

Lecture 4 ML & Optimization

- Independence $\Rightarrow$ Law of Large Numbers (LLN)
- LLN justifies ERM aka learning from data

Independent random variables & their sums

$$X \perp\!\!\!\perp Y \iff p(x, y) = p(x)p(y)$$

Fact:

$$\mathbb{E}[XY] = \mathbb{E}[X] \mathbb{E}[Y]$$

Prf

$$\begin{aligned} \mathbb{E}[XY] &= \int_{\mathbb{R}^2} xy p(x, y) dx dy = \int_{\mathbb{R}^2} x p(x) y p(y) dx dy \\ &= \int x p(x) dx \int y p(y) dy = \mathbb{E}[X] \mathbb{E}[Y] \end{aligned}$$

Actually,

$$\begin{aligned} \forall f, g: \mathbb{E}[f(X)g(Y)] &= \mathbb{E}[f(X)]\mathbb{E}[g(Y)] \\ \iff X \perp\!\!\!\perp Y \end{aligned}$$

Fact if $X_1, \dots, X_K$ are r.v.s, then

$$\mathbb{E}[X_1 + \dots + X_K] = \mathbb{E}[X_1] + \dots + \mathbb{E}[X_K]$$

in fact

$$\mathbb{E}[a_1 X_1 + \dots + a_K X_K] = a_1 \mathbb{E}[X_1] + \dots + a_K \mathbb{E}[X_K]$$

Notation:

We say $X_1, \dots, X_K$ are i.i.d. (independent and identically distributed) if:

- $\forall i, X_i \sim \mathbb{P}$: identically distributed
- $p(x_1, \dots, x_K) = p(x_1) \dots p(x_K)$: independence

Supervised ML learns models f using i.i.d. samples
 $(x_i, y_i) \sim \mathbb{P}$

Variance of sums of r.v.s

① Fact: $X_1, \dots, X_K$ are independent

$$\Rightarrow \text{Var}(X_1 + \dots + X_K) = \text{Var}(X_1) + \dots + \text{Var}(X_K)$$

② Fact:

$$\text{Var}(aX) = a^2 \text{Var}(X)$$

Prf 2

$$\begin{aligned} \text{Var}(aX) &= \mathbb{E}(aX - a\mathbb{E}X)^2 = \mathbb{E}[a^2(X - \mathbb{E}X)^2] \\ &= a^2 \mathbb{E}(X - \mathbb{E}X)^2 = a^2 \text{Var}(X) \end{aligned}$$

Prf (1)

$$\text{Var}(X_1 + \dots + X_k) = \mathbb{E}(X_1 + \dots + X_k)^2 - \left[\mathbb{E}(X_1 + \dots + X_k) \right]^2$$

$$= \mathbb{E} \left[\sum_{i=1}^k \sum_{j=1}^k X_i X_j \right] - \left[\sum_{i=1}^k \mathbb{E}[X_i] \right]^2$$

$$= \sum_{i=1}^k \mathbb{E}[X_i^2] + \sum_{i=1}^k \sum_{\substack{j=1 \\ j \neq i}}^k \mathbb{E}[X_i] \mathbb{E}[X_j]$$

$$- \sum_{i=1}^k \sum_{j=1}^k \mathbb{E}[X_i] \mathbb{E}[X_j]$$

$$= \underbrace{\sum_{i=1}^k \mathbb{E}[X_i^2] - \sum_{i=1}^k (\mathbb{E}[X_i])^2}_{\text{Var}(X_i)} + \sum_{i=1}^k \sum_{\substack{j=1 \\ j \neq i}}^k \mathbb{E}[X_i] \mathbb{E}[X_j]$$

$$- \sum_{i=1}^k \sum_{\substack{j=1 \\ j \neq i}}^k \mathbb{E}[X_i] \mathbb{E}[X_j] = \underbrace{\sum_{i=1}^k \text{Var}(X_i)}$$

Behavior of Sample Mean

Sample $X_1, \dots, X_n$ i.i.d. from $\mathbb{P}$

and take
$$\mu_n = \frac{1}{n} \sum_{i=1}^n X_i$$

Fact 1

$$\mathbb{E} \mu_n = \mathbb{E} X_1 = \dots = \mathbb{E} X_n = \mu \quad (\text{where } \mu = \mathbb{E}_{X \sim \mathbb{P}} X)$$

Pf

$$\mathbb{E} \mu_n = \frac{1}{n} \sum_{i=1}^n \mathbb{E} X_i = \frac{1}{n} \sum_{i=1}^n \mu = \mu$$

Fact 2

$$\text{Var}(\mu_n) = \frac{1}{n} \sigma^2 \quad (\text{where } \sigma^2 = \text{Var}(X) \text{ when } X \sim \mathbb{P})$$



(cont'd)

$$\begin{aligned}\text{Var}(\mu_n) &= \text{Var}\left(\sum_{i=1}^n \frac{1}{n} X_i\right) = \sum_{i=1}^n \text{Var}\left(\frac{1}{n} X_i\right) \\ &= \sum_{i=1}^n \frac{1}{n^2} \text{Var}(X_i) = \sum_{i=1}^n \frac{\sigma^2}{n^2} \\ &= \frac{\sigma^2}{n}\end{aligned}$$

Cor LLN (see Chebyshev's ineq)

If $X_1, \dots, X_n \sim P$ are i.i.d. and $\varepsilon > 0$ is arbitrary

$$P(|\mu_n - \mu| > \varepsilon) \rightarrow 0 \text{ as } n \rightarrow \infty$$

LLN justifies ERM

In ML we want to minimize $\mu(f) = \mathbb{E}_{(x,y) \sim P} \ell(f(x), y)$
but we can't because we don't know P or we can't solve that optim problem.

