\Theta} -\sum_{t=1} \log P_{\Theta}(\mathbf{y}_t | \mathbf{x}_t) \\ &= \arg \min_{\Theta} \sum_{t=1} \left(\log \sum_{\mathbf{y}'} \exp(z_{\mathbf{y}'}) - z_{\mathbf{y}_t} \right), \end{aligned}$$

i.e., assign as much probability mass as possible to the correct labels!

Logistic Regression

- This objective function is **convex**
- Therefore any local minimum is a global minimum
- No closed form solution, but lots of numerical techniques
 - Gradient methods (gradient descent, conjugate gradient)
 - Quasi-Newton methods (L-BFGS, ...)

Recap: Convex functions

Pro: Guarantee of a global minima ✓



Figure: Illustration of a convex function. The line segment between any two points on the graph lies entirely above the curve.

Recap: Gradients

A gradient of a function $f(\Theta)$ wrt parameters $\Theta = [w_1, \dots, w_P]$ is:

$$\nabla_{\Theta} f(\Theta) = \left[\frac{\partial}{\partial w_1} f, \dots, \frac{\partial}{\partial w_P} f \right]$$

I.e., the vector of partial derivatives of f , which is the derivative of f wrt to each variable w_i

Our parameters are a matrix $\mathbf{W}$ and a vector $\mathbf{b}$, so the gradient **has the same shape as the parameters**: $\nabla_{\mathbf{W}} f$ is $|\mathcal{Y}| \times D$ and $\nabla_{\mathbf{b}} f$ is $|\mathcal{Y}| \times 1^3$

Gradient = direction and fastest rate of increase of f at point Θ

When a gradient is zero we are at a stationary point of f . For convex functions that means global minima.

³Nothing changes if you stack all the entries of $\mathbf{W}$ and $\mathbf{b}$ into one long vector.

Recap: Iterative Descent Methods

Goal: find the minimum/minimizer of $f : \mathbb{R}^d \rightarrow \mathbb{R}$

- Proceed in **small steps** in the **optimal direction** till a **stopping criterion** is met (usually norm of gradient is small)
- **Gradient descent (GD)** updates: $\Theta \leftarrow \Theta - \eta \nabla_{\Theta} f(\Theta)$
- η is the step size / learning rate

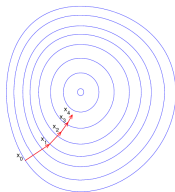


Figure: Illustration of gradient descent. The red lines correspond to steps taken in the negative gradient direction.

Logistic Regression: Gradient Descent (GD)

- Let $L(\Theta; (x, y)) = \left(\log \sum_{y'} \exp(z_{y'}) - z_y \right)$, with $\mathbf{z} = \mathbf{W}\phi(x) + \mathbf{b}$
- Call this our **loss function** for instance x, y
 - We want to minimize over $\mathcal{D} = \{(x_t, y_t)\}_{t=1}$ with GD
 - I.e., Find $\arg \min_{\Theta} \sum_{t=1} L(\Theta; (x_t, y_t))$
 - Logistic-regressions loss function often called **log-loss** or **cross-entropy**
- GD update will look like

$$\begin{aligned}\Theta &= \Theta - \eta \nabla_{\Theta} \left(\sum_{t=1} L(\Theta; (x_t, y_t)) \right) \\ &= \Theta - \eta \sum_{t=1} \nabla_{\Theta} L(\Theta; (x_t, y_t))\end{aligned}$$

- Need to calculate $\nabla_{\Theta} L(\Theta; (x, y))$: gradient of L w.r.t. Θ
- This is a **batch optimization**: updates are over whole dataset

Stochastic Gradient Descent (SGD)

SGD is like perceptron – update every instance:

- Pick $(\mathbf{x}_t, \mathbf{y}_t)$ randomly
- Update $\Theta = \Theta - \eta \nabla_{\Theta} L(\Theta; (\mathbf{x}_t, \mathbf{y}_t))$
- i.e. we approximate the true gradient with a noisy, unbiased, gradient, based on a single sample
- Variants exist in-between (mini-batches)
- GD and SGD guaranteed to find the optimal Θ (for suitable step sizes)

Logistic Regression: Simple SGD Algorithm

input: labeled data $\mathcal{D}$, step size η

initialize: $\mathbf{W} = \mathbf{0}$, $\mathbf{b} = \mathbf{0}$

repeat

observe example $(\mathbf{x}_t, \mathbf{y}_t) \in \mathcal{D}$

Update $\Theta = \Theta - \eta \nabla_{\Theta} L(\Theta; (\mathbf{x}_t, \mathbf{y}_t))$

until stopping criterion

output: model parameters $\Theta = (\mathbf{W}, \mathbf{b})$

- Picking step sizes example of **hyperparameter tuning**
- Stopping criterion usually gradient is small: $\|\nabla_{\Theta} L(\Theta; (\mathbf{x}_t, \mathbf{y}_t))\| < \epsilon, \forall t$
- Small (or zero) gradient is stationary point – global minimum

Computing the Gradient: $\nabla_{\Theta} L(\Theta; (x, y))$

- We need $\nabla_{\mathbf{W}} L$ and $\nabla_{\mathbf{b}} L$, where

$$L(\Theta; (x, y)) = \log \sum_{y'} \exp(z_{y'}) - z_y, \quad \mathbf{z} = \mathbf{W}\phi(x) + \mathbf{b}$$

- Θ only touches L through the logits $\mathbf{z}$, so use the chain rule in two steps:
 1. $\nabla_{\mathbf{z}} L$ (how does the loss change with the scores?)
 2. $\mathbf{z} = \mathbf{W}\phi(x) + \mathbf{b}$ (how do the scores change with $\mathbf{W}$ and $\mathbf{b}$?)
- Some reminders:
 - 1 $\frac{\partial}{\partial u} \log F(u) = \frac{1}{F(u)} \frac{\partial}{\partial u} F(u)$
 - 2 $\frac{\partial}{\partial u} \exp F(u) = \exp(F(u)) \frac{\partial}{\partial u} F(u)$

Step 1: Gradient w.r.t. the Logits z

Take the partial derivative w.r.t. one logit z_k :

$$\begin{aligned}\frac{\partial L}{\partial z_k} &= \frac{\partial}{\partial z_k} \left(\log \sum_{y'} \exp(z_{y'}) - z_y \right) \\&= \frac{1}{\sum_{y'} \exp(z_{y'})} \cdot \frac{\partial}{\partial z_k} \sum_{y'} \exp(z_{y'}) - \frac{\partial z_y}{\partial z_k} \\&= \frac{1}{Z_x} \sum_{y'} \exp(z_{y'}) \frac{\partial z_{y'}}{\partial z_k} - \frac{\partial z_y}{\partial z_k} \quad \frac{\partial z_{y'}}{\partial z_k} = \mathbb{I}[y' = k]: \text{ only one term survives} \\&= \frac{\exp(z_k)}{Z_x} - \mathbb{I}[k = y] = P_{\Theta}(k|x) - \mathbb{I}[k = y]\end{aligned}$$

Stacking these for all k gives one clean vector:

$$\nabla_z L = \mathbf{p} - \mathbf{e}_y \quad (\text{predicted distribution} - \text{the truth})$$

where $\mathbf{p} = \text{softmax}(\mathbf{z})$ and $\mathbf{e}_y$ is the one-hot vector of the correct label.

