

Lecture 2-1

频率 interpretation: prob are long term relative frequencies

Bayesian interpretation: logically consistent degrees of belief 信念. 主观层面

$$L(H|E) = P(E|H) \quad L(M|D) = P(D|M) \leftarrow \text{model explains data 怎么样}$$

$$P(H|E) = \frac{P(H)P(E|H)}{P(E)} \propto P(H)L(H|E)$$

↑ 假设多个假设 $P(H_i|E)$ 分母都是同一个 $P(E)$ 可以忽略.

1. MLE. 假设 θ 固定. 完全依赖 Data 来找最能解释的参数. $m_h + m_t = m$

$$D = \{D_1, \dots, D_m\} \quad L(\theta|D) = \prod_{i=1}^m \theta^{m_h} (1-\theta)^{m_t} \quad \log\text{-下} \quad m_h \log \theta + m_t \log(1-\theta)$$

$$\text{对 } \theta \text{ 求导} \quad \frac{m_h}{\theta} - \frac{m_t}{1-\theta} = 0 \quad \Rightarrow \hat{\theta}_{MLE} = \frac{m_h}{m_t + m_h} = \frac{m_h}{m} \quad L(\theta|D) = \theta^{m_h} (1-\theta)^{m_t}$$

1.1 sufficient statistics: function 整合所有信息. 如 $st(D) = L$

(m_h, m_t) 是 sufficient statistic

MLE 无法提供置信度. 10 (3:7) 认为随机波动. 因为有强烈生活中 bias. 因该 coin 是 unbiased 的.

↓
Bayesian: 小数据时先验重要. 大数据时数据压倒一切.

{ MLE: $\theta^* = \arg \max P(D|\theta)$ 点估计. X 可信程度 X uncertainty
Bayesian: 把 θ 当成随机变量 得到的是 $p(\theta|D)$ 有 prior $p(\theta)$

$$\text{Bayesian: } P(D_{m+1} = H|D) = \int P(D_{m+1} = H, \theta|D) d\theta = \int P(D_{m+1} = H|\theta, D) p(\theta|D) d\theta$$

↑
在知道 θ 后, 与 D 无关.

$$\text{Bayesian MAP: } P(D_{m+1} = H|D) = \theta^* = \arg \max_{\theta} p(\theta|D) \quad \text{退化成} \quad \int \theta p(\theta|D) d\theta$$

↑ 简化成求后验最大的

$$\text{posterior distribution } p(\theta|D) \propto p(\theta)L(\theta|D) = \theta^{m_h} (1-\theta)^{m_t} p(\theta)$$

$$\text{有 Beta 分布做 } p(\theta) \text{ 先验} \quad p(\theta) = \frac{\Gamma(\alpha_t + \alpha_h)}{\Gamma(\alpha_t)\Gamma(\alpha_h)} \theta^{\alpha_h-1} (1-\theta)^{\alpha_t-1}$$

$$\Gamma(\alpha) = (\alpha-1)!$$

这里的 α_h 和 α_t 是虚构 (prior) 认为在 $\alpha_h + \alpha_t$ 次 game 中有 α_h 次 head α_t 次 tail.

$$\text{于是变成} \quad \theta^{m_h + \alpha_h - 1} (1-\theta)^{m_t + \alpha_t - 1}$$

α 越大, 对 prior 越自信

$$P(D_{m+1} = H|D) = \int \theta p(\theta|D) d\theta = c \int \frac{\theta^{m_h + \alpha_h}}{m + \alpha} d\theta$$

↑ 还是 Beta. Beta 期望 $E[\theta] = \frac{m_h + \alpha_h}{m + \alpha}$

Lecture 2-2

1. Jensen's Inequality

若 f 是定义在 I 上凹函数, $\forall x, y \in I, \lambda \in [0, 1]$ 满足 $\lambda f(x) + (1-\lambda)f(y) \leq f(\lambda x + (1-\lambda)y)$

对于概率分布 $p_i \in [0, 1], \sum_{i=1}^n p_i = 1, x_i \in I$.

$$\sum_{i=1}^n p_i f(x_i) \leq f\left(\sum_{i=1}^n p_i x_i\right) \quad \text{若 } f \text{ 强凹}$$

证明: 显然 $n=2$ 时成立. 那么扩展到 n . 有 $\sum_{i=1}^n p_i f(x_i) \leq f\left(\sum_{i=1}^n p_i x_i\right)$

$$\sum_{i=1}^{n+1} p_i f(x_i) = \sum_{i=1}^n p_i f(x_i) + p_{n+1} f(x_{n+1}) \quad \text{设 } \sum_{i=1}^n p_i = \lambda \in [0, 1].$$

$$\leq f\left(\sum_{i=1}^n p_i x_i\right)$$

变成 $1 \times \frac{1}{\lambda} \cdot \lambda$

$$= \lambda \sum_{i=1}^n \frac{1}{\lambda} p_i f(x_i) + p_{n+1} f(x_{n+1})$$

$$\leq \lambda f\left(\sum_{i=1}^n \frac{p_i}{\lambda} x_i\right) + p_{n+1} f(x_{n+1}) \Rightarrow \lambda + p_{n+1} = 1$$

$$\leq f\left(\lambda \sum_{i=1}^n \frac{p_i}{\lambda} x_i + (1-\lambda) x_{n+1}\right) = f\left(\sum_{i=1}^{n+1} p_i x_i\right)$$

