

Maximum Likelihood Estimate (MLE) and the Gaussian model

S. Marchand-Maillet

Computer Science – University of Geneva

Given $\mathcal{X} = \{\mathbf{x}_1, \dots, \mathbf{x}_N\}$ where $\mathbf{x}_i \in \mathbb{X} \subseteq \mathbb{R}^D$, we wish to model the generative process of $\mathcal{X}$. The Maximum Likelihood criterion:

1. assumes a model $\mathbf{p}_\theta$ with parameters θ for the generative density (of which the data is supposed to be a sample)
2. selects its optimal parameters θ^* such that the likelihood of $\mathcal{X}$ as a sample of $\mathbf{p}_{\theta^*}$ is maximized over all possible parameters.

Definition 1 (Maximum Likelihood Estimate of parameters). Given a dataset $\mathcal{X} \subset \mathbb{X}$, a density model $\mathbf{p}_\theta$ with parameters $\theta \in \Theta$, the Maximum Likelihood Estimate (MLE) parameters θ^* maximize the likelihood of data $\mathcal{X}$ under sampling of density $\mathbf{p}_\theta$ and are given by:

$$\theta^* = \operatorname{argmax}_{\theta \in \Theta} \mathbf{p}_\theta(\mathcal{X})$$

Now,

$$\mathbf{p}_\theta(\mathcal{X}) = \mathbf{p}_\theta(\mathbf{x}_1, \dots, \mathbf{x}_N) = \mathbf{p}_\theta(\mathbf{x}_1) \prod_{i=2}^N \mathbf{p}_\theta(\mathbf{x}_i | \mathbf{x}_{j < i})$$

We therefore assume that the data is identically and independently distributed (i.i.d) in order to alleviate the conditional dependence between data points:

$$\mathbf{p}_\theta(\mathbf{x}_i | \mathbf{x}_{j < i}) \stackrel{\text{i.i.d}}{=} \mathbf{p}_\theta(\mathbf{x}_i) \quad \Rightarrow \quad \mathbf{p}_\theta(\mathcal{X}) \stackrel{\text{i.i.d}}{=} \prod_{i=1}^N \mathbf{p}_\theta(\mathbf{x}_i)$$

Hence:

$$\theta^* = \operatorname{argmax}_{\theta \in \Theta} \mathbf{p}_\theta(\mathcal{X}) \stackrel{\text{i.i.d}}{=} \operatorname{argmax}_{\theta \in \Theta} \prod_{i=1}^N \mathbf{p}_\theta(\mathbf{x}_i) \quad (1)$$

Assuming a Gaussian generative model:

$$\mathbf{p}_\theta(\mathbf{x}_i) = \mathcal{N}(\mathbf{x}_i | \mu, \Sigma) = \frac{1}{(2\pi)^{\frac{D}{2}} |\Sigma|^{\frac{1}{2}}} e^{-\frac{1}{2} (\mathbf{x}_i - \mu)^\top \Sigma^{-1} (\mathbf{x}_i - \mu)} \quad \theta = [\mu, \Sigma]$$

which, in the case where $D = 1$ reduces to

$$\mathbf{p}_\theta(\mathbf{x}_i) = \mathcal{N}(\mathbf{x}_i | \mu, \sigma) = \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(\mathbf{x}_i - \mu)^2}{\sigma^2}} \quad \theta = [\mu, \sigma]$$

Inserting into (1):

$$\theta^* = \operatorname{argmax}_{\theta \in \Theta} \prod_{i=1}^N \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(\mathbf{x}_i - \mu)^2}{\sigma^2}}$$

We further write

$$\theta^* = \operatorname{argmax}_{\theta \in \Theta} \log \left(\prod_{i=1}^N \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(\mathbf{x}_i - \mu)^2}{\sigma^2}} \right) \quad \text{since } \log(\cdot) \text{ is a monotonically increasing function}$$

Denote:

$$\begin{aligned} \mathcal{L}_\theta^{\log}(\mathcal{X}) &= \log \left(\prod_{i=1}^N \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(\mathbf{x}_i - \mu)^2}{\sigma^2}} \right) \\ &= \sum_{i=1}^N \log \left(\frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{1}{2} \frac{(\mathbf{x}_i - \mu)^2}{\sigma^2}} \right) \\ &= \sum_{i=1}^N -\log(\sqrt{2\pi\sigma^2}) - \frac{1}{2} \frac{(\mathbf{x}_i - \mu)^2}{\sigma^2} \\ &= -\frac{N}{2} \log(2\pi\sigma^2) - \frac{1}{2} \sum_{i=1}^N \frac{(\mathbf{x}_i - \mu)^2}{\sigma^2} \\ &= -\frac{N}{2} \log(2\pi) - N \log \sigma - \frac{1}{2} \sum_{i=1}^N \frac{(\mathbf{x}_i - \mu)^2}{\sigma^2} \end{aligned}$$

so that

$$\boldsymbol{\theta}^* = \operatorname{argmax}_{\boldsymbol{\theta} \in \boldsymbol{\Theta}} \mathcal{L}_{\boldsymbol{\theta}}^{\log}(\mathcal{X}) = \operatorname{argmax}_{\boldsymbol{\theta} \in \boldsymbol{\Theta}} \left[-\frac{N}{2} \log(2\pi) - N \log \sigma - \frac{1}{2} \sum_{i=1}^N \frac{(x_i - \mu)^2}{\sigma^2} \right]$$

and

$$\left. \frac{\partial \mathcal{L}_{\boldsymbol{\theta}}^{\log}(\mathcal{X})}{\partial \boldsymbol{\theta}} \right|_{\boldsymbol{\theta}=\boldsymbol{\theta}^*} = 0$$

Now,

$$\frac{\partial \mathcal{L}_{\boldsymbol{\theta}}^{\log}(\mathcal{X})}{\partial \mu} = -\frac{1}{2} \sum_{i=1}^N 2(x_i - \mu) = -\sum_{i=1}^N (x_i - \mu)$$

so that

$$\frac{\partial \mathcal{L}_{\boldsymbol{\theta}}^{\log}(\mathcal{X})}{\partial \mu} = 0 \quad \Leftrightarrow \quad \sum_{i=1}^N (x_i - \mu) = 0 \quad \Leftrightarrow \quad \mu = \frac{1}{N} \sum_{i=1}^N x_i$$

Similarly,

$$\frac{\partial \mathcal{L}_{\boldsymbol{\theta}}^{\log}(\mathcal{X})}{\partial \sigma} = -\frac{N}{\sigma} + \frac{1}{\sigma^3} \sum_{i=1}^N (x_i - \mu)^2$$

so that

$$\frac{\partial \mathcal{L}_{\boldsymbol{\theta}}^{\log}(\mathcal{X})}{\partial \sigma} = 0 \quad \Leftrightarrow \quad -\frac{N}{\sigma} + \frac{1}{\sigma^3} \sum_{i=1}^N (x_i - \mu)^2 = 0 \quad \Leftrightarrow \quad \sigma^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2$$

We see that the classical empirical mean (μ) and variance (σ^2) computation correspond to the Maximum Likelihood Estimate of the data against a Gaussian model.