Instead:

- sample $(x_i, y_i) \sim P$ i.i.d. for $i=1, \dots, n$
- form $\mu_n(f) = \frac{1}{n} \sum_{i=1}^n \ell(f(x_i), y_i)$
and minimize μ_n

We know $\mu_n(f) \rightarrow \mu(f)$ as $n \rightarrow \infty$ by LLN, so

ERM gives same minimizer as population risk

Parameterized ML models

In ML we want to learn

$$\mathbb{E}[y|x] = f(x) \text{ — regression function}$$

$$p(y|x) = g(x, y) \text{ — conditional distributive (predictive)}$$

In general, these are very complicated, so we use simplified **but useful** parameterized ML models

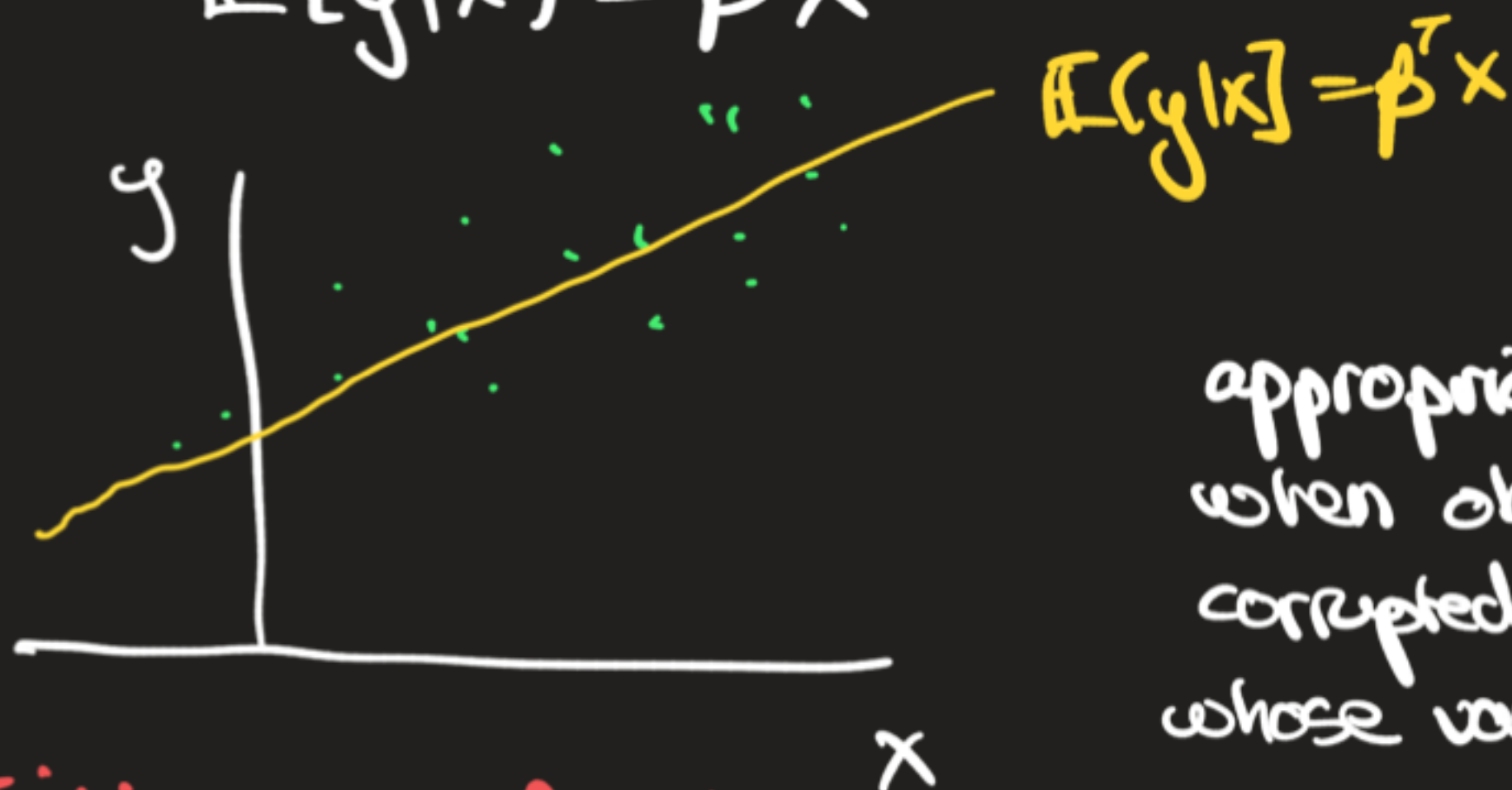
e.g. to learn

$$\mathbb{E}[y|x] = f(x; \theta) \text{ where } \theta \in \mathbb{R}^p \text{ is a set of parameters we optimize over}$$