特例. 对数函数 $\log(x)$ 在 $(0, +\infty)$ 严格凹. $\sum_{i=1}^n p_i \log(x_i) \leq \log\left(\sum_{i=1}^n p_i x_i\right)$

$$\text{即 } E[\log(X)] \leq \log E[X]$$

2. Entropy $H(X) = -\sum_{x \in X} P(x) \log P(x) = -E[\log P(X)]$ 约定 $0 \log 0 = 0$

性质: 1. $H(X) \geq 0$

2. $H(X) = 0$ iff $\exists x \in X, P(X=x) = 1$

3. $H(X) \leq \log |X|$

取值空间基数 iff 均匀分布. 用 Jensen.

$$H(X) = \sum_x P(x) \log \frac{1}{P(x)} \leq \log\left(\sum_x P(x) \frac{1}{P(x)}\right) = \log |X|$$

这里 f 是 $\log$ x_i 是 $\frac{1}{P(x)}$

条件熵

$$H(Y|X=x) = \sum_y P(y|x) \log \frac{1}{P(y|x)}$$

这可行

其是就是把 $P(x)$ 换成 $P(Y|X=x)$

$H(Y|X=x)$ 可能大于 $H(Y)$ 因为 $X=x$ 既可能带来确定性.

也可能带来不确定性.

$H(Y|X) = \sum_x P(x) H(Y|X=x)$ 平均条件熵: 已知 X 后, Y 剩余的平均不确定性.

$$H(X, Y) = -\sum_{x,y} P(x,y) \log(P(y)P(x|y)) = H(Y) + H(X|Y)$$

$$H(X) + H(X|Y) = -\sum_y \log P(Y) + \sum_y P(y) H(X|Y=y)$$

联合熵

$$H(X, Y) = -\sum_{x,y} P(x,y) \log P(x,y) = -\sum_y P(y) \log P(y) + \sum_{x,y} P(y) \sum_x P(x|Y=y) \log \frac{1}{P(x|Y=y)}$$

3. Divergence 衡量分布差异

1. $KL(P||Q) = \sum_x P(x) \log \frac{P(x)}{Q(x)}$; 若 $P(x) > 0$ 且 $Q(x) = 0$. 则 $KL = \infty$

$KL(P||Q) \geq 0$. $\log x$ 有负号. 所以大于等于

proof $\sum_x P(x) \log \frac{P(x)}{Q(x)} = - \sum_x P(x) \log \frac{Q(x)}{P(x)} \geq - \log \left(\sum_x P(x) \frac{Q(x)}{P(x)} \right)$

Cross entropy $H(P, Q) = - \sum_x P(x) \log Q(x) = -E[\log Q(X)]$
分布

$KL(P||Q) = E[\log P(X)] - E[\log Q(X)] = H(P, Q) - H(P)$

推理: $H(P, Q) \geq H(P)$ 因为 $KL(P||Q) \geq 0$.

即 $\sum_x P(x) \log Q(x) \leq \sum_x P(x) \log P(x)$ 负号.

令 $f(x)$ 为非负函数. 有 $\sum_x f(x) \log Q(x) \leq \sum_x f(x) \log P(x)$

$P^*(x) = f(x) / \sum_x f(x)$ $\uparrow$ 可以归一化回去

2. JS Divergence. $JS(P||Q) = \frac{1}{2} KL(P||M) + \frac{1}{2} KL(Q||M)$

其中 $M = \frac{P+Q}{2}$

对称. 有界 $[0, \log 2]$

$JS(P||Q) = 0 \Leftrightarrow P = Q$

$JS(P||Q) = \log 2$ 当 P 和 Q 完全不相交.

$\uparrow$ 算 $P(x)$ 时. $m(x)$ 退化为 $\frac{1}{2}P(x)$ 整个退化为 $\sum_x P(x) \log 2 = \log 2$.

非 supervised-learning: 最小化 $KL(P||Q) = H(P, Q) - H(P)$

$H(P, Q) = - \int P(x) \log Q(x) dx \approx - \frac{1}{N} \sum_{i=1}^N \log Q(x_i) = - \frac{1}{N} \log Q(D)$ $\uparrow$ 固定.

supervised-learning. $P(x, y) \xrightarrow{\text{sample}} D \xrightarrow{\text{学}} Q(y|x)$ $\uparrow$ 最大化 $Q(D)$

希望 缩小误差, 即 $H(P, Q)$ 变成 最大化 $\sum_{i=1}^N \log Q(y_i|x_i)$

4. Mutual information

$I(X; Y) = H(X) - H(X|Y)$ 在得知 Y 后对 X 减少的不确定度

$= \sum_x P(x) \log \frac{1}{P(x)} - \sum_{x,y} P(x,y) \log \frac{1}{P(x,y)} = \sum_{x,y} P(x,y) \frac{P(x,y)}{P(x)} = \sum_{x,y} P(x,y) \frac{P(x)}{P(y)}$
 $\uparrow$
 $P(x) = \sum_y P(x,y)$ $= KL(P(X, Y) || P(X)P(Y))$
对称

$= H(Y) - H(Y|X)$ 从 KL 性质来

$\Rightarrow I(X, Y) \geq 0 = 0 \text{ iff } X \perp Y$ $H(X|Y) \leq H(X) \text{ iff } X \perp Y$

