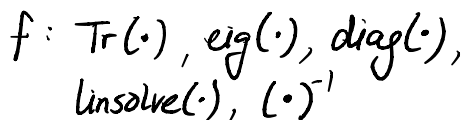
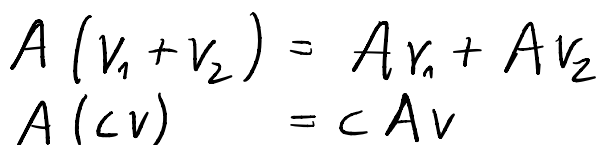
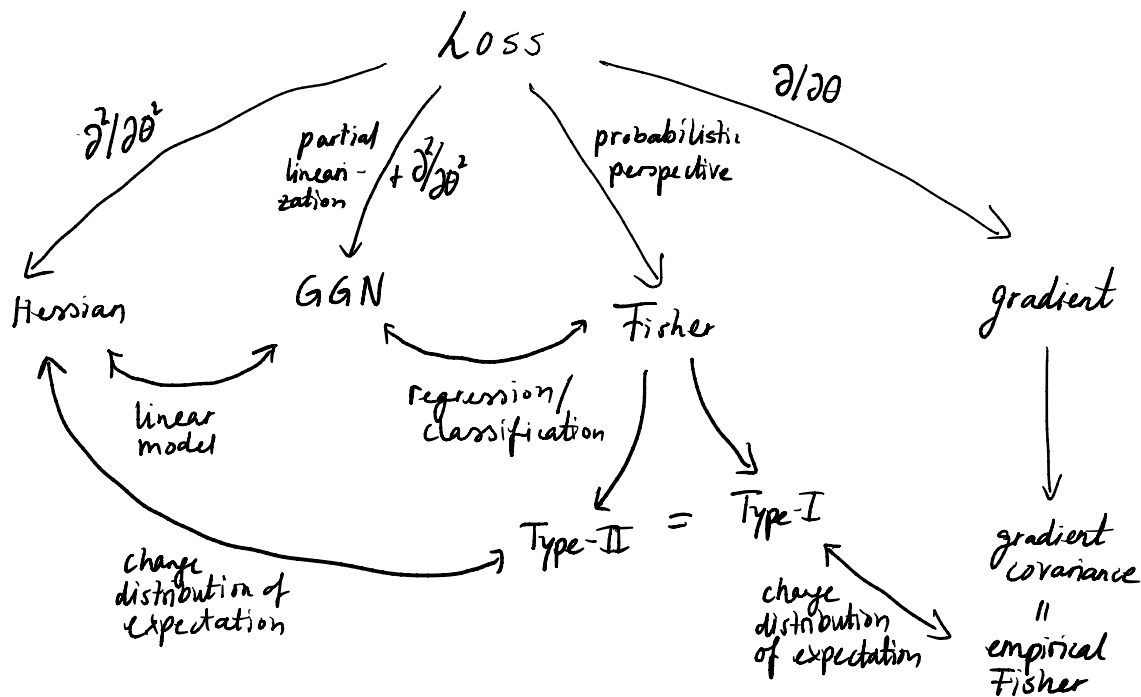
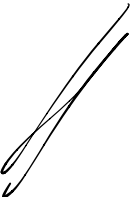


Overview: linear operators



Overview: DL curvature matrices and their relation





DL

Curvature

matrices



Loss, gradient, Hessian

$$\begin{array}{ccc} \theta \in \mathbb{R}^D & & y \\ \downarrow & & \downarrow \\ x \longrightarrow f \in \mathbb{R}^C & \longrightarrow & \ell \end{array} \quad + \quad \mathcal{D} = \left\{ (x_n, y_n) \right\}_{n=1}^N$$