Step 2: From Logits to $\mathbf{W}$ and b

Each logit is $z_k = \mathbf{w}_k^\top \phi(\mathbf{x}) + b_k$, so

$$\frac{\partial z_k}{\partial \mathbf{W}_{k,i}} = \phi_i(\mathbf{x}), \quad \frac{\partial z_k}{\partial b_k} = 1$$

Chain rule, entry by entry:

$$\frac{\partial L}{\partial \mathbf{W}_{k,i}} = \underbrace{\frac{\partial L}{\partial z_k}}_{p_k - \mathbb{I}[k=y]} \cdot \underbrace{\frac{\partial z_k}{\partial \mathbf{W}_{k,i}}}_{\phi_i(\mathbf{x})} = (p_k - \mathbb{I}[k=y]) \phi_i(\mathbf{x})$$

Every entry at once is an **outer product**:

$$\nabla_{\mathbf{W}} L = (\mathbf{p} - \mathbf{e}_y) \phi(\mathbf{x})^\top, \quad \nabla_b L = \mathbf{p} - \mathbf{e}_y$$

- $\underbrace{(\mathbf{p} - \mathbf{e}_y)}_{|\mathcal{Y}| \times 1} \underbrace{\phi(\mathbf{x})^\top}_{1 \times D}$ is $|\mathcal{Y}| \times D$ – the same shape as $\mathbf{W}$, as it must be

What does the update look like?

A single SGD step on example (x, y) is one outer product:

$$\mathbf{W} = \mathbf{W} - \eta (\mathbf{p} - \mathbf{e}_y) \phi(x)^\top, \quad \mathbf{b} = \mathbf{b} - \eta (\mathbf{p} - \mathbf{e}_y)$$

Written out, the outer product is an $|\mathcal{Y}| \times D$ matrix – one row per label:

$$(\mathbf{p} - \mathbf{e}_y) \phi(x)^\top = \begin{bmatrix} p_1 \phi(x)^\top \\ \vdots \\ (p_y - 1) \phi(x)^\top \\ \vdots \\ p_{|\mathcal{Y}|} \phi(x)^\top \end{bmatrix} \quad \text{row } y: \text{ the true label}$$

- Row y : $p_y - 1 \leq 0$, so it moves **towards** $\phi(x)$
- Every other row: $p_{y'} \geq 0$, so it moves **away** from $\phi(x)$
- Unlike the perceptron, **every** row moves – each by its own probability

SGD for Logistic Regression

input: labeled data $\mathcal{D}$, step size η

initialize: $\mathbf{W} = \mathbf{0}$, $\mathbf{b} = \mathbf{0}$

repeat

observe example $(\mathbf{x}_t, \mathbf{y}_t) \in \mathcal{D}$

$\mathbf{p} = \text{softmax}(\mathbf{W}\phi(\mathbf{x}_t) + \mathbf{b})$

$\mathbf{W} = \mathbf{W} - \eta(\mathbf{p} - \mathbf{e}_{\mathbf{y}_t})\phi(\mathbf{x}_t)^\top$

$\mathbf{b} = \mathbf{b} - \eta(\mathbf{p} - \mathbf{e}_{\mathbf{y}_t})$

until stopping criterion

output: model parameters $\Theta = (\mathbf{W}, \mathbf{b})$

Logistic Regression Summary

- Define conditional probability

$$\mathbf{p} = \text{softmax}(\mathbf{W}\phi(\mathbf{x}) + \mathbf{b}), \quad P_{\Theta}(\mathbf{y}|\mathbf{x}) = p_{\mathbf{y}}$$

- Set parameters to minimize negative conditional log-likelihood:

$$\Theta = \arg \min_{\Theta} \sum_t -\log P_{\Theta}(\mathbf{y}_t|\mathbf{x}_t) = \arg \min_{\Theta} \sum_t L(\Theta; (\mathbf{x}_t, \mathbf{y}_t))$$

- Gradient is an outer product: $\nabla_{\mathbf{W}} L = (\mathbf{p} - \mathbf{e}_{\mathbf{y}}) \phi(\mathbf{x})^{\top}$
- Can run gradient descent (or any gradient-based optimization algorithm)

The Story So Far

- Logistic regression is **probabilistic**: minimizes **neg log-likelihood**
 - also called log-linear model and max-entropy classifier
 - no closed form solution
 - For training instance (x, y) , the SGD update is

$$\mathbf{W} = \mathbf{W} - \eta (\mathbf{p} - \mathbf{e}_y) \phi(x)^\top$$

- Perceptron is non-probabilistic; minimizes **training errors**
 - Assumes linearly separable, though works well anyways
 - For training instance (x, y) , the update is

$$\mathbf{W} = \mathbf{W} - (\mathbf{e}_{y'} - \mathbf{e}_y) \phi(x)^\top$$

Same outer product! The only difference: logistic regression uses the **soft** distribution $\mathbf{p}$ where the perceptron uses the **hard** prediction $\mathbf{e}_{y'}$.

Outline

① Terminology, notation and feature representations

② Linear Classifiers

Perceptron

Logistic Regression

③ Regularization

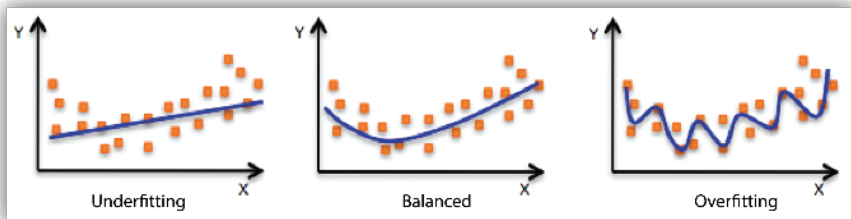
④ Dense Representations

⑤ Similarity-based Learning

⑥ Neural Networks

Overfitting

If the model is too complex (too many parameters) and the data is scarce, we run the risk of **overfitting**:



Regularization

In practice, we **regularize** models to prevent overfitting

$$\arg \min_{\Theta} \sum_{t=1} L(\Theta; (x_t, y_t)) + \lambda \Omega(\Theta),$$

where $\Omega(\Theta)$ is the regularization function, and λ controls how much to regularize.

- Gaussian prior (ℓ_2), promotes smaller weights⁴:

$$\Omega(\Theta) = \|\mathbf{W}\|_F^2 = \sum_{\mathbf{y}} \sum_i \mathbf{w}_{\mathbf{y},i}^2.$$

- Laplacian prior (ℓ_1), promotes **sparse** weights!

$$\Omega(\Theta) = \|\mathbf{W}\|_1 = \sum_{\mathbf{y}} \sum_i |\mathbf{w}_{\mathbf{y},i}|$$

⁴ $\|\cdot\|_F$ is the **Frobenius norm**: the sum of squares over both indices of the matrix.

Empirical Risk Minimization

$$\arg \min_{\Theta} \sum_{t=1} L(\Theta; (x_t, y_t)) + \lambda \Omega(\Theta),$$

- This formulation is generally called *Empirical Risk Minimization*.
- It consists of a Loss/Risk function that we want to minimize;
- And (optionally) a regularization function that stops us from overfitting to the data.
- Useful abstraction that covers most modern ML algorithms.
- Linear classifiers: different choices of L and Ω invoke specific models

Logistic Regression with ℓ_2 Regularization

- Still optimize with GD or SGD
- What is the new gradient?

$$\begin{aligned} & \sum_{t=1} \nabla_{\mathbf{W}} L(\Theta; (\mathbf{x}_t, \mathbf{y}_t)) + \nabla_{\mathbf{W}} \lambda \Omega(\Theta) \\ &= \sum_{t=1} \nabla_{\mathbf{W}} L(\Theta; (\mathbf{x}_t, \mathbf{y}_t)) + \nabla_{\mathbf{W}} \frac{\lambda}{2} \|\mathbf{W}\|_F^2 \end{aligned}$$

- We know $\nabla_{\mathbf{W}} L(\Theta; (\mathbf{x}_t, \mathbf{y}_t)) = (\mathbf{p} - \mathbf{e}_{y_t}) \phi(\mathbf{x}_t)^\top$
- Just need $\nabla_{\mathbf{W}} \frac{\lambda}{2} \|\mathbf{W}\|_F^2 = \lambda \mathbf{W}$ – again the same shape as $\mathbf{W}$

Loss Function

Should match as much as possible the metric we want to optimize at test time

Should be well-behaved (continuous, maybe smooth) to be amenable to optimization (this rules out the 0/1 loss)

Some examples:

- Squared loss for regression
- Negative log-likelihood (cross-entropy): multinomial logistic regression
- Hinge loss: support vector machines
- A bunch more ...