Lecture 2-3

$X, Y, c: X \rightarrow Y$: target concept to learn C : concept class

D, S, H : set of hypotheses
 $\uparrow$ target distribution training samples

$\downarrow$ hypothesis, 学到的

$$1. \begin{cases} \text{泛化误差} & R(h) = \Pr_{x \sim D} [h(x) \neq c(x)] = E_{x \sim D} [I_{h(x) \neq c(x)}] \\ \text{经验误差} & \hat{R}_S(h) = \frac{1}{m} \sum_{i=1}^m I_{h(x_i) \neq c(x_i)} \end{cases}$$

$$E_{S \sim D^m} [\hat{R}_S(h)] = R(h) \quad \text{经验期望为泛化}$$

2. Probably Approximately Correct Learning (PAC)

如果说 concept class C 是 PAC-learnable, 那么 $\exists$ 算法 A 和 $\text{poly}(\dots, \dots)$

s.t. $\forall \epsilon, \delta > 0, \forall D$ on $X, \exists \forall c \in C$ 都有

$$m \geq \text{poly}(1/\epsilon, 1/\delta, n, \text{size}(c)):$$

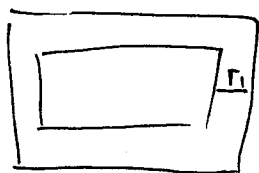
$$\text{等价于 } \Pr[R(h) > \epsilon] \leq \delta$$

$$\Pr_{S \sim D^m} [R(h) \leq \epsilon] \geq 1 - \delta$$

当算法 A 存在, 则称为一个 PAC learning

最多错误的概率大于 $1 - \delta$ also for C

2.1 若 $H = C$ 即假设空间包含 target class



希望 $\Pr(r_i) \leq \frac{\epsilon}{4}$ 这样 $\Pr[U_i; r_i] \leq \sum_i \Pr[r_i]$

Bad case 有些超过 $\frac{\epsilon}{4}$

放大了, 有重叠

$$\Pr[R(h) > \epsilon] = \Pr[U_{i=1}^m \Pr[r_i] > \epsilon/4] \leq 4(1 - \epsilon/4)^m \leq 4e^{-m\epsilon/4}$$

某个大 strip 一个样本没进去的 prob 小于 $1 - \frac{\epsilon}{4}$

$\uparrow$ strip

$$(1-x)^n \leq e^{-nx}$$

2.2 $H \neq C$ $R(h_S) = 0$ consistent: 在训练数据上能做到 $\hat{R}_S(h) = 0$ 即 $\epsilon_{RM} = 0$

$$\Pr_{S \sim D^m} [R(h) \leq \epsilon] \geq 1 - \delta \text{ holds if } m \geq \frac{1}{\epsilon} (\log |H| + \log \frac{1}{\delta})$$

$$R(h) \leq \frac{1}{m} (\log |H| + \log \frac{1}{\delta})$$

Proof: $\Pr[\exists h \in H: \hat{R}_S(h) = 0 \wedge R(h) > \epsilon]$

$$\leq \sum_{h \in H} \Pr[\hat{R}_S(h) = 0 \wedge R(h) > \epsilon]$$

单个样本不犯错 $\leq 1 - \epsilon$

$$= \sum_{h \in H} \Pr[\hat{R}_S(h) = 0 | R(h) > \epsilon] \Pr[R(h) > \epsilon]$$

放成 1

$$\leq |H| (1 - \epsilon)^m \leq |H| e^{-m\epsilon} = \delta \Leftrightarrow m \geq \frac{1}{\epsilon} (\log |H| + \log \frac{1}{\delta}) \Leftrightarrow R(h) \leq \frac{1}{m} (\log |H| + \log \frac{1}{\delta})$$

bound $O(\frac{1}{m})$

只要 hypothesis set finite. ERM 就是 PAC 算法

2.3. Inconsistent. 常见. $H \neq C$ 且 $R(hs) \neq 0$

Hoeffding's inequality. X_i 独立. $\in [a_i, b_i]$ $\Rightarrow \Pr(\sum X_i - E[\sum X_i] \geq \varepsilon) \leq e^{-2\varepsilon^2 / \sum (b_i - a_i)^2}$

显然 $E[\sum X_i]$ 就是 $R(h)$. 前者是 $\hat{R}(h)$ $\forall$ hypothesis $X \rightarrow \{0, 1\}$
代入 $a=0, b=1$

$$\Rightarrow \Pr[|\hat{R}(h) - R(h)| \geq \varepsilon] \leq 2e^{-2m\varepsilon^2} \Leftrightarrow 2e^{-2m\varepsilon^2} \leq \delta \Leftrightarrow |\hat{R}(h) - R(h)|$$

$$\forall h \in H, |\hat{R}(h) - R(h)| \leq \sqrt{\frac{\log |H| + \log \frac{2}{\delta}}{2m}} \leq \sqrt{\frac{\log \frac{2}{\delta}}{2m}}$$

$$\forall h \in H, R(h) \leq \hat{R}(h) + O\left(\sqrt{\frac{\log |H|}{m}}\right) \text{ error bound in } O\left(\frac{1}{\sqrt{m}}\right)$$

3. Infinite Hypothesis Set (之前 $\log |H| \rightarrow \infty$, 失效)

引入 a. growth function

b. VC dimension 用来 compute grow function

$\Pi_H: N \rightarrow N$ for a hypothesis H

$$\forall m \in N, \Pi_H(m) = \max_{\{x_1, \dots, x_m\} \subseteq X} |\{h(x_1), \dots, h(x_m)\}, h \in H|$$

最多的用在 H 中的 hypothesis h 把 m 个点分成不同方式的数

若 H 为二分类, 则 $\Pr[|R(h) - \hat{R}(h)| \geq \varepsilon] \leq 4 \Pi_H(2m) e^{-\frac{m\varepsilon^2}{8}}$

$$\text{即 } R(h) \leq \hat{R}(h) + \underbrace{\sqrt{\frac{2 \log \Pi_H(m)}{m}}}_{\text{复杂度}} + \underbrace{\sqrt{\frac{\log \frac{1}{\delta}}{2m}}}_{\text{置信度}}$$

