

Machine Learning I

Lucas Fabian Naumann, David Stein, Bjoern Andres

Machine Learning for Computer Vision
TU Dresden



<https://mlcv.cs.tu-dresden.de/courses/25-winter/ml1/>

Winter Term 2025/2026

Contents. This part of the course is about the supervised learning of linear functions, more specifically, about logistic regression.

- ▶ We introduce the problem by defining labeled data, a family of functions and a probability measure whose maximization motivates a regularizer and a loss function.
- ▶ We show: This supervised learning problem is convex. It can be solved, e.g., by the steepest descent algorithm.

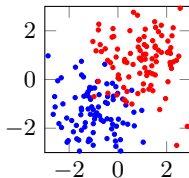
Logistic regression

We consider labeled data with **real features**. More specifically, we consider some finite, non-empty set V , called the set of features, and labeled data $T = (S, X, x, y)$ such that $X = \mathbb{R}^V$. Hence:

$$x: S \rightarrow \mathbb{R}^V \quad (1)$$

$$y: S \rightarrow \{0, 1\} \quad (2)$$

Example.

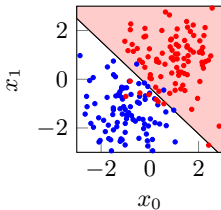


Logistic regression

We consider **linear functions**. More specifically, we consider $\Theta = \mathbb{R}^V$ and $f : \Theta \rightarrow \mathbb{R}^X$ such that

$$\forall \theta \in \Theta \quad \forall \hat{x} \in X: \quad f_{\theta}(\hat{x}) = \langle \theta, \hat{x} \rangle = \sum_{v \in V} \theta_v \hat{x}_v \quad (3)$$

Example.



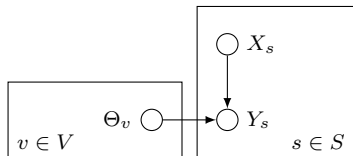
Logistic regression

We introduce a probabilistic model:

- ▶ For any sample $s \in S$, let X_s be a random variable whose value is a vector $x_s \in \mathbb{R}^V$, the **feature vector** of s
- ▶ For any sample $s \in S$, let Y_s be a random variable whose value is a binary number $y_s \in \{0, 1\}$, the **label** of s
- ▶ For any $v \in V$, let Θ_v be a random variable whose value is a real number $\theta_v \in \mathbb{R}$, a **parameter** of the linear function we seek to learn

We assume that the joint probability factorizes according to:

$$P(X, Y, \Theta) = \prod_{s \in S} (P(Y_s | X_s, \Theta) P(X_s)) \prod_{v \in V} P(\Theta_v) \quad (4)$$



Logistic regression

We attempt to learn parameters by maximizing the conditional probability

$$\begin{aligned} P(\Theta \mid X, Y) &= \frac{P(X, Y, \Theta)}{P(X, Y)} \\ &= \frac{P(Y \mid X, \Theta) P(X) P(\Theta)}{P(X, Y)} \\ &\propto P(Y \mid X, \Theta) P(\Theta) \\ &= \prod_{s \in S} P(Y_s \mid X_s, \Theta) \prod_{v \in V} P(\Theta_v) . \end{aligned}$$

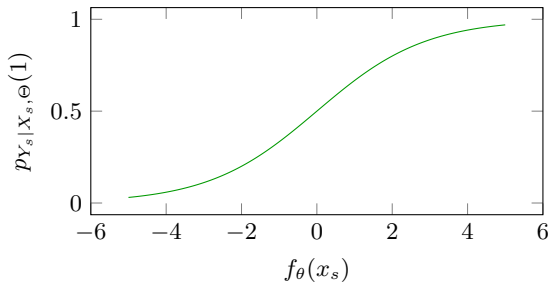
We attempt to infer labels by maximizing the conditional probability

$$P(Y \mid X, \Theta) = \prod_{s \in S} P(Y_s \mid X_s, \Theta) .$$

Logistic regression

► Sigmoid distribution

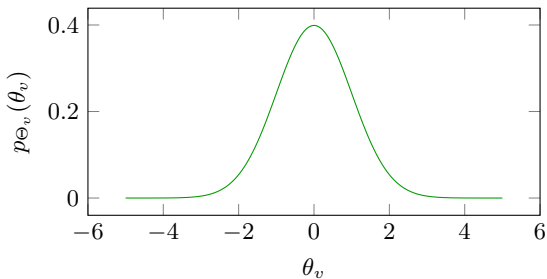
$$\forall s \in S : \quad p_{Y_s|X_s, \Theta}(1) = \frac{1}{1 + 2^{-f_{\theta}(x_s)}} \quad (5)$$



Logistic regression

- **Normal distribution** with $\sigma \in \mathbb{R}^+$:

$$\forall v \in V : \quad p_{\Theta_v}(\theta_v) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\theta_v^2/2\sigma^2} \quad (5)$$



Logistic regression

Lemma. Estimating maximally probable parameters θ , given attributes x and labels y , i.e.,

$$\operatorname{argmax}_{\theta \in \mathbb{R}^m} p_{\Theta|X,Y}(\theta, x, y)$$

is equivalent to the supervised learning problem

$$\min_{\theta \in \Theta} \lambda R(\theta) + \sum_{s \in S} L(f_{\theta}(x_s), y_s) \quad (6)$$

with L , R and λ such that

$$\forall r \in \mathbb{R} \ \forall \hat{y} \in \{0, 1\}: \quad L(r, \hat{y}) = -\hat{y}r + \log(1 + 2^r) \quad (7)$$

$$\forall \theta \in \Theta: \quad R(\theta) = \|\theta\|_2^2 \quad (8)$$

$$\lambda = \frac{\log e}{2\sigma^2} . \quad (9)$$

It is called the l_2 -regularized **logistic regression problem** with respect to x , y and σ .