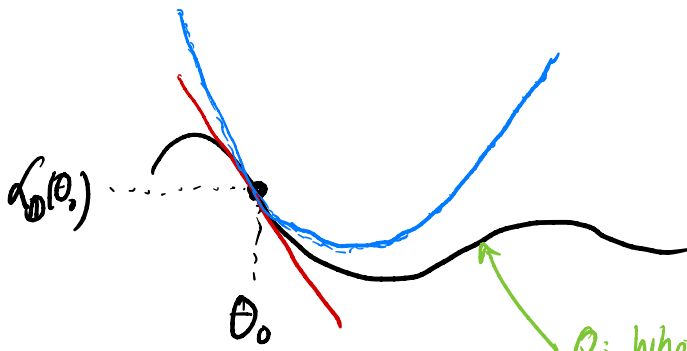
$$\mathcal{L}_{\mathcal{D}}(\theta) = \frac{1}{N} \sum_{n=1}^N \ell(f(x_n, \theta), y_n) \in \mathbb{R}$$

$$g_{\mathcal{D}}(\theta) = \frac{1}{N} \sum_{n=1}^N \nabla_{\theta} \ell(f(x_n, \theta), y_n) \in \mathbb{R}^D$$

$$H_{\mathcal{D}}(\theta) = \frac{1}{N} \sum_{n=1}^N \nabla_{\theta}^2 \ell(f(x_n, \theta), y_n) \in \mathbb{R}^{D \times D}$$

Motivation: Taylor expansion

$$\begin{aligned} \mathcal{L}_{\mathcal{D}}(\theta - \theta_0) &= \mathcal{L}_{\mathcal{D}}(\theta_0) + g_{\mathcal{D}}(\theta_0)^T (\theta - \theta_0) \\ &\quad + \frac{1}{2} (\theta - \theta_0)^T H_{\mathcal{D}}(\theta_0) (\theta - \theta_0) + o((\theta - \theta_0)^3) \end{aligned}$$



Q: What happens here if we apply Newton's method?

Use cases:

- Optimization (fit local quadratic, jump to minimum)
- Pruning (find weight that least affects loss when removed)
- Sharpness & influence functions (sensitivity of loss w.r.t. perturbations in the data/parameters)
- Laplace approximations (fit a Gaussian to a more complicated distribution)

Pros and cons of the Hessian:

- Cannot touch the full matrix explicitly ($D \times D$)
- Can multiply with the Hessian efficiently using first-order autodiff! (double-backward)

$$\left(\frac{\partial^2}{\partial \theta^2} \mathcal{L}\right) v = \frac{\partial}{\partial \theta} \left[\left(\frac{\partial}{\partial \theta} \mathcal{L}\right) \cdot v \right] = \frac{\partial}{\partial \theta} [g \cdot v]$$

- Hessian can be indefinite because $\underset{\substack{\uparrow \\ \text{convex}}}{\ell} \circ \underset{\substack{\uparrow \\ \text{non-convex}}}{f}$ is usually non-convex

Generalized Gauss-Newton (GGN) matrix

- Compositions of convex functions are convex
- can we convexify f ? Yes, through linearization around some anchor θ_0 :

$$f(\theta) \rightarrow f_{\theta_0}^{\text{lin}}(\theta) = f(\theta_0) + \underbrace{J_{\theta_0} f(\theta_0)}_{C \times D} (\theta - \theta_0)$$

$[\cdot]_{i,j} = \frac{\partial f_i}{\partial \theta_j}$

$$H_{\theta}(\theta) = \frac{1}{N} \sum_{n=1}^N \nabla_{\theta}^2 (\ell_n \circ f_n^{\text{lin}})(\theta)$$

Hessian of the partially linearized loss

$$G_{\theta}(\theta) = \frac{1}{N} \sum_{n=1}^N \left[\nabla_{\theta}^2 (\ell_n \circ f_{\theta_0,n}^{\text{lin}})(\theta) \right] \Big|_{\theta_0=\theta}$$