Could not possible cover everything.

Please look at Andre Martins excellent lecture for LXMLS:

- http://lxmls.it.pt/2019/LINEAR_LEARNERS.pdf
- Also covers
 - Naive Bayes
 - Support Vector Machines (SVMs)
 - SVMs v Perceptron v Log Reg
 - Non-Linear Classifiers $\neq$ Neural Networks
 - K-Nearest neighbors (we'll do this)
 - Kernel methods (SVMs)

Outline

① Terminology, notation and feature representations

② Linear Classifiers

Perceptron

Logistic Regression

③ Regularization

④ Dense Representations

⑤ Similarity-based Learning

⑥ Neural Networks

Recap: Sparse Feature Representations

- x is a name

$$\phi_0(x) = \begin{cases} 1 & \text{if } x \text{ contains "George"} \\ 0 & \text{otherwise} \end{cases}$$

$$\phi_1(x) = \begin{cases} 1 & \text{if } x \text{ contains "Washington"} \\ 0 & \text{otherwise} \end{cases}$$

$$\phi_2(x) = \begin{cases} 1 & \text{if } x \text{ contains "Bridge"} \\ 0 & \text{otherwise} \end{cases}$$

$$\phi_3(x) = \begin{cases} 1 & \text{if } x \text{ contains "General"} \\ 0 & \text{otherwise} \end{cases}$$

$$\phi_4(x) = \begin{cases} 1 & \text{if } x \text{ contains an unknown word} \\ 0 & \text{otherwise} \end{cases}$$

- x =General George Washington $\rightarrow \phi(x) = [1 \ 1 \ 0 \ 1 \ 0]^\top$
- x =George Washington Bridge $\rightarrow \phi(x) = [1 \ 1 \ 1 \ 0 \ 0]^\top$
- x =George Washington University $\rightarrow \phi(x) = [1 \ 1 \ 0 \ 0 \ 1]^\top$
- x =George George George of the Jungle $\rightarrow \phi(x) = [1 \ 0 \ 0 \ 0 \ 1]^\top$

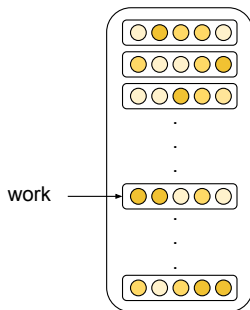
Dense Feature Representations

- $\phi(x) \in \mathbb{R}^D$
- But ϕ is dense real valued vector, i.e., zero values not frequent
- E.g.,

$$\phi(x) = [0.123, 0.439, -0.213, 0.692, -0.002]^\top$$

- But **where does $\phi(x)$ come from?**
- We learn it from data!
- Long history: tf-idf, vector space models, ..., **Word2Vec**, Glove, ...

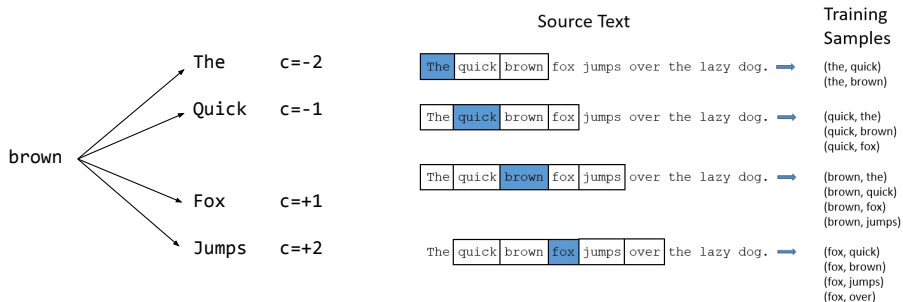
Embedding / Lookup Table



- Input is a word x with embedding $\phi(x) \in \mathbb{R}^D$, for all $x \in \mathcal{V}$
- $\mathcal{V}$ is a fixed vocabulary of words
- We store these in an **embedding matrix** $\mathbf{E} \in \mathbb{R}^{|\mathcal{V}| \times D}$ – **one row per word**
 - These are the model *word embeddings*
 - AKA embedding layer; word look-up table; ...
 - A lookup is just a matrix product: $\phi(x) = \mathbf{E}^\top \mathbf{e}_x$, with $\mathbf{e}_x$ the one-hot vector of x

Word2Vec

"You shall know a word by the company it keeps"



Example from McCormick <http://mccormickml.com/2016/04/19/word2vec-tutorial-the-skip-gram-model/>

Word2Vec

- Corpus $\mathcal{C} = \{\mathcal{X}_1, \dots, \mathcal{X}_{|\mathcal{C}|}\}$
- With sentences $\mathcal{X} = x_1, \dots, x_{|\mathcal{X}|}$
- Vocab $\mathcal{V} = \{x_i | x_i \in \mathcal{X} \text{ and } \mathcal{X} \in \mathcal{C}\}$
- Goal: learn vector/embedding for all $x_i \in \mathcal{V}$ (embedding table)
- **word2vec** (Mikolov et al., 2013)
 - Define vector/embedding per word: $\phi(x_i)$ ⁵
 - word2vec optimizes/maximizes (SkipGram model):

$$\sum_j \sum_i \sum_{-c \leq k \leq c, k \neq 0} \log p(x_{i+k} | x_i) = \sum_j \sum_i \sum_{-c \leq k \leq c, k \neq 0} \log \frac{e^{\phi(x_i)^\top \phi(x_{i+k})}}{\sum_{x_l \in \mathcal{V}} e^{\phi(x_i)^\top \phi(x_l)}}$$

Maximize the probability word embedding can predict neighbours in some context window (of size c)

⁵Usually two embeddings used: word and word as context. Simplified here.

Re-writing the equation:

$$\left(\sum_i \sum_j \sum_{-c \leq k \leq c, k \neq 0} \log e^{\phi(x_j)^\top \phi(x_{j+k})} \right) - \left(\sum_i \sum_j \sum_{-c \leq k \leq c, k \neq 0} \log \sum_{x_l \in \mathcal{V}} e^{\phi(x_j)^\top \phi(x_l)} \right)$$

- On the left: Sum over positive contexts
- On the right: Sum over negative contexts
 - Not feasible to sum over entire vocabulary

- Solution: negative sampling

$$\left(\sum_i \sum_j \sum_{-c \leq k \leq c, k \neq 0} \log e^{\phi(x_j)^\top \phi(x_{j+k})} \right) - \left(\sum_i \sum_j \sum_{-c \leq k \leq c, k \neq 0} \log \sum_{x_l \in \mathcal{V}_s} e^{\phi(x_j)^\top \phi(x_l)} \right)$$

- $\mathcal{V}_s$ is randomly sampled, i.e., $\mathcal{V}_s \subset \mathcal{V}$ and $|\mathcal{V}_s| \ll |\mathcal{V}|$ (often 1)

$$\left(\sum_i \sum_j \sum_{-c \leq k \leq c, k \neq 0} \log e^{\phi(x_j)^\top \phi(x_{j+k})} \right) - \left(\sum_i \sum_j \sum_{-c \leq k \leq c, k \neq 0} \log \sum_{x_l \in \mathcal{V}_s} e^{\phi(x_j)^\top \phi(x_l)} \right)$$

- $\phi(x_i)$ are used as final word embeddings
- Usually optimized with SGD⁶

⁶Negate function to make it a loss and minimize.