$A \subseteq B$ 那么 $d_{VC}(A) \leq d_{VC}(B)$

3.1 VC Dimension:

① Dichotomy: 一种 label 点的方式 ② shattering 能够被打散说明 H 能实现任意 Dichotomy

$$VC \dim(H) = \max \{m: \Pi_H(m) = 2^m\}$$

$$VC \dim(R^d \text{ 的超平面}) = d+1$$

$$VC \dim(\sin) = +\infty$$

↑ 最大能被 H 打散
的点集大小

$$\Pi_H(m) = 2^m$$

3.2 Sauer's lemma: 令 H with $VC \dim(H) = d, \forall m \in N$ 这里 $m < d$

$$\Pi_H(m) \leq \underbrace{\sum_{i=0}^d \binom{m}{i}}_{\text{①}} \leq \sum_{i=0}^d \binom{m}{i} \left(\frac{m}{d}\right)^{d-i} = \binom{m}{d} \sum_{i=0}^d \binom{m}{i} \left(\frac{d}{m}\right)^i \stackrel{\text{②}}{\leq} \binom{m}{d} \sum_{i=0}^d \binom{m}{i} \left(\frac{d}{m}\right)^i = \left(\frac{m}{d}\right)^d \left(1 + \frac{d}{m}\right)^m \leq \left(\frac{em}{d}\right)^d$$

要么 $VC \dim(H) = d < +\infty$ and $\Pi_H(m) = O(m^d)$

or $VC \dim(H) = +\infty$ and $\Pi_H(m) = 2^m$

3.3 令 H 为 $\{-1, +1\}$ $VC \dim(H) = d, \forall \delta > 0$. 至少有 $1-\delta$ 的概率 $|R(h) - \hat{R}(h)| \leq \sqrt{\frac{8d \log \frac{2em}{d} + 8 \log \frac{4}{\delta}}{m}}$

3.4 For $|R(h) - \hat{R}(h)| \leq \varepsilon$ to hold for all $h \in H$ with p at least $1-\delta$. 满足 $m = O(d)$
模型越复杂, 数据量要求越大

DL Lecture 3

1. 从传统 $\rightarrow$ DL

传统 supervised $\{x_i, y_i\}_{i=1}^N \rightarrow P(y|x)$

DL: $\{x_i, y_i\}_{i=1}^N \rightarrow h = f(x), P(y|h) \Rightarrow P(y|x, \theta)$

多层线性复合仍线性

2. Weight Initialization (避免前向信号, GD BP 层间爆/消失)

Xavier $w \sim U[-\frac{1}{\sqrt{n}}, \frac{1}{\sqrt{n}}]$ n 输入维度: 每层输入方差大致稳定

He $w \sim N[0, \frac{2}{n}]$ 或 $w \sim N(0, \frac{2}{n})$ 通常用于 ReLU. (负值截断, 需更大方差)
 $\hookrightarrow$ 标准差

3. Activation

3.1 sigmoid $g(z) = \sigma(z) = \frac{1}{1 + \exp(-z)}$ $\sigma'(z) = \sigma(z)(1 - \sigma(z))$

$z \gg 0$ or $z \ll 0$ $\sigma(z) \rightarrow 0/1$ 导数接近 0, vanishing gradient
不推荐作为 hidden 的 act. 适合 = 分类输出层

3.2 ReLU $g(z) = \max\{0, z\}$

pros: $z > 0, \sigma = 1$ 快训 cons: 若某 neuron 对所有样本都 $z \leq 0, \sigma = 0$ 无法更新

缓解: 把 bias 初始化为小的正数 dead neuron

3.2.1 ReLU variant: 负半边留 α

$g(z, \alpha) = \max\{0, z\} + \alpha \min\{0, z\}$ $\begin{cases} \alpha = -1 & \text{absolute value rectification} \\ \alpha = 0.01 & \text{Leaky ReLU} \\ \alpha \text{ 可学} & \text{Parametric ReLU} \end{cases}$

3.3 Hyperbolic Tangent (Tanh) 零中心, 但仍会 saturate

$g(z) = \tanh(z) = \frac{\exp(z) - \exp(-z)}{\exp(z) + \exp(-z)} \in [-1, 1]$

相比 sigmoid, gradient 通常更大. 但当 $|z|$ 很大, 仍饱和

3.4 Swish, Mish, GeLU

Swish: $g(z) = z \sigma(\beta z)$ β 固定或可学

Mish: $g(z) = z \tanh(\log(1 + \exp(z))) = z \tanh(S(z))$ softplus $S(z) = \log(1 + \exp(z))$

GeLU: $g(z) = z \phi(z)$ $\phi(z)$ 是标准高斯的 CDF

共特: 平滑, 负半非 0, 非单调, 上方无界, 一般比 ReLU 稳

3.5 Softplus $g(z) = \zeta(z) = \log(1 + \exp(z))$

4. 输出层概率模型 $\text{logit } z = W^T h + b$

4.1 Linear Gaussian 当 y 是实值 $p(y|x) = \mathcal{N}(y|z, I)$ $L(x, y, \theta) = -\log P(y|x, \theta) \propto \frac{1}{2} \|y - z\|_2^2$

$$\frac{\partial L}{\partial z_k} = z_k - y_k$$

4.2 Sigmoid 输出 当 $y \in \{0, 1\}$ 二分类

$$P(y|x) = \text{Bernoulli}(y|6(z)) = 6(z)^y (1-6(z))^{1-y}$$

$$L(x, y, \theta) = -[y \log 6(z) + (1-y) \log (1-6(z))]$$

Saturation 只有在 model 正确且非常确信时有

$$\begin{aligned} \frac{\partial L}{\partial z} &= \frac{\partial L}{\partial 6(z)} \frac{\partial 6(z)}{\partial z} = \left(-\frac{y}{6(z)} + \frac{1-y}{1-6(z)}\right) \frac{\partial 6(z)}{\partial z} \\ &= \left(-\frac{y}{6(z)} + \frac{1-y}{1-6(z)}\right) 6(z)(1-6(z)) \\ &= -y(1-6(z)) + (1-y)6(z) \\ &= 6(z) - y \end{aligned}$$