"GGN"

$$= \frac{1}{N} \sum_{n=1}^N \underbrace{(\nabla_{\theta} f_n)^T}_{D \times C} \underbrace{(\nabla_{f_n}^2 \ell_n)}_{C \times C} \underbrace{(\nabla_{\theta} f_n)}_{C \times D}$$

- Positive semi-definite
- GGN-vector products through JVP, HVP, VJP
- = Hessian if f is linear in θ (e.g. linear regression)

// Prelude to Fisher: Risk minimization

$$\mathcal{L}(\theta) = \frac{1}{N} \sum_{n=1}^N \ell(f(x_n, \theta), y_n) \quad \left(\begin{array}{l} \text{massage into} \\ \text{an expectation} \end{array} \right)$$

Assume $\exists$ data-generating process $p_{\text{data}}(x, y)$
and we want to minimize the expected risk

$$\mathcal{L}_{p_{\text{data}}}(\theta) = \mathbb{E}_{p_{\text{data}}(x, y)} [\ell(f(x, \theta), y)] \quad \text{w.r.t. } \theta$$

$\hookrightarrow p_{\text{data}}$ is intractable though. That's why we collected a data set, which approximates
 $p_{\text{data}}(x, y) \approx p_{\mathcal{D}}(x, y) := \frac{1}{N} \sum_{n=1}^N \delta(x - x_n) \delta(y - y_n)$

↑
"empirical data distribution", tractable

We then have

$$\mathcal{L}_D(\theta) = \mathcal{L}_{P_\theta}(\theta) = \mathbb{E}_{P_D(x,y)} [\ell(f(x,\theta), y)]$$

for the empirical risk.

// Connection to maximum likelihood

Want to learn a parametrized density $P_\theta(x,y)$ that approximates $P_{data}(x,y)$:

$$\underset{\theta}{\text{minimize}} \quad KL(P_{data}(x,y) \parallel P_\theta(x,y))$$

$$\Leftrightarrow \underset{\theta}{\text{minimize}} \quad \mathbb{E}_{P_{data}(x,y)} [-\log P_\theta(x,y)]$$

(we will only model $y|x$ with parameters, so $P_\theta(x,y) = P_\theta(y|x) P_{data}(x)$, we don't have access to P_{data} , so $P_{data} \hookrightarrow P_D$)

$$\underset{\theta}{\text{minimize}} \quad \frac{1}{N} \sum_{n=1}^N \underbrace{-\log p_{\theta}(y_n | x_n)}$$

This looks like $\mathcal{L}(f(x_n, \theta), y_n)$
from above, where $\mathcal{L}$ is a
negative log-likelihood

Rewrite $p_{\theta}(y|x) = r(y|f(x, \theta))$, then we
can show

$$1) \quad r(y|f(x, \theta)) = \mathcal{N}(y | \mu=f(x, \theta), \Sigma=I)$$

$$\Leftrightarrow -\log r = \mathcal{L}(\text{MSELoss})$$

$$2) \quad r(y|f(x, \theta)) = \text{Cat}\left(y \mid \pi = \text{softmax}(f(x, \theta))\right)$$

$\in \{1, \dots, C\}$ $\uparrow$

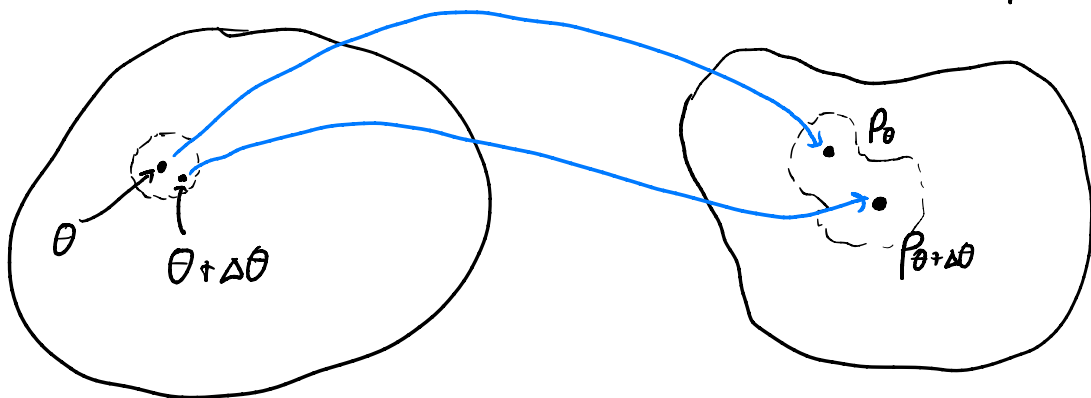
$$\Leftrightarrow -\log r = \mathcal{L}(\text{CrossEntropyLoss})$$