Gaussian noise model for regression

Assume: $y|x \sim \mathcal{N}(\beta^T x, \sigma^2)$

$$\Rightarrow \mathbb{E}[y|x] = \beta^T x$$



Fitting: learn β using
OLS

Predicting: given x_{new} , take
 $y_{\text{pred}} = \beta^T x_{\text{new}}$

appropriate for regression
when observations look
corrupted by Gaussian noise
whose variance does not depend
on x , and there is a linear
trend in x

Classification w/ Bernoulli noise model

Assumption: $y|x \sim \text{Bern}(p(x))$

useful for binary classification (e.g. spam detection)

$$\mathbb{E}[y|x] = p(x)$$

where we require $p(x) \in [0, 1]$. To ensure this, we take

$p(x) = \sigma(\underbrace{\Theta^T x}_{\text{logit}})$ where σ is the logistic sigmoid



$$\sigma(u) = \frac{1}{1 + e^{-u}}$$

Very used for binary classification (maybe w/ a different logit function)

Fitting: logistic regression (next class)

Predicting: if $p(x_{\text{new}}) > \text{threshold}$, return 1 else 0

Categorical regression (multiclass classification) model

A r.v. $z \sim \text{Categorical}(p_1, \dots, p_k)$ if

$$\mathbb{P}(z=i) = p_i \quad \text{for } i=1, \dots, k$$

where $p_i \geq 0$ and $\sum p_i = 1$

We use this dist in ML for classification

Assumption: $y|x \sim \text{Categorical}(p_1(x), \dots, p_k(x))$

s.t. $p_i(x) \geq 0$ and $\sum p_i(x) = 1$

To ensure this, take

$$p_i(x) = \frac{e^{\theta_i^T x}}{\sum_{j=1}^k e^{\theta_j^T x}}$$

logit of class i

$$= \frac{e^{\theta_i^T x}}{z_{\theta}(x)}$$

$z_{\theta}(x) = \sum_{j=1}^k e^{\theta_j^T x}$
is called the
partition function or
normalizing constant

Fitting: multinomial logistic regression

Predicting: given x_{new} , predict $y_{\text{pred}} = \arg\max_i p(x_{\text{new}})_i$

Some notation:

- Instead of letting $y = 1, \dots, k$, encode the classes with one-hot encoding $y = e_1, \dots, e_k$ where $e_i = \begin{bmatrix} 0 \\ \vdots \\ 1 \\ \vdots \\ 0 \end{bmatrix} \in \mathbb{R}^k$ 
- Given a vector x , define $e^x = \begin{bmatrix} e^{x_1} \\ \vdots \\ e^{x_k} \end{bmatrix}$
- Given $\theta_i \in \mathbb{R}^d$ for $i = 1, \dots, k$, define $\Theta = \begin{bmatrix} \theta_1^T \\ \vdots \\ \theta_k^T \end{bmatrix} \in \mathbb{R}^{k \times d}$ as our parameter matrix
- Let $\mathbf{1} = \begin{bmatrix} 1 \\ \vdots \\ 1 \end{bmatrix} \in \mathbb{R}^k$ denote the vector of all ones

Now we can write $p(x) = p_{\theta}(x)$

$$p_{\theta}(x) = \begin{bmatrix} p_1(x) \\ \vdots \\ p_k(x) \end{bmatrix} = \begin{bmatrix} e^{\theta_1^T x} / z_{\theta}(x) \\ \vdots \\ e^{\theta_k^T x} / z_{\theta}(x) \end{bmatrix} = \frac{1}{z_{\theta}(x)} \begin{bmatrix} e^{\theta_1^T x} \\ \vdots \\ e^{\theta_k^T x} \end{bmatrix}$$

$$= \frac{e^{\theta x}}{z_{\theta}(x)} = \frac{e^{\theta x}}{\mathbf{1}^T e^{\theta x}} = \text{softmax}(\theta x)$$

$$(\text{softmax}: v \rightarrow \frac{e^v}{\mathbf{1}^T e^v})$$

Ex

$$\begin{bmatrix} 30 \\ 1 \\ 0 \end{bmatrix} \xrightarrow{\text{softmax}} \begin{bmatrix} 1 \\ 2.54 \times 10^{-13} \\ 9.36 \times 10^{-14} \end{bmatrix}$$

$$\begin{bmatrix} 0.9 \\ 0.05 \\ 0.05 \end{bmatrix} \xrightarrow{\text{softmax}} \begin{bmatrix} 0.539 \\ 0.23 \\ 0.23 \end{bmatrix}$$

Prediction: $y_{\text{pred}} = \underset{i}{\operatorname{argmax}} (\theta x)_i$

Notice

$$\mathbb{E}[y|x] = \sum_{i=1}^k e_i [p_{\theta}(x)]_i = p_{\theta}(x)$$