4.3 Softmax 输出 $\{1, 2, \dots, C\}$ 多分类

$$P(y=k|x) = \frac{\exp(z_k)}{\sum_i \exp(z_i)}$$

$$L = -\log P(y=k|x) \quad \frac{\partial L}{\partial z_k} = -1 + \frac{\log S}{S} \frac{\partial S}{\partial z_k}$$

$$\Rightarrow \frac{\partial L}{\partial z_k} = P(y=k) - I(y=k) \quad \text{但 } k \neq y_{\text{true}}$$

$$\text{或 } \frac{\partial L}{\partial z_k} = \hat{y}_k - I(y=k) \quad \hat{y}_k = P(y=k)$$

$$\begin{aligned} \frac{\partial L}{\partial z_m} &= \frac{1}{S} \exp(z_m) P(y=k) \\ &= P(y=m) \end{aligned}$$

$$J(\theta) = \frac{1}{N} \sum_{i=1}^N L(x_i, y_i, \theta) = \frac{1}{N} \sum_{i=1}^N -\log P(y_i | x_i, \theta)$$

$$\text{SGD} \quad \theta \leftarrow \theta - \alpha \frac{1}{|B|} \sum_{i \in B} \nabla L(x_i, y_i, \theta)$$

6. BP

$$z_k = \sum_j u_j W_{kj}, \quad u_j = g(z_j) \quad z_j = \sum_i u_i W_{ji}$$

$$\frac{\partial L}{\partial W_{kj}} = \frac{\partial L}{\partial z_k} \frac{\partial z_k}{\partial W_{kj}} = u_j \delta_k \quad \delta_k = \frac{\partial L}{\partial z_k}$$

$$\frac{\partial L}{\partial W_{ji}} = \sum_k \frac{\partial L}{\partial z_k} \frac{\partial z_k}{\partial W_{ji}} = \sum_k \frac{\partial L}{\partial z_k} \frac{\partial z_k}{\partial u_j} \frac{\partial u_j}{\partial z_j} \frac{\partial z_j}{\partial W_{ji}} \downarrow u_i$$

$$= \sum_k \delta_k W_{kj} g'(z_j) u_i = u_i (g'(z_j) \sum_k \delta_k W_{kj})$$

$$\text{定义 } \delta_j = g'(z_j) \sum_k \delta_k W_{kj}$$

$$\frac{\partial L}{\partial W_{ji}} = u_i \delta_j \quad \frac{\partial L}{\partial b_j} = \delta_j$$

输入激活 目标单元误差信号 从 L 到输出的信号
 W_{ji} 只和 u_i 相乘 输出 z_j

7. Optimization

7.1 SGD 只看当前 ∇ noise 大. 路径震荡. 狭长 valley 收敛慢

7.2 Momentum $g \leftarrow \frac{1}{m} \nabla_{\theta} \sum_i L(f(x^{(i)}; \theta), y^{(i)})$

SGD
 $v \leftarrow \alpha v - \epsilon g$ velocity v
 $\theta \leftarrow \theta + v$ exponentially decaying the average of past gradients
 $\alpha = 0.5, 0.9, 0.99$

7.3 AdaGrad $r \leftarrow r + g \circ g$
 $\Delta \theta \leftarrow -\frac{\epsilon}{\sqrt{r}} \circ g$ 对稀疏特征友好. 历史又小的参数 lr 大
 $\theta \leftarrow \theta + \Delta \theta$ cons: r 单增. $lr \downarrow$ 后期学不动

7.4 RMSProp: 遗忘太久远历史

$r \leftarrow \rho r + (1-\rho) g \circ g$ 指数滑动平均. 不会无限累积
 $\Delta \theta \leftarrow -\frac{\epsilon}{\sqrt{r}} \circ g$ 有 adaptive scaling 但无显式 momentum
 $\theta \leftarrow \theta + \Delta \theta$ noise 仍大

7.5 Adam: Momentum + RMSProp

$s \leftarrow \rho_1 s + (1-\rho_1) g$
 $r \leftarrow \rho_2 r + (1-\rho_2) g \circ g$ 类似 RMSProp $\epsilon = 0.001$ $\rho_1 = 0.9$ $\rho_2 = 0.999$ $\epsilon = 10^{-8}$
类似动量 $\hat{s} = \frac{s}{1-\rho_1^t}$ $\hat{r} = \frac{r}{1-\rho_2^t}$
 $\Delta \theta \leftarrow -\epsilon \frac{\hat{s}}{\sqrt{\hat{r}} + \epsilon}$ $\theta \leftarrow \theta + \Delta \theta$

8. LR Scheduling

固定 | step decay | exponential decay | cosine annealing | warm restart 周期降低再升 模拟重新开搜. 但有好 base

9. Overfit Dropout

$$\tilde{J}(\theta) = J(\theta) + \lambda \|\theta\|_2^2$$

9.1 Bagging 训多个投票. 费劲

9.2 Dropout 廉价 Bagging. $m_j \sim \text{Bernoulli}(p)$ $\tilde{u}_j = m_j u_j$
大网中随机训小网. 防止神经元间产生过强 co-adaption

convex $O(\frac{1}{\sqrt{T}})$ strongly convex $O(\frac{1}{T})$

DL Week 4 CNN

1. MLP处理图时参数爆炸. 丢失空间结构, 缺乏平移泛化能力
2. CNN基本架构 INPUT $\rightarrow$ CONV $\rightarrow$ RELU $\rightarrow$ POOL $\rightarrow$ FC
3. Convolution