Variable Length Inputs

- What is input x is not just a single word?
- Or x is not of fixed length, e.g., sentences or documents?
- E.g., $x = x_1 \dots x_n$

COMMON SOLUTIONS

- Average: $\phi(x) = \frac{1}{|x|} \sum_{x \in x} \phi(x)$
 - Other pointwise operations, e.g., max, sum, ...
- Truncation+concatentation: $\phi(x) = [\phi(x_1); \dots; \phi(x_k)] \in \mathbb{R}^{kD}$ for a fixed k (stack the column vectors)
- Sparse-dense: Whole look-up table is input, but
 - Zero out rows of words that are not present
 - Usually not practical

Outline

① Terminology, notation and feature representations

② Linear Classifiers

Perceptron

Logistic Regression

③ Regularization

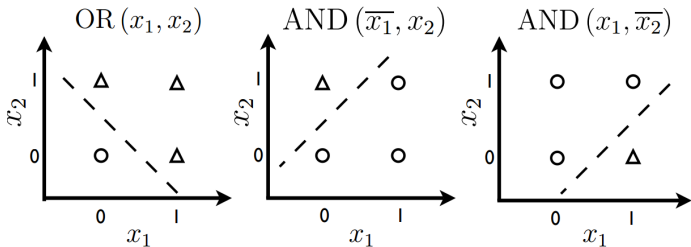
④ Dense Representations

⑤ Similarity-based Learning

⑥ Neural Networks

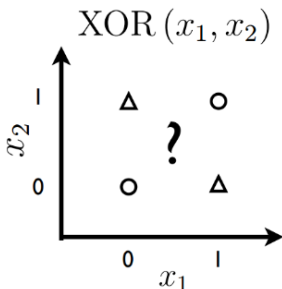
Recap: What a Linear Classifier Can Do

- It **can** solve linearly separable problems (OR, AND)



Recap: What a Linear Classifier **Can't** Do

- ... but it **can't** solve **non-linearly separable** problems such as simple XOR:



What If Data Are Not Linearly Separable?

What If Data Are Not Linearly Separable?

Engineer better features (often works!)



What If Data Are Not Linearly Separable?

Engineer better features (often works!)



Similarity-based / **Kernel methods:**

- define a similarity / kernel function between points
- use it to classify new instances; need a good function



What If Data Are Not Linearly Separable?

Engineer better features (often works!)



Similarity-based / **Kernel methods:**

- define a similarity / kernel function between points
- use it to classify new instances; need a good function



Neural networks (up next)

- embrace non-convexity and local minima

Two Views of Machine Learning

There's two big ways of building machine learning systems:

- ① **Feature-based**: describe objects' properties (features) and build models that manipulate them
 - everything that we have seen so far.
- ② **Similarity-based**: don't describe objects by their properties; rather, build systems based on **comparing** objects to each other
 - k -th nearest neighbors; kernel methods; Gaussian processes.

Sometimes the two are equivalent!

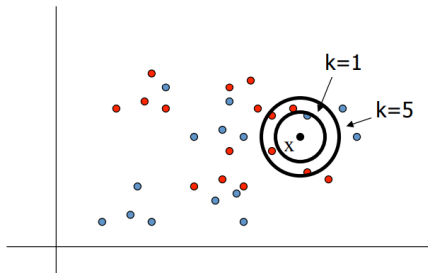
Parametric vs. Non-parametric

Another way to classify machine learning systems:

- 1 **Parametric:** Fix the number of parameters, model structure and hypothesis space
 - Goal: find parameters to optimize objective in the hypothesis space
 - Everything so far plus NNs
- 2 **Non-parametric:** No/little assumptions about form of solution; size of parameters not fixed
 - Goal: find the function/solution to best fit data
 - Similarity methods (k -th nearest neighbors); kernel methods; decision trees, random forests, gaussian processes.

(K-th) Nearest Neighbor Classifier

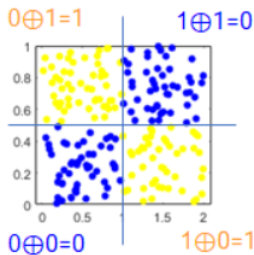
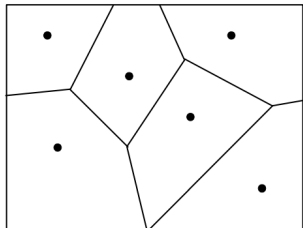
- Memorize (training) data $\mathcal{D} = \{\mathbf{x}_t, \mathbf{y}_t\}_{t=1}$
- To classify a new datapoint $\mathbf{x}$
 - Find the k closest data points in the data (e.g., training set)
 - Assign the most frequent class in those k points



(K-th) Nearest Neighbor Classifier

- We are not learning any parameters
- Hypothesis space can change by adding more data points
- Power is greater than linear classifier (xor on right)

1-NN Decision Surface



Similarity Functions and Inference

- Common Similarity functions

- Euclidean: $\|\phi(x) - \phi(x')\| = \sqrt{\sum_i (\phi_i(x) - \phi_i(x'))^2}$

- Inner product: $\phi(x)^\top \phi(x')$

- Cosine: $\frac{\phi(x)^\top \phi(x')}{\|\phi(x)\| \|\phi(x')\|}$

- Function can also be learned!

- Searching a K-NN database

- Dense retrieval! Used in RAG, search, etc.

- Brute-force

- Branch and bound / k-d tree

- Approx: Greedy proximity graph; locality sensitive hashing

Outline

① Terminology, notation and feature representations

② Linear Classifiers

Perceptron

Logistic Regression

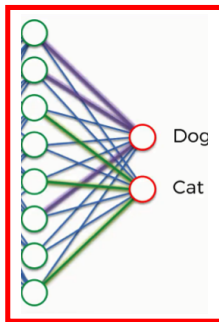
③ Regularization

④ Dense Representations

⑤ Similarity-based Learning

⑥ Neural Networks

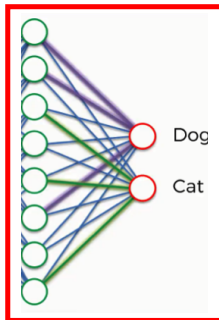
Recap: Weights and Biases



Linear Classifier

$$z = \mathbf{W}\phi(x) + \mathbf{b}, \quad y' = \operatorname{argmax}(z), \quad \mathbf{W} = \begin{bmatrix} \vdots \\ w_y^\top \\ \vdots \end{bmatrix}, \quad \mathbf{b} = \begin{bmatrix} \vdots \\ b_y \\ \vdots \end{bmatrix}$$

