

Linear Regression

Model: $f(\mathbf{x}) = \mathbf{w}^\top \mathbf{x}$ (+ w_0 with bias).

OLS: $L(\mathbf{w}) = \|\mathbf{X}\mathbf{w} - \mathbf{y}\|^2$ ($\mathbf{X} \in \mathbb{R}^{n \times d}$) is quadratic $\Rightarrow$ cvx $\Rightarrow$ any local min is global.

Normal eq. $\mathbf{X}^\top \mathbf{X}\mathbf{w} = \mathbf{X}^\top \mathbf{y}$: always solvable, minimizer set $\mathcal{W} \neq \emptyset$; $|\mathcal{W}| \in \{1, \infty\}$.

Unique $\Leftrightarrow \mathbf{X}^\top \mathbf{X}$ invertible $\Leftrightarrow \text{rank } \mathbf{X} = d$ (full col. rank: $n \geq d$ & no collinear cols; $n \geq d$ alone *not* enough): $\hat{\mathbf{w}} = (\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{y}$.

Else ($\mathbf{X}^\top \mathbf{X}$ singular): ∞ many = $\hat{\mathbf{w}}$ + null($\mathbf{X}^\top \mathbf{X}$). $n < d \Rightarrow \mathbf{X}^\top \mathbf{X}$ singular, and $\exists \mathbf{w}: \mathbf{X}\mathbf{w} = \mathbf{y}$ (interpolates).

Adding features: train. loss never $\uparrow$.

1D ($w_0 + w_1x$): $w_1^* = \frac{\text{Cov}(x,y)}{\text{Var}(x)}$, $w_0^* = \bar{y} - w_1^* \bar{x}$.

Stat model: $\mathbf{y} = \mathbf{X}\mathbf{w}^* + \varepsilon$, $\varepsilon \sim \mathcal{N}(0, \sigma^2 I) \Rightarrow \hat{\mathbf{w}} \sim \mathcal{N}(\mathbf{w}^*, \sigma^2(\mathbf{X}^\top \mathbf{X})^{-1})$ (unbiased).

Loss Functions: **Sqrd:** outliers;

Abs: robust; not diff. at 0.

Huber: ($\ell_H/\ell_{\text{abs}} \rightarrow \delta$); robust, smooth.

Optimization

GD (batch, all N): $\mathbf{w}^{t+1} = \mathbf{w}^t - \eta \nabla f(\mathbf{w}^t)$; for smooth f , small enough $\eta > 0$ always $\downarrow f$. Exact but $O(N)$ /step

SGD (1 example): unbiased est of ∇L , but noisy: can $\uparrow$ the loss, leave a stationary pt., be large even at an ε -stationary pt. of L . $O(1)$ /step $\Rightarrow$ cheap; **Mini-batch SGD:** compromise; cost $\approx$ SGD if parallelized.

Momentum: ($\beta \in \mathbb{R}$):

$\mathbf{w}^{t+1} - \mathbf{w}^t = \beta(\mathbf{w}^t - \mathbf{w}^{t-1}) - \eta \nabla_{\mathbf{w}} L(\mathbf{w}^t)$

Convexity (cvx): f cvx $\Leftrightarrow D^2 f \succeq 0$ (PSD) $\forall \mathbf{w}$. If diff. & $\nabla f(\mathbf{w}) = 0 \Rightarrow \mathbf{w}$ global min. α -strong cvx $\Leftrightarrow D^2 f \succeq \alpha I$ ($\alpha > 0$, PD) $\Rightarrow$ unique global min.

Preserved: (f, g cvx); $\alpha f + \beta g$ ($\alpha, \beta \geq 0$); $\max(f, g)$; affine compos. $f(A\mathbf{x} + \mathbf{b})$.

Not preserved: $\min(f, g)$; $f \cdot g$; $f(g(\mathbf{x}))$ (cvx iff g affine OR f cvx & non-decr).

Cvx any $\|\cdot\|$; ReLU; linear/affine. **Strict:** $\|\cdot\|^2$ (*strongly*); e^{-z} , $\log(1+e^{-z})$ (*not strongly*).

Model Evaluation & Selection

True model: f^* , $Y = f^* + \epsilon$, $\text{Var}(\epsilon) = \sigma^2$

Gen. error: $L(\hat{f}; \mathbb{P}) = \mathbb{E}_{X,Y} \ell(\hat{f}(X), Y) = \mathbb{E}_X[(\hat{f}(X) - f^*(X))^2] + \sigma^2$ (estim. err. + noise.) BUT **Training err. underestimates gen. err.:** use indep. test set: **test err.** $\rightarrow L(\hat{f}; \mathbb{P})$ (LLN) **Pitfall:** select a model by test err. $\Rightarrow$ test err. no longer

unbiased est. of gen. err.