Logistic regression

Proof. Firstly,

$$\begin{aligned} & \operatorname{argmax}_{\theta \in \mathbb{R}^m} p_{\Theta|X,Y}(\theta, x, y) \\ &= \operatorname{argmax}_{\theta \in \mathbb{R}^m} \prod_{s \in S} p_{Y_s|X_s, \Theta}(y_s, x_s, \theta) \prod_{v \in V} p_{\Theta_v}(\theta_v) \\ &= \operatorname{argmax}_{\theta \in \mathbb{R}^m} \sum_{s \in S} \log p_{Y_s|X_s, \Theta}(y_s, x_s, \theta) + \sum_{v \in V} \log p_{\Theta_v}(\theta_v) \end{aligned} \quad (10)$$

Secondly,

$$\begin{aligned} & \log p_{Y_s|X_s, \Theta}(y_s, x_s, \theta) \\ &= y_s \log p_{Y_s|X_s, \Theta}(1, x_s, \theta) + (1 - y_s) \log p_{Y_s|X_s, \Theta}(0, x_s, \theta) \\ &= y_s \log \frac{p_{Y_s|X_s, \Theta}(1, x_s, \theta)}{p_{Y_s|X_s, \Theta}(0, x_s, \theta)} + \log p_{Y_s|X_s, \Theta}(0, x_s, \theta) \end{aligned} \quad (11)$$

Thus, with (5) and (4):

$$\operatorname{argmin}_{\theta \in \mathbb{R}^m} \sum_{s \in S} \left(-y_s \langle \theta, x_s \rangle + \log \left(1 + 2^{\langle \theta, x_s \rangle} \right) \right) + \frac{\log e}{2\sigma^2} \|\theta\|_2^2 \quad (12)$$

Lemma. The objective function

$$\varphi(\theta) = \lambda R(\theta) + \frac{1}{|S|} \sum_{s \in S} L(f_{\theta}(x_s), y_s) \quad (13)$$

of the l_2 -regularized logistic regression problem is convex.

Proof. Exercise!

The l_2 -regularized logistic regression problem can be solved, e.g., by the steepest descent algorithm with a tolerance parameter $\epsilon \in \mathbb{R}_0^+$:

Algorithm. Steepest descent with line search

```
 $\theta := 0$ 
repeat
   $d := \nabla \varphi(\theta)$ 
   $\eta := \operatorname{argmin}_{\eta' \in \mathbb{R}} \varphi(\theta - \eta' d)$  (line search)
   $\theta := \theta - \eta d$ 
  if  $\|d\| < \epsilon$ 
    return  $\theta$ 
```

Lemma: Estimating maximally probable labels y , given attributes x' and parameters θ , i.e.,

$$\operatorname{argmax}_{y \in \{0,1\}^S} p_{Y|X,\Theta}(y, x', \theta) \quad (14)$$

is equivalent to the inference problem

$$\min_{y' \in \{0,1\}^S} \sum_{s \in S} L(f_\theta(x_s), y'_s) \quad (15)$$

It has the solution

$$\forall s \in S' : \quad y_s = \begin{cases} 1 & \text{if } f_\theta(x'_s) > 0 \\ 0 & \text{otherwise} \end{cases} \quad (16)$$

Logistic regression

Proof. Firstly,

$$\begin{aligned} & \operatorname{argmax}_{y \in \{0,1\}^{S'}} p_{Y|X,\Theta}(y, x', \theta) \\ &= \operatorname{argmax}_{y \in \{0,1\}^{S'}} \prod_{s \in S'} p_{Y_s|X_s,\Theta}(y_s, x'_s, \theta) \\ &= \operatorname{argmax}_{y \in \{0,1\}^{S'}} \sum_{s \in S'} \log p_{Y_s|X_s,\Theta}(y_s, x'_s, \theta) \\ &= \operatorname{argmax}_{y \in \{0,1\}^{S'}} \sum_{s \in S'} \left(y_s \log \frac{p_{Y_s|X_s,\Theta}(1, x'_s, \theta)}{p_{Y_s|X_s,\Theta}(0, x'_s, \theta)} + \log p_{Y_s|X_s,\Theta}(0, x'_s, \theta) \right) \\ &= \operatorname{argmin}_{y \in \{0,1\}^{S'}} \sum_{s \in S'} \left(-y_s f_\theta(x'_s) + \log \left(1 + 2^{f_\theta(x'_s)} \right) \right) \\ &= \operatorname{argmin}_{y \in \{0,1\}^{S'}} \sum_{s \in S'} L(f_\theta(x'_s), y_s) . \end{aligned}$$

Secondly,

$$\min_{y \in \{0,1\}^{S'}} \sum_{s \in S'} \left(-y_s f_\theta(x'_s) + \log \left(1 + 2^{f_\theta(x'_s)} \right) \right) = \sum_{s \in S'} \max_{y_s \in \{0,1\}} y_s f_\theta(x'_s) .$$

Summary.

- ▶ The l_2 -regularized logistic regression problem is a supervised learning problem wrt. the family of linear functions.
- ▶ It can be derived from a statistical model with the sigmoid distribution as the likelihood as the normal distribution as the prior.
- ▶ It is a convex optimization problem that can be solved, e.g., by the steepest descent algorithm.