Recap: Weights and Biases

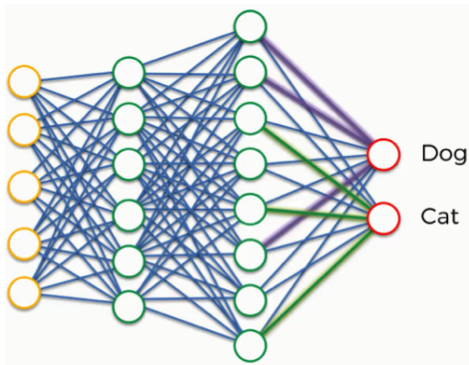


Linear Classifier

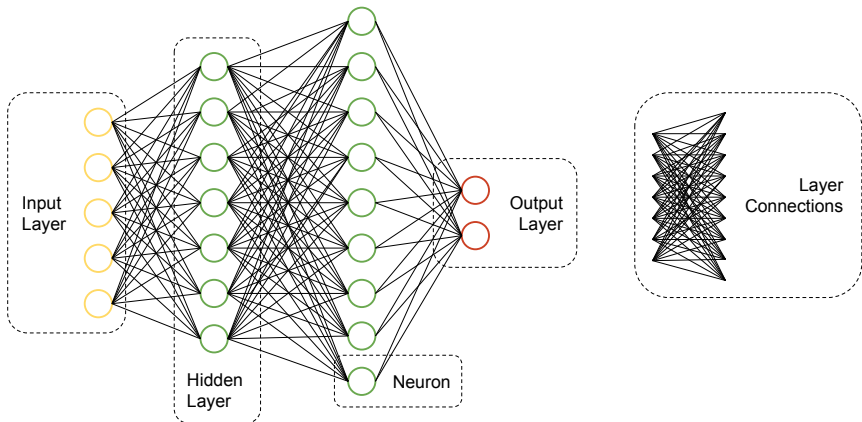
$$\mathbf{z} = \mathbf{W}\mathbf{x} + \mathbf{b}, \quad y' = \operatorname{argmax}(\mathbf{z}), \quad \mathbf{W} = \begin{bmatrix} \vdots \\ w_y^\top \\ \vdots \end{bmatrix}, \quad \mathbf{b} = \begin{bmatrix} \vdots \\ b_y \\ \vdots \end{bmatrix}$$

Notation: input is now a **dense** vector $\mathbf{x}$ instead of sparse features $\phi(\mathbf{x})$.

On to neural networks!

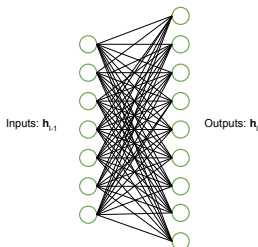


Neurons, Layers and Connections



- A (dense / fully-connected) feed-forward neural network (FF-NN)
 - AKA fully connected network (FCN) / multilayer perceptron (MLP)
- Input and output layers are special (more on this)
- However connections between layers take a similar form

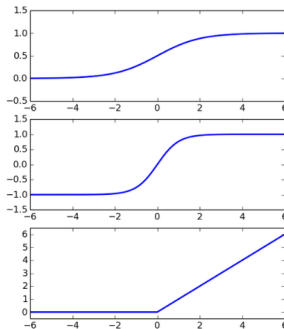
Hidden Layer Connections



- Let $\mathbf{h}_l \in \mathbb{R}^{D_l}$ be the l^{th} hidden layer with D_l dimensions/neurons
- $\mathbf{h}_0 = \mathbf{x}$ the input
- $\mathbf{h}_l = \sigma_l(\mathbf{W}_l \mathbf{h}_{l-1} + \mathbf{b}_l)$ ← same form as the linear classifier
- $\mathbf{W}_l \in \mathbb{R}^{D_l \times D_{l-1}}$ and $\mathbf{b}_l \in \mathbb{R}^{D_l}$ are layer parameters
- σ_l is the layer's (**non-linear**) activation function
- Row j of $\mathbf{W}_l$ = the weights of **one neuron** (instead of one label)

Activation Functions

- Non-linearity by transforming/projecting the data
- Squashes output to finite range
- Examples ...



Sigmoid

$$\phi(z) = \frac{1}{1 + e^{-z}}$$

Hyperbolic Tangent

$$\phi(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}}$$

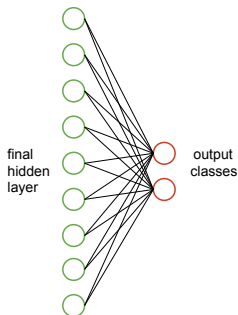
Rectified Linear

$$\phi(z) = \begin{cases} 0 & \text{if } z < 0 \\ z & \text{if } z \geq 0 \end{cases}$$

From Hughes and Correll 2016

Output Layer

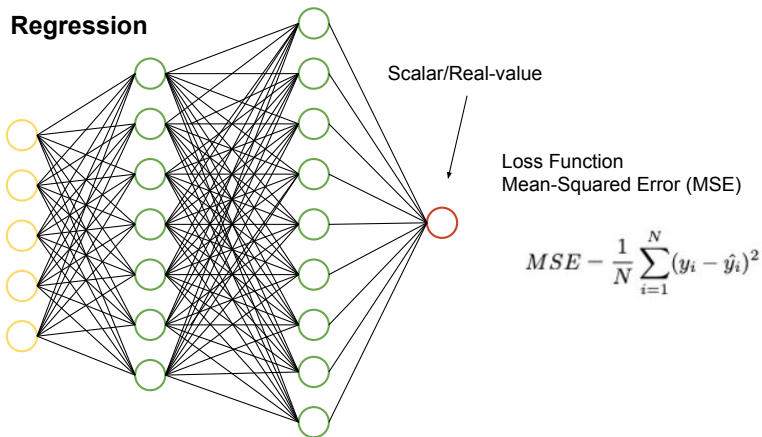
- This was in the last lecture
- $\mathbf{z} = \mathbf{W}_L \mathbf{h}_{L-1} + \mathbf{b}_L$; $y' = \text{argmax}(\mathbf{z})$
- z_y is the **logit**: the raw score of label y
- **Training?**
 - Common paradigm: take the softmax of the logits, $\mathbf{p} = \text{softmax}(\mathbf{z})$
 - Gives probability distribution $P(y|x) = p_y, \forall p_y \in \mathbf{p}$
 - Minimize some loss



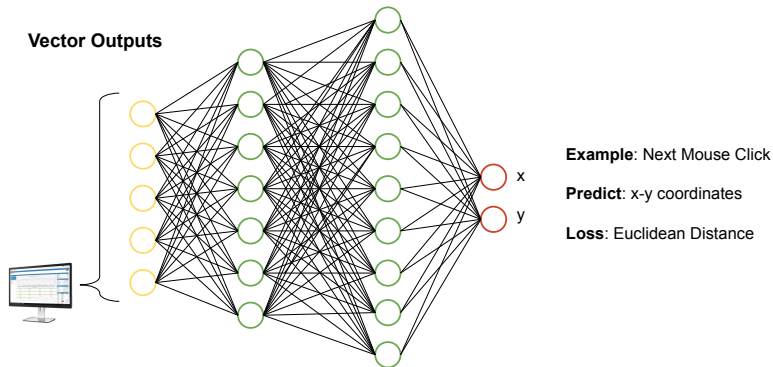
Example NN Loss: Cross-Entropy

- Can use any differentiable loss function.
- Let $\Theta = \{(\mathbf{W}_l, \mathbf{b}_l)\}$, i.e., parameters for all layers l
- **Goal:** Find Θ to maximize $P(\mathbf{y}|\mathbf{x}; \Theta)$ for all $(\mathbf{x}, \mathbf{y}) \in \mathcal{D}$
- **Cross-Entropy (CE)**
 - Intuition: Avg. bits needed to distinguish two distributions
 - Let $\bar{P}$ be the truth: $\bar{P}(\mathbf{y}'|\mathbf{x}_t) = 1$ if $\mathbf{y}' = \mathbf{y}_t$ and 0 otherwise
 - $CE = -E_{\bar{P}} \log P = -\sum_{(\mathbf{x}, \mathbf{y}) \in \mathcal{D}} \bar{P}(\mathbf{y}|\mathbf{x}) \log P(\mathbf{y}|\mathbf{x}; \Theta)$
 - Therefore: $CE = -\sum_{(\mathbf{x}, \mathbf{y}) \in \mathcal{D}} \log P(\mathbf{y}|\mathbf{x}; \Theta)$
- Min of $CE = -\sum_{(\mathbf{x}, \mathbf{y})} \log P(\mathbf{y}|\mathbf{x}; \Theta)$
- **Cross-entropy = Min Neg Log-likelihood / log-loss**
- **In other words: logistic regression!**