K-fold CV: $\text{CV}_K(M) = \frac{1}{K} \sum_k L(\hat{f}_{M,k}; \mathcal{D}_k)$; pick $\arg \min_M \text{CV}_K$; refit on all data, eval. on $\mathcal{D}_{\text{test}}$. $\uparrow K$: sub-models $\hat{f}_k \rightarrow$ full model, but each $\hat{f}_k$ less precise & **more** compute. LOOCV ($K=n$): max cost.

$\uparrow$ val. set $\Rightarrow$ lower-variance gen-err estimate; $\uparrow$ train set $\Rightarrow$ lower gen err. Train err $\leq$ val err usually, **not always**.

CV/LOOCV estimates the **method** M , not one fixed $\hat{f}$ (each fold a different f); deterministic if M is.

Regularization & Model Selection

$\downarrow$ model complexity $\Rightarrow \uparrow$ bias and gen. err. $\uparrow$ complexity $\Rightarrow \downarrow$ train err (always), but $\uparrow$ variance and gen. err. $\Rightarrow$ bias-variance tradeoff

Bias-Var Decomp: $\mathbb{E}[\text$

(image size)*c $\Rightarrow$ same # params
Channels/N $\Rightarrow$ params/N.
Regularization: **Dropout:** zero out activations with prob $1 - p$ at train. At test: **multiply weights by p .**
Weight decay = L2 reg: $\theta \leftarrow \theta(1 - \alpha) - \eta \nabla \mathcal{L}_1 \Leftrightarrow$ L2 with $\lambda = \frac{\alpha}{2\eta}$.
 Early stopping, batch norm.
Universal Approx.: 1 hidden layer of fixed width p : does **NOT** approximate all continuous functions. Need to let $p \rightarrow \infty$.
Clustering & PCA
K-Means: k centers $\mu_1, \dots, \mu_k$. Cost $\hat{R}(\mu) = \sum_{i=1}^n \min_{j \in \{1, \dots, k\}} \|x_i - \mu_j\|_2^2$. Each $\|\cdot\|_2^2$ cvx, but $\arg \min_{\mu} \hat{R}$ non-cvx, NP-hard (matrix factor. $X \approx WZ$, Z one-hot cols). Assign $z_i = \arg \min_j \|x_i - \mu_j\| \Rightarrow$ Voronoi partition of $\mathbb{R}^d$; model-based (extends to new points), unlike hierarchical/spectral.
Lloyd's: alternate (1) assign z_i ; (2) $\mu_j \leftarrow$ mean of cluster j . $O(nkd)$ /iter. Cost never $\uparrow \Rightarrow$ converges, but to a **local** min only.
Init: uniform (bad if unbalanced); furthest-point (outlier-sensitive); **k-means++:** 1st center uniform, then $p(x) \propto \min_j \|x - \mu_j\|_2^2$; expected cost within $O(\log k)$ of optimal.
Choosing k : train & val. cost both $\downarrow$ monotonically $\Rightarrow$ CV unusable; use elbow.
PCA: **Centered** data ($\sum_i x_i = 0$). Linear, model-based dim. reduction $x_i \in \mathbb{R}^d \rightarrow z_i \in \mathbb{R}^k$, $k \ll d$ (manifold hypothesis). $W^*, z_1^*, \dots, z_n^* = \arg \min_{W^\top W=I, z_i} \frac{1}{n} \sum_i \|x_i - Wz_i\|_2^2$. Min. recon. error $\equiv$ max. proj. var. $w^\top \Sigma w$, $\Sigma = \frac{1}{n} \sum_i x_i x_i^\top = \frac{1}{n} X^\top X = \sum_{i=1}^d \lambda_i v_i v_i^\top$ (sym., PSD, spectral decomp.).
Solution: $W^* = [v_1 | \dots | v_k]$ top- k eigenvectors of Σ ($\lambda_1 \geq \dots \geq \lambda_d \geq 0$); $z_i^* = W^{*\top} x_i$; recon = $W^* W^{*\top} x_i$. 1st PC = max-variance dir.; PCs orthonormal; sign of v_i ambiguous.
 Variance along PC $j = \lambda_j$. Recon. error $C^{(k)} = \sum_{j=k+1}^d \lambda_j$; $(C^{(k-1)} - C^{(k)}) = \lambda_k$, $C^{(d)}=0$. Var. explained = $\sum_{j \leq k} \lambda_j / \sum_j \lambda_j$.
 $P = WW^\top$ orthogonal projection: $P^2 = P$.
Choosing k : retain e.g. 95% var, or screeplot elbow; CV if downstream supervised.
SVD: $X = USV^\top \Rightarrow n\Sigma = X^\top X = V S^\top S V^\top$; top- k PCs = first k cols of V