$\Rightarrow$ The NN models a Gaussian/categorical distribution

Probabilistic perspective

parameter space

statistical manifold



$$d(\theta, \theta + \Delta\theta) = \|\Delta\theta\|_2^2$$

$$\begin{aligned} d(p_{\theta + \Delta\theta}, p_\theta) &= \left(\text{second-order approximation of} \right) \\ &= \text{KL}(p_{\theta + \Delta\theta} \| p_\theta) \\ &= -\frac{1}{2} \mathbb{E}_{p_\theta(x,y)} \left[\nabla_\theta^2 \log p_\theta(x,y) \right] \end{aligned}$$

"Fisher information matrix"

(use $p_\theta(x,y) = p_\theta(y|x) p_\theta(x)$)

$$F(\theta) = \frac{1}{N} \sum_{n=1}^N \mathbb{E}_{p_\theta(y|x_n)} \left[-\nabla_\theta^2 \log p_\theta(y|x_n) \right]$$

$$= \frac{1}{N} \sum_{n=1}^N \mathbb{E}_{p_\theta(y|x_n)} \left[-\nabla_\theta^2 \ell(f(x_n, \theta), y) \right]$$

looks similar to Hessian

$$= \frac{1}{N} \sum_{n=1}^N (\nabla_\theta f_n)^T \mathbb{E}_{p_\theta(y|x_n)} \left[\nabla_f^2 \ell(f_n, y) \right] \nabla_\theta f_n$$

Type-II
looks like GGN

$$= \frac{1}{N} \sum_{n=1}^N (\nabla_\theta f_n)^T \mathbb{E}_{p_\theta(y|x_n)} \left[\nabla_f \ell(f_n, y) (\nabla_f \ell(f_n, y))^T \right] \nabla_\theta f_n$$

Type-I
motivation for ET

Similarities and differences

$$(l = -\log p)$$

$$H(\theta) = \mathbb{E}_{p_{\mathcal{D}}(x,y)} \left[\nabla_{\theta}^2 l \right] = \mathbb{E}_{p_{\mathcal{D}}(x)} \mathbb{E}_{p_{\mathcal{D}}(y|x)} \left[\nabla_{\theta}^2 l \right] \quad \text{Hessian}$$

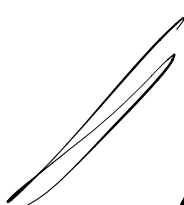
$$F(\theta) = \mathbb{E}_{p_{\mathcal{D}}(x,y)} \left[\nabla_{\theta}^2 l \right] \quad \text{Fisher}$$

$$F_{\text{II}}(\theta) = \mathbb{E}_{p_{\mathcal{D}}(x)} \left[(\mathcal{J}_{\theta} f)^T \mathbb{E}_{p_{\mathcal{D}}(y|x)} \left[\nabla_f^2 l \right] \mathcal{J}_{\theta} f \right] \quad \text{Fisher type-II}$$