Examples: NN for non-classification tasks



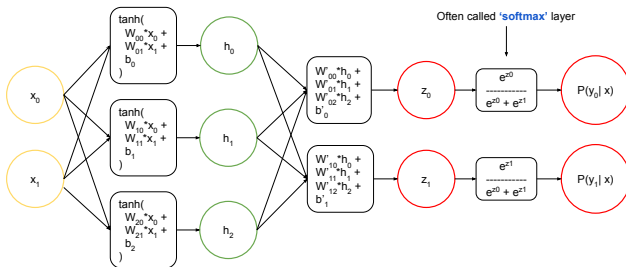
Examples: NN for non-classification tasks



A Wee Example

- $x \in \mathbb{R}^2$
- $h = \tanh(Wx + b)$ with $W \in \mathbb{R}^{3 \times 2}$ and $b \in \mathbb{R}^3$
- $|y| = 2$ with $z = W'h + b'$, $W' \in \mathbb{R}^{2 \times 3}$ and $b' \in \mathbb{R}^2$
- Cross-entropy:

- $L(\Theta; (x, y)) = -\log(P(y|x)) = -\log \frac{e^{z_y}}{\sum_{y' \in y} e^{z_{y'}}}$



Neural Networks So Far

- Neural network structure (FF-NN; FCN; MLP)
 - Input layer: for now, assume given to us $\mathbf{x} \in \mathbb{R}^D$
 - Outputs: $\mathbf{y} \in \mathcal{Y}$
 - Hidden layers: $\mathbf{h}_l \in \mathbb{R}^{D_l}$; with $\mathbf{h}_l = \sigma_l(\mathbf{W}_l \mathbf{h}_{l-1} + \mathbf{b}_l)$
 - Thus, model parameters $\Theta = \{\mathbf{W}_l, \mathbf{b}_l \mid \forall l\}$
 - Including the last (output) layer, $\mathbf{z} = \mathbf{W}_L \mathbf{h}_{L-1} + \mathbf{b}_L$
 - Loss function: $L(\Theta; (\mathbf{x}, \mathbf{y}))$ – usually log-loss/cross-entropy

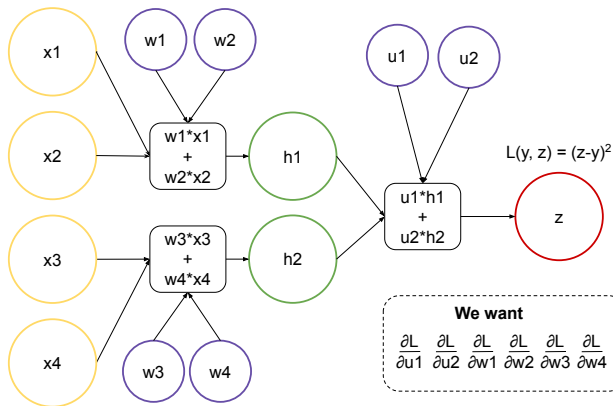
Neural Networks: Optimization

- Hidden layers make model non-convex!
- No single global optimum. Must settle for a local one.
- If loss function and activation functions are differentiable, then can be optimized with gradient-based techniques (e.g., gradient descent)
- Gradient computation a little trickier
 - Solution: **backpropagation** (Rumelhart et al., 1988)

Backpropagation and the Chain Rule

- We need to compute $\nabla_{\Theta} L(\Theta; \mathcal{D}) = [\frac{\partial L}{\partial w_0}, \frac{\partial L}{\partial w_1}, \dots]$, $\forall w_i \in \Theta$
 - For linear classifiers, Θ was a single $(\mathbf{W}, \mathbf{b})$
 - For NNs, Θ is one $(\mathbf{W}_l, \mathbf{b}_l)$ **per layer**: $\Theta = \{\mathbf{W}_l, \mathbf{b}_l \mid \forall l\}$
 - Each $\nabla_{\mathbf{W}_l} L$ has the same shape as $\mathbf{W}_l$
- Chain rule: $\frac{\partial z}{\partial y} \frac{\partial y}{\partial x}$

Toy Example: Analytical Partial Derivatives



We want

$$\frac{\partial L}{\partial u_1} \quad \frac{\partial L}{\partial u_2} \quad \frac{\partial L}{\partial w_1} \quad \frac{\partial L}{\partial w_2} \quad \frac{\partial L}{\partial w_3} \quad \frac{\partial L}{\partial w_4}$$

All base derivatives

$$\frac{\partial L}{\partial z} = 2(z - y)$$

$$\frac{\partial z}{\partial h_1} = u_1 \quad \frac{\partial z}{\partial h_2} = u_2$$

$$\frac{\partial z}{\partial u_1} = h_1 \quad \frac{\partial z}{\partial u_2} = h_2$$

$$\frac{\partial h_1}{\partial w_1} = x_1 \quad \frac{\partial h_1}{\partial w_2} = x_2$$

$$\frac{\partial h_1}{\partial x_1} = w_1 \quad \frac{\partial h_1}{\partial x_2} = w_2$$

$$\frac{\partial h_2}{\partial w_3} = x_3 \quad \frac{\partial h_2}{\partial w_4} = x_4$$

$$\frac{\partial h_2}{\partial x_3} = w_3 \quad \frac{\partial h_2}{\partial x_4} = w_4$$

Full derivation examples

$$\frac{\partial L}{\partial u_1} = \frac{\partial L}{\partial z} \frac{\partial z}{\partial u_1} = 2(z - y) * h_1$$

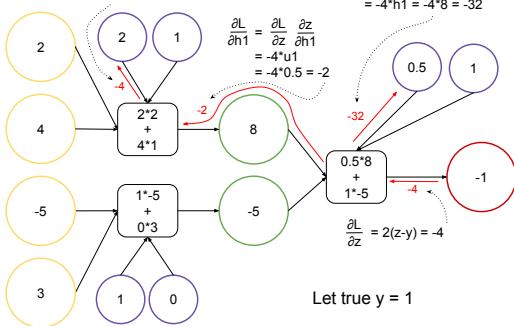
$$\frac{\partial L}{\partial w_1} = \frac{\partial L}{\partial h_1} \frac{\partial h_1}{\partial w_1} = \frac{\partial L}{\partial z} \frac{\partial z}{\partial h_1} \frac{\partial h_1}{\partial w_1} = 2(z - y) * u_1 * x_1$$

Toy Example: Backpropagation at Work

- Analytically computing chain rule in deep networks is onerous
- Backpropagation
 - Forward pass: compute values at neurons and final loss
 - Backward pass: compute $\frac{\partial L}{\partial w_i}$ at each neuron
 - $\frac{\partial L}{\partial w_i}$ of parameter neurons form gradient

$$\frac{\partial L}{\partial w_1} = \frac{\partial L}{\partial h_1} \frac{\partial h_1}{\partial w_1} = -2 \times 1 = -2 \times 2 = -4$$

$$\begin{aligned} \frac{\partial L}{\partial u_1} &= \frac{\partial L}{\partial z} \frac{\partial z}{\partial u_1} \\ &= -4 \times h_1 = -4 \times 8 = -32 \end{aligned}$$



Neuron derivatives

$$\frac{\partial L}{\partial z} = 2(z-y)$$

$$\frac{\partial L}{\partial h_1} = \frac{\partial L}{\partial z} \frac{\partial z}{\partial h_1}$$

$$\frac{\partial L}{\partial h_2} = \frac{\partial L}{\partial z} \frac{\partial z}{\partial h_2}$$

$$\frac{\partial L}{\partial u_1} = \frac{\partial L}{\partial z} \frac{\partial z}{\partial u_1}$$

$$\frac{\partial L}{\partial u_2} = \frac{\partial L}{\partial z} \frac{\partial z}{\partial u_2}$$