3.1 连续卷积 function f, g

$$(f * g)(t) = \int_{-\infty}^{+\infty} f(\tau) g(t-\tau) d\tau = \int_{-\infty}^{+\infty} f(t-\tau) g(\tau) d\tau$$

输入信号 $\uparrow$ filter $\uparrow$ 输出位置

由 f 邻域加权求和, 权重来自 g

3.2 离散卷积 weighted smoothing

$$f * g(n) = \sum_{m=-\infty}^{+\infty} f[m] g[n-m] = \sum_{m=-\infty}^{\infty} f[n-m] g[m]$$

4. 2D conv 和 Cross-Correlation

4.1 Real 2D conv

$$Z(i, j) = (I * k)(i, j) = \sum_m \sum_n I(m, n) k(i-m, j-n) = \sum_m \sum_n I[i-m, j-n] k[m, n]$$

Image $\uparrow$ kernel

卷积核被翻转

4.2 Cross correlation

$$Z(i, j) = (I * k)(i, j) = \sum_m \sum_n I(i+m, j+n) k(m, n)$$

在CNN中, kernel 是可学习, 是否翻转不影响表达能力

5. CNN 三大优势
- ① Sparse Connectivity: 一个神经元只连接 receptive field 里的 region 计算可行
 - ② Parameter Sharing: 不同空间相同权重 是等变性前提
 - ③ Equivariance: 平稳等变性 先移再卷 = 先卷再移

Invariance 是 $f(g(x)) = f(x)$

$f(g(x)) = g(f(x))$

6. 卷积层公式

$$Z_{j,k}^{(i)} = \sum_{l,m,n} I_{j+m-1, k+n-1}^{(l)} K_{m,n}^{(i,l)} + b^{(i)}$$

feature map $\uparrow$ 是从1开始

是 channel $\uparrow$ i 是对应输出 channel 的

加入 Stride

$$Z_{j,k}^{(i)} = \sum_{l,m,n} I_{(j-1)s+m, (k-1)s+n}^{(l)} K_{m,n}^{(i,l)} + b^{(i)}$$

也是从1开始 对齐位置

$$H_{out} = \left\lfloor \frac{H + 2P - F_h}{s} \right\rfloor + 1$$

$$W_{out} = \left\lfloor \frac{W + 2P - F_w}{s} \right\rfloor + 1$$

参数量 $C_{in} F_h F_w C_{out} + C_{out}$

一个 filter 为 $C_{in} F_h F_w + 1$

7. 反向传播 $J(k, b)$ 设上一层传回梯度张量 G

$$G_{j,k}^{(i)} = \frac{\partial J(k, b)}{\partial Z_{j,k}^{(i)}}$$

7.1 对 kernel 梯度 $\frac{\partial J(k, b)}{\partial k_{j,k}^{(i,l)}} = \sum_{m,n} \frac{\partial J}{\partial Z_{m,n}^{(i)}} \frac{\partial Z_{m,n}^{(i)}}{\partial k_{j,k}^{(i,l)}} = \sum_{m,n} G_{m,n}^{(i)} I_{(m-1)s+j, (n-1)s+k}^{(l)}$

对 input $\frac{\partial J(k, b)}{\partial I_{p,q}^{(l)}} = \sum_i \sum_{m,n} K_{p-(m-1)s, q-(n-1)s}^{(i,l)} G_{m,n}^{(i)}$

对 bias $\frac{\partial J(k, b)}{\partial b^{(i)}} = \sum_{m,n} G_{m,n}^{(i)}$

$$k_{j,k}^{(i,l)} = k_{j,k}^{(i,l)} - \alpha \frac{\partial}{\partial k_{j,k}^{(i,l)}} J(k, b)$$

8. Pooling

- Max Pooling
- Average Pooling 平滑
- l_2 -norm pooling 强调能量
- prob weighted pooling 较少见

$$b^{(i)} = b^{(i)} - \alpha \frac{\partial}{\partial b^{(i)}} J(k, b)$$

- 作用 {
- ① 对小平移不敏感
 - ② 降低计算、显存要求
 - ③ 改善统计效率，关注是否存在某特征，而不是精确位置
 - ④ 处理可变尺寸输入

8.1 反传 upsampling operation that ~~increases~~ inverses the subsampling in the forward pass

max pooling. 只取获得 max response 区域

如 $\begin{bmatrix} 1 & 3 \\ 2 & 0 \end{bmatrix} \rightarrow 3$ 反梯为 $\begin{bmatrix} 0 & 1 \\ 0 & 0 \end{bmatrix}$

9. Data Norm; Batch Norm

9.1 Why?

9.1.1 特征尺度不同 $\Rightarrow$ GD 会对某些 weight 猛打正

9.1.2 梯度尺度不同 $\Rightarrow$ loss contour 狭长，梯度下降 zig-zag 训练变慢

令 $X = [x_{ij}]_{N \times D}$ $\mu_j = \frac{1}{N} \sum_{i=1}^N x_{ij}$ $\sigma_j^2 = \frac{1}{N} \sum_{i=1}^N (x_{ij} - \mu_j)^2$ $\hat{x}_{ij} = \frac{x_{ij} - \mu_j}{\sigma_j}$