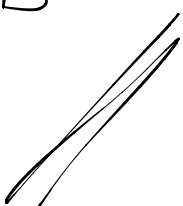
$$G(\theta) = \mathbb{E}_{p_{\mathcal{D}}(x)} \left[(\mathcal{J}_{\theta} f)^T \mathbb{E}_{p_{\mathcal{D}}(y|x)} \left[\nabla_f^2 l \right] \mathcal{J}_{\theta} f \right] \quad \text{GGN}$$

$$F_{\text{I}}(\theta) = \mathbb{E}_{p_{\mathcal{D}}(x)} \left[(\mathcal{J}_{\theta} f)^T \mathbb{E}_{p_{\mathcal{D}}(y|x)} \left[\nabla_f l (\nabla_f l)^T \right] \mathcal{J}_{\theta} f \right] \quad \text{Fisher type-I}$$

$$EF(\theta) = \mathbb{E}_{p_{\mathcal{D}}(x)} \left[(\mathcal{J}_{\theta} f)^T \mathbb{E}_{p_{\mathcal{D}}(y|x)} \left[\nabla_f l (\nabla_f l)^T \right] \mathcal{J}_{\theta} f \right] \quad \text{empirical Fisher}$$



Matrix-free
linear algebra
with linear operators



Hutchinson trace/diagonal estimation

$$\text{Tr}(A) = \text{Tr}(A I)$$

Trace

$$= \text{Tr}(A \mathbb{E}_{p(v)}[v v^T])$$

$$= \mathbb{E}_{p(v)}[\text{Tr}(A v v^T)]$$

$$= \mathbb{E}_{p(v)}[\text{Tr}(v^T A v)]$$

$$= \mathbb{E}_{p(v)}[v^T A v]$$

$$\approx \frac{1}{S} \sum_{s=1}^S v_s^T A v_s \quad \text{where } v_s \stackrel{\text{i.i.d.}}{\sim} p(v)$$

(can also use to estimate $\|A\|_2^2 = \text{Tr}(A^T A)$)

$$\text{diag}(A) = \text{diag}(A I)$$

$$= \text{diag}(A \mathbb{E}_{p(v)}[v v^T])$$

$$= \mathbb{E}_{p(v)}[\text{diag}(A v v^T)]$$

$$= \mathbb{E}_{p(v)}[(A v) \odot v]$$

$$\approx \frac{1}{S} \sum_{s=1}^S A v_s \odot v_s$$

Diagonal

where $v_s \stackrel{\text{i.i.d.}}{\sim} p(v)$

Power iteration

eig

$\lambda_1(A), e_1(A)$

$\left| \begin{array}{l} \lambda_1 \leftarrow \infty \\ e_1 \leftarrow v \\ \text{until converged} \end{array} \right.$

v random, $\|v\|_2 = 1$

$\left| \begin{array}{l} e_1 \leftarrow A e_1 \\ \lambda_1 \leftarrow \|e_1\|_2 \\ e_1 \leftarrow e_1 / \|e_1\|_2 \end{array} \right.$

matvec

update eigval

- (more sophisticated methods, e.g. implicitly restarted Arnoldi iterations (ARPACK), `scipy.sparse.linalg.eigsh`)
- Also: Spectral density estimation

$$\rho(\lambda) = \frac{1}{D} \sum_{d=1}^D \delta(\lambda - \lambda_d)$$

//

linear systems and inversion

$\text{linsolve}(A, b) = x$ such that $Ax = b$.

↳ Iterative solvers like conjugate gradients (CG)
(details not important, only uses $[v \mapsto Av]$).

//

CG inversion: compute $A^{-1}b = x$

$$\Leftrightarrow \text{Find } x : b = Ax$$

$$\Rightarrow x = \text{linsolve}(A, b)$$

//

Neumann series:

$$A^{-1} = \sum_{k=0}^{\infty} (I - A)^k \quad (\text{if } 0 < \lambda(A) < 2)$$

$$A^{-1}v = v + (I - A)v + (I - A)^2v + \dots \quad (\text{truncate})$$

curvlinops library: offers linear operators for H , $\tilde{F}_H, \tilde{F}_I, G, EF$, and more + functionality to estimate their properties