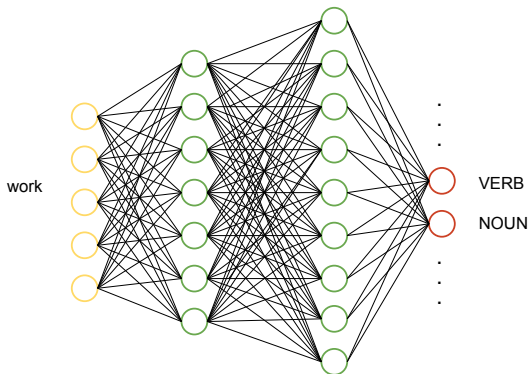
$$\frac{\partial L}{\partial w_1} = \frac{\partial L}{\partial h_1} \frac{\partial h_1}{\partial w_1}$$

$$\frac{\partial L}{\partial w_2} = \frac{\partial L}{\partial h_1} \frac{\partial h_1}{\partial w_2}$$

$$\frac{\partial L}{\partial w_3} = \frac{\partial L}{\partial h_2} \frac{\partial h_2}{\partial w_3}$$

$$\frac{\partial L}{\partial w_4} = \frac{\partial L}{\partial h_2} \frac{\partial h_2}{\partial w_4}$$

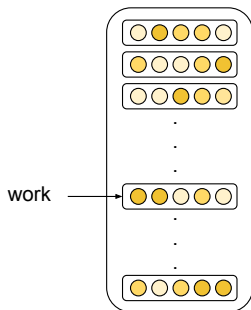
Input layer



- Consider classifying a word in isolation with a part-of-speech tag⁷
- The input is a word from a fixed, finite vocabulary $\mathcal{V}$
- But the network needs a vector: $\mathbf{h}_0 = \mathbf{x} \in \mathbb{R}^D$

⁷This is contrived. We usually use context.

Input layer = Embedding layer

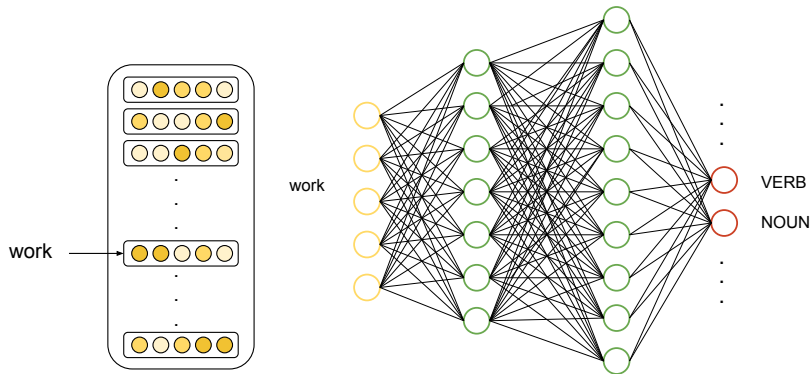


- Each word in $\mathcal{V}$ gets a vector in $\mathbb{R}^D$
- We store these in the **embedding matrix** $\mathbf{E} \in \mathbb{R}^{|\mathcal{V}| \times D}$ – one row per word
 - These are the model *word embeddings*
 - AKA embedding layer; word look-up table; ...
 - The lookup is again a matrix product: $\mathbf{h}_0 = \mathbf{E}^\top \mathbf{e}_x$

Input layer = Embedding layer

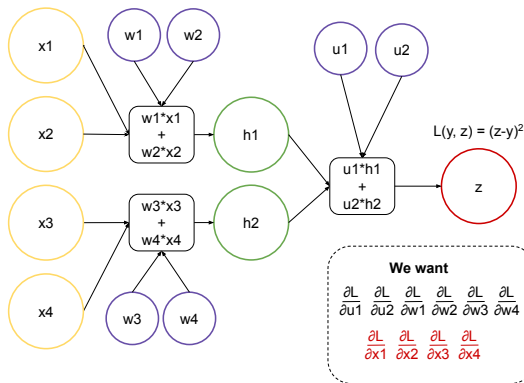
- Static embedding layer
 - Fixed word embeddings; not updated during training
 - Examples: SVD; **word2vec**; glove; ...
- Dynamic embedding layer
 - Randomly initialize word embeddings
 - Learn during training of the full network
 - Updated like any other layer during backpropagation
- Static + Dynamic
 - Initialize model with static embeddings; update dynamically
 - Combination: part of embedding layer is static; part is learned

Input layer



- Static (e.g., word2vec) or dynamic word embeddings give us input layer

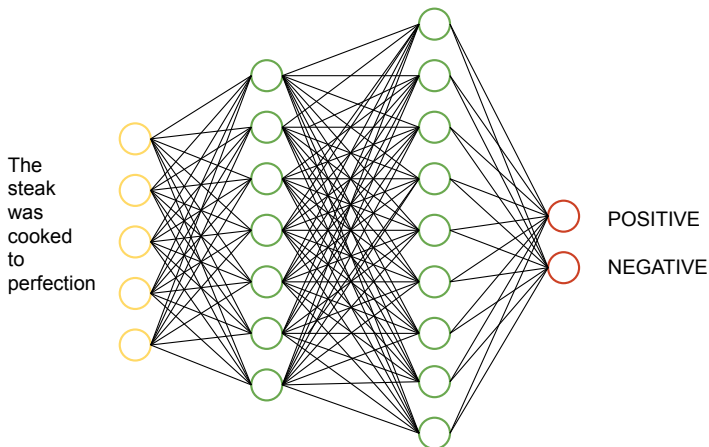
Dynamic Input layer



- Gradient now includes the input neurons, $\frac{\partial L}{\partial x_i}$, and hence $\nabla_{\mathbf{E}} L$
- Every value in the entire lookup table is a parameter:
 $\Theta = \{\mathbf{E}\} \cup \{\mathbf{W}_l, \mathbf{b}_l \mid \forall l\}$

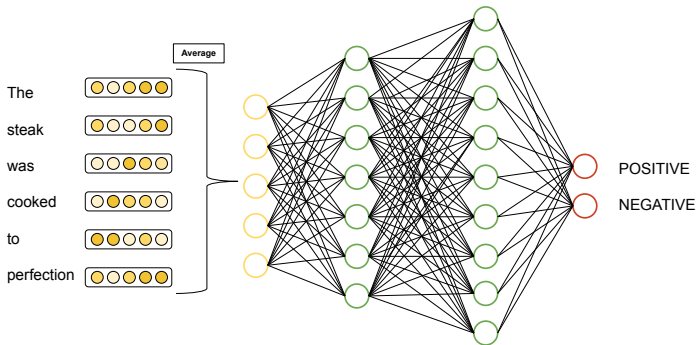
Variable Length Inputs

- But what if input is a whole document and not just a single word?
- Feed-forward neural networks assume a fixed-length input, $x \in \mathbb{R}^D$
- Documents are not fixed length