每个特征 均值为 0，方差为 1

9.2 Batch Normalization 前层参数改后，后层输入分布也改 $\Rightarrow$ covariate shift. 不仅规范化原始输入，也规范化每一层输入/激活

$$\hat{x} = \frac{x - \mu}{\sigma}$$

$$y = \gamma \hat{x} + \beta$$

测试时 μ, σ 来自 moving average

是否恢复某些均值和尺度

Pros: ① 加速训练 ② 降低初始化敏感度 ③ 轻微正则化 ④ 稳定深层网络

Lec 5 RNN

1. Vanilla RNN

1.1 更新 $h^{(t)} = \tanh(b + W h^{(t-1)} + U x^{(t)})$

$$o^{(t)} = c + V h^{(t)}$$

$$\hat{y}^{(t)} = \text{softmax}(o^{(t)})$$

$$W, U, V, b, c, \theta_{em}$$

1.2 Loss Input $\{x^{(1)}, x^{(2)}, \dots, x^{(T)}\}$ $\{y^{(1)}, \dots, y^{(T)}\}$

$$= - \sum_{t=1}^T \log P(y^{(t)} | x_{1:t}) \quad \text{对所有 timestamp 求和}$$

1.3 BPTT $\nabla_w L, \nabla_u L, \nabla_v L, \nabla_b L, \nabla_c L, \nabla_{\theta_{em}} L$

$$\hat{y}_i^{(t)} = \frac{\exp(o_i^{(t)})}{\sum_j \exp(o_j^{(t)})} \quad L^{(t)} = -\log \hat{y}_{y^{(t)}}^{(t)}$$

time + 不同输出

$$\frac{\partial L^{(t)}}{\partial o_i^{(t)}} = \hat{y}_i^{(t)} - 1_{i, y^{(t)}} \quad \text{GT} \Rightarrow \text{预测 prob - 真实 one-hot 标签}$$

$\because o^{(t)} = c + V h^{(t)} \Rightarrow \nabla_c L = \sum_{t=1}^T \nabla_{o^{(t)}} L, \nabla_v L = \sum_{t=1}^T (\nabla_{o^{(t)}} L) (h^{(t)})^T$

$h^{(t)}$ 到总 loss 可以通过 $\begin{cases} \rightarrow o^{(t)} \rightarrow L^{(t)} \\ \rightarrow h^{(t+1)} \dots \rightarrow L^{(t+1)} \dots \end{cases}$

$$\nabla_{h^{(t)}} L = V^T \nabla_{o^{(t)}} L + W^T \delta^{(t+1)}$$

定义 $\delta^{(t)} \equiv \nabla_{a^{(t)}} L$

$\frac{\partial h^{(t)}}{\partial a^{(t)}} \Rightarrow \text{matrix}$

$\frac{\partial h^{(t)}}{\partial a^{(t)}} = \text{diag}(1 - (h^{(t)})^2) \Rightarrow \delta^{(t)} = \text{diag}(1 - (h^{(t)})^2) \nabla_{h^{(t)}} L$

tanh 是 element wise 导致的

chain rule

$\Rightarrow \delta^{(t)} = \text{diag}(1 - (h^{(t)})^2) [V^T \nabla_{o^{(t)}} L + W^T \delta^{(t+1)}]$ 且 $\delta^{(T+1)} = 0$ 边界条件

$\Downarrow \nabla_b L = \sum_{t=1}^T \delta^{(t)} \quad \nabla_w L = \sum_{t=1}^T \delta^{(t)} (h^{(t+1)})^T \quad \nabla_v L = \sum_{t=1}^T \delta^{(t)} (x^{(t)})^T$

1.4 Teacher forcing: 把 $y^{(t)}$ 作为输入

output recurrence $\log P(y^{(1)}, y^{(2)} | x^{(1)}, x^{(2)}) = \log P(y^{(2)} | y^{(1)}, x^{(1)}, x^{(2)}) + \log P(y^{(1)} | x^{(1)}, x^{(2)})$

pros: $\begin{cases} \text{① 时间步之间完全 decouple} \Rightarrow \text{可以并行} \end{cases}$

lack hidden to hidden recurrence: less powerful.

solu: ~~random~~ 随机生成 | 实际 value 当 input

2. 梯度消失与爆炸

2.1 线性

$$h^{(t)} = W h^{(t-1)} + V x^{(t)} + b$$

$$\Rightarrow \frac{\partial h^{(t)}}{\partial h^{(k)}} = \prod_{j=k+1}^t \frac{\partial h^{(j)}}{\partial h^{(j-1)}} = \prod_{j=k+1}^t W^T = (W^T)^{t-k}$$

$$\text{若 } W^T = V \text{diag}(\lambda) V^{-1} \text{ 则 } = V \text{diag}(\lambda)^{t-k} V^{-1}$$

若 $|\lambda_i| > 1$ 则 $\rightarrow \infty$ 梯度爆炸

若 $|\lambda_i| < 1$ $|\lambda_i|^{t-k} \rightarrow 0$ 消失

非线性 RNN

$$h^{(t)} = f(W h^{(t-1)} + V x^{(t)} + b)$$

$$\frac{\partial h^{(t)}}{\partial h^{(k)}} = \prod_{j=k+1}^t \frac{\partial h^{(j)}}{\partial h^{(j-1)}} = \prod_{j=k+1}^t W^T \text{diag}(f'(h^{(j-1)}))$$

范数满足

$$\left\| \frac{\partial h^{(t)}}{\partial h^{(k)}} \right\| \leq \|W^T\| \cdot \|\text{diag}(f'(h^{(j-1)}))\|$$

非线性激活作用 { ① bound recurrent dynamics 缓解 exploding gradient
② 若 f' 很小, 如 sigmoid / tanh 饱和 $\Rightarrow$ 更严重 gradient vanish

2.2 Solution

解决爆炸

huge difference

2.2.1 Gradient clipping

$$\hat{g} = \nabla \theta \varepsilon \quad \hat{g} \leftarrow \begin{cases} \hat{g}, & \|\hat{g}\| < \tau \\ \frac{\tau}{\|\hat{g}\|} \hat{g}, & \|\hat{g}\| \geq \tau \end{cases}$$

限 gradient 长度 防止更新过大

2.2.2 Identity Init + ReLU

huge difference.

初始化 W 为 $W \approx I$ + ReLU

3. LSTM

3.1 forget gate $f^{(t)} = \sigma(W^f h^{(t-1)} + U^f x^{(t)} + b^f)$ 对之前 $c^{(t-1)}$ 留多少

input gate $i^{(t)} = \sigma(W^i h^{(t-1)} + U^i x^{(t)} + b^i)$ 更新什么 value

$$\text{output gate } o^{(t)} = \sigma(W^o h^{(t-1)} + U^o x^{(t)} + b^o)$$

★ cell state $c^{(t)} = f^{(t)} \circ c^{(t-1)} + i^{(t)} \circ \tanh(W^c h^{(t-1)} + U^c x^{(t)} + b^c)$ candidate values
final hidden state $h^{(t)} = o^{(t)} \circ \tanh(c^{(t)})$ 让 gradient 能沿 cell state 加法路径传播

缓解长程依赖中的梯度消失

3.2 GRU Update gate $z^{(t)} = \sigma(W^z h^{(t-1)} + U^z x^{(t)} + b^z)$ 留多少旧 hidden state

Reset gate: $r^{(t)} = \sigma(W^r h^{(t-1)} + U^r x^{(t)} + b^r)$ 生成候选记忆是否忽略过去

New Memory Content: $\bar{h}^{(t)} = \tanh(W(r^{(t)} \circ h^{(t-1)}) + U x^{(t)} + b)$ 候选记忆

Final Memory: $h^{(t)} = z^{(t)} \circ h^{(t-1)} + (1 - z^{(t)}) \circ \bar{h}^{(t)}$

GRU 参数比 LSTM 少, 训练快, 结构简单. 但 LSTM 在大任务上性能更强且长程依赖 $\uparrow$

4. 双向 RNN

$$\vec{h}_t = f(\vec{W} x_t + \vec{V} \vec{h}_{t-1} + \vec{b}) \quad \overleftarrow{h}_t = f(\overleftarrow{W} x_t + \overleftarrow{V} \overleftarrow{h}_{t+1} + \overleftarrow{b})$$

$$\vec{y} = g(\vec{U} h_t + c) = g(\vec{U} [\vec{h}_t, \overleftarrow{h}_t] + c) \quad \text{分类, tagging, encoding}$$

Deep Bidirectional RNN: 将上者堆多层 $\vec{h}_t = f(\vec{W}_i \vec{h}_t^{i-1} + \vec{V}_i \overleftarrow{h}_{t+1} + b_i^{\vec{h}})$ $\overleftarrow{h}_t = f(\overleftarrow{W}_i \overleftarrow{h}_t^{i-1} + \overleftarrow{V}_i h_{t+1} + b_i^{\overleftarrow{h}})$
 $\vec{y} = g(\vec{U} [\vec{h}_t^L, \overleftarrow{h}_t^L] + c)$

5. Attention

Decoder 预测公式 $p(y_t | y_1, \dots, y_{t-1}, x) = g(y_{t-1}, s_t, c_t)$

$$s_t = f(s_{t-1}, y_{t-1}, c_t)$$

$$c_t = \sum_{j=1}^T \alpha_{t,j} h_j$$

上一个 context vector t 对 j 个输入位置 注意力权重

$$\alpha_{t,j} = \frac{\exp(e_{t,j})}{\sum_{k=1}^T \exp(e_{t,k})} \quad \text{其中 } e_{t,j} = a(s_{t-1}, h_j)$$

6. RNN 问题 { 难以并行
BP through sequence
局部信息与全局信息都通过 hidden state 这个 bottle-neck 传递

	RNN	self-Attention	convolution
单层	$O(nd^2)$	$O(n^2 d)$	
顺序操作	$O(n)$		
最大路长度	$O(n)$	$O(1)$	
总 work 仍是	$O(T)$	$O(1)$	

7. Mamba

$$h_t = A(x_t) h_{t-1} + B(x_t) x_t$$

$$y_t = C(x_t) h_t$$

重要 token 强烈更新 memory

不重要 token 几乎不改 memory.

inside a structured state-space evolution

Lec 6

1. Generalization gap $S = \{(x_i, y_i)\}_{i=1}^m$ $y_i \in \mathbb{R}$

空间 $\mathcal{H} = \{h(x) = w^T \phi(x) | w\}$

训练误差 $\hat{\epsilon}(h) = \frac{1}{m} \sum_{i=1}^m (y_i - h(x_i))^2$ 训练过程 $\hat{h} = \operatorname{argmin}_{h \in \mathcal{H}} \hat{\epsilon}(h)$

泛化误差 $\epsilon(h) = E_{(x,y) \sim D} [(y - h(x))^2]$

generalization gap $\epsilon(h) - \hat{\epsilon}(h)$

overfit: train error 低, generalization high

model 对 train set 随机波动敏感

2. Bias-Variance Decomposition

$$\epsilon = E_S [\epsilon(h_S)] = E_S [E_{(x,y) \sim D} [(y - h_S(x))^2]]$$

$$\text{代入 } y - h_S = y - E_S[