(right singular vectors); numerically stable. PCA = best rank- k approx of X (Eckart-Young). Duplicate feature $\Rightarrow$ reduce dim by ≥ 1 , zero recon. error.
Linear autoencoder (identity act., sq. recon. loss, bottleneck m) $\equiv$ PCA- m (Eckart-Young); depth irrelevant (linear act.).
PCA vs. k-means: both $\min \sum_i \|x_i - Wz_i\|_2^2$; differ by constraint on z_i — PCA $W^\top W=I$, $z_i \in \mathbb{R}^k$; k-means $z_i \in \{e_1, \dots, e_k\}$.
Probabilistic Modeling
1. Model: **Discr.:** model only $p(y | x; \theta)$ — regr.: $Y | X=x \sim \mathcal{N}(f(x; \theta), \sigma^2)$; class.: $p(y | x; \theta) = \sigma(y\theta^\top x)$ (logistic; predict $Y=1$ iff $\theta^\top x > 0$). **Gen.:** model $p(x, y) = p(y)p(x | y)$ — NB: $\hat{y} = \arg \max_k p(k) \prod_i p(x_i | k)$ (cond. indep. given class).
GBC hierarchy ($\Sigma_y := \text{Cov}(x | Y=y)$): **QDA** per-class full Σ_y , unrestricted $\Rightarrow$ quadratic boundary; **LDA** $\Sigma_{+1} = \Sigma_{-1}$ shared $\Rightarrow$ linear; **Fisher's LDA** = LDA + $P(y) = \frac{1}{2}$; **GNB** features of x i.i.d. given $y \Rightarrow \Sigma_y = \text{diag}(\sigma_{y,1}^2, \dots)$.
2. Inference (fit θ): **MLE:** $\hat{\theta} = \arg \max_{\theta} \log p(\mathcal{D} | \theta)$. **MAP:** $\arg \max_{\theta} [\log p(\mathcal{D} | \theta) + \log p(\theta)]$. **Bayes:** full posterior $p(\theta | \mathcal{D}) \propto p(\mathcal{D} | \theta)p(\theta)$. MAP = MLE iff uniform prior (Gaussian prior, $\sigma \rightarrow \infty$: MAP $\rightarrow$ MLE). MLE $\neq$ posterior mean in general. MLE minimizes $D_{\text{KL}}(\hat{\pi} || \nu)$ (empirical vs. model; $D_{\text{KL}} \geq 0$, not sym.).
Neg. log-lik $\equiv$ loss: Gaussian noise $\Rightarrow$ squared loss; logistic model $\Rightarrow$ logistic regr. MAP Gaussian prior $\Rightarrow$ ridge; Laplace prior $\Rightarrow$ LASSO.
Known-MLEs: Bern: $\hat{\theta} = \bar{x}$; Gaus: $\hat{\mu} = \bar{x}$, $\hat{\Sigma} = \frac{1}{N} \sum (x_i - \hat{\mu})(x_i - \hat{\mu})^\top$; Exp: $\hat{\lambda} = \frac{N}{\sum x_i}$; Poisson: $\hat{\lambda} = \bar{x}$; Unif(0, θ): $\hat{\theta} = \max_i x_i$.
Conjugate prior: Beta-Bernoulli — $\theta \sim \text{Beta}(\alpha, \beta)$, n_1 heads, n_0 tails $\Rightarrow \theta | \mathcal{D} \sim \text{Beta}(\alpha + n_1, \beta + n_0)$; Beta(1, 1) = Unif[0, 1].
3. Decision (predict): **Bayes optimal predictor:** $f^*(x) = \arg \min_a \mathbb{E}[\ell(a, Y) | X=x]$ — best possible (no function class, min. gen. err.). Square loss $\Rightarrow$ cond. mean $\mathbb{E}[Y | X=x]$; 0-1 loss $\Rightarrow$ posterior mode

$\arg \max_a p(a | x)$. With learned $\hat{\mathbb{P}}$: only an *estimate* of the optimal rule.
Expectation Maximization & GMMs
GMM: $p(x) = \sum_{k=1}^K \pi_k \mathcal{N}(x; \mu_k, \Sigma_k)$; $\sum_k \pi_k = 1$. Latent: $z_i \sim \text{Cat}(\pi)$; $x_i | z_i = k \sim \mathcal{N}(\mu_k, \Sigma_k)$.