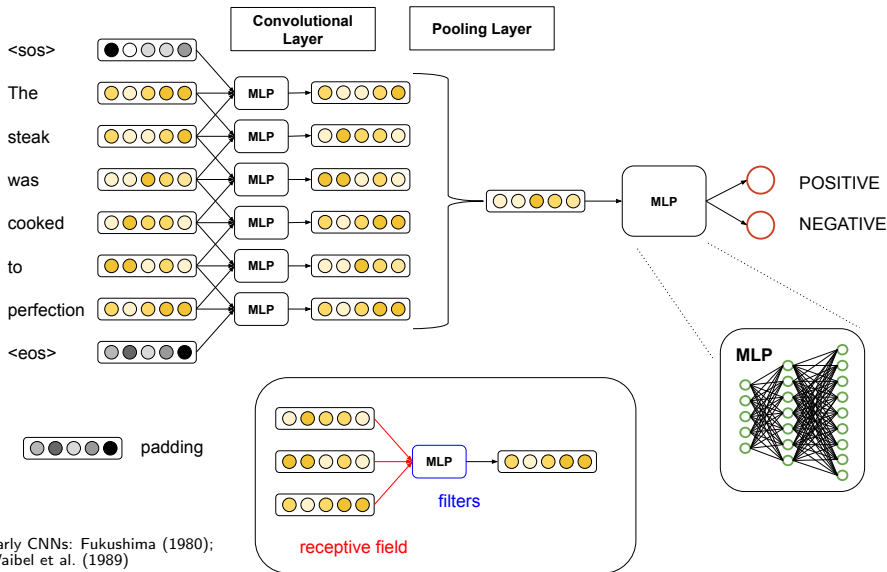
Variable Length Inputs: Options

- 1 Truncate document at fixed length K and stack the embeddings, $x \in \mathbb{R}^{KD}$
- 2 Average embeddings (**below**), $x \in \mathbb{R}^D$
- 3 **convolutional** (CNNs) and recurrent neural networks (RNNs)⁸



⁸ RNNs not covered. See <https://athnlp.github.io/2024/ppts/Day02AthNLP2024-Lec2-BPlank-handout.pdf>

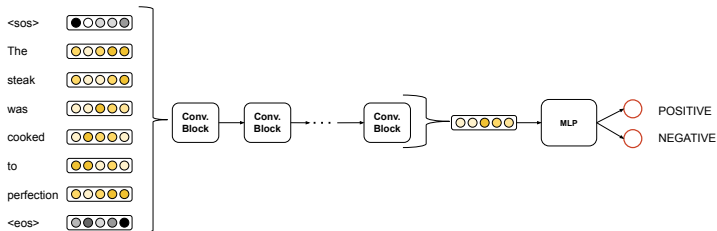
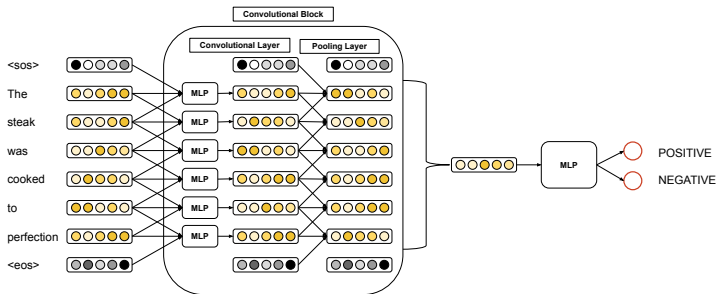
Convolutional Neural Networks



Convolutional Neural Networks

- Convolutional layer
 - A NN sub-architecture
 - Slides over input at a fixed **stride**, usually 1
 - **Receptive field**: fixed size input (e.g., n -gram)
 - **Filter**: MLP that creates a single vector output per position
 - Can be multiple filters: Almost always shared positionally; sometimes even per layer
- Pooling layer
 - Converts convolutional output to a single fixed-length vector
 - **Average pooling**: average outputs of convolutional layers
 - **Max pooling**: position-wise max over outputs of convolutional layers
 - **NN**: Can also learn this, e.g., attention.

Deep Convolutional Neural Networks



Core Neural Network Summary

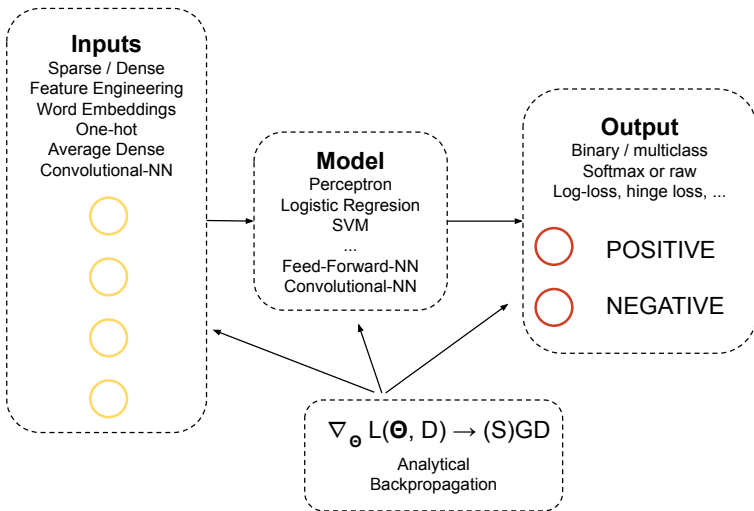
- Fully-connected Feed-forward Neural Networks
- Neurons, layers and connections
- Output layers (linear) and losses
- Back propagation
- Input layers
 - Static vs dynamic vs mixed
- Variable length inputs: CNNs
- Deep NNs – keep adding layers

Where Does Network Structure Come From?

- Hyperparameters: input/hidden dimensions; activation functions; ...
 - Usually empirical but becoming standardized (e.g., transformers)
- **Deep Learning** = lot's of layers. How many? Empirical accuracy vs. resources.
- Fully-connected/dense required?
 - No!
 - Sparse layers / chunks. Especially in LLMs
 - But for FF-NN components usually full-connected
 - Any efficiency concerns lessened by modern architectures (GPU, TPU)

Main Points (Parametric ML)

The
steak
was
cooked
to
perfection



... in Words

- Sparse (binary) vs. dense (embeddings) features
- Optimization: Use gradient-based techniques
- Linear Classifiers
 - Scores are a matrix-vector product: $\mathbf{z} = \mathbf{W}\phi(x) + \mathbf{b}$, one row of $\mathbf{W}$ per label
 - Usually sparse features
 - Loss functions define model
 - Regularization necessary for good performance
- Sparse vs. Dense representations
- Parameteric (e.g., linear cls) vs. Non-parametric (e.g., knn) ML
- Neural Networks
 - Final layer = linear classifier: same $\mathbf{z} = \mathbf{W}_L \mathbf{h}_{L-1} + \mathbf{b}_L$
 - Hidden layers = linear + non-lin activation
 - Compute gradient with backpropagation
 - Input layer: static (e.g., word2vec) vs. dynamic (backprop)
 - (Deep) CNNs for variable length inputs

References I

- Fukushima, K. (1980). Neocognitron: A self-organizing neural network model for a mechanism of pattern recognition unaffected by shift in position. *Biological cybernetics*, 36(4):193–202.
- Mikolov, T., Sutskever, I., Chen, K., Corrado, G. S., and Dean, J. (2013). Distributed representations of words and phrases and their compositionality. In *Advances in neural information processing systems*, pages 3111–3119.
- Minsky, M. and Papert, S. (1969). Perceptrons.
- Novikoff, A. B. (1962). On convergence proofs for perceptrons. In *Symposium on the Mathematical Theory of Automata*.
- Rosenblatt, F. (1958). The perceptron: A probabilistic model for information storage and organization in the brain. *Psychological review*, 65(6):386.
- Rumelhart, D. E., Hinton, G. E., and Williams, R. J. (1988). Learning representations by back-propagating errors. *Cognitive modeling*, 5(3):1.
- Waibel, A., Hanazawa, T., Hinton, G., Shikano, K., and Lang, K. J. (1989). Phoneme recognition using time-delay neural networks. *IEEE transactions on acoustics, speech, and signal processing*, 37(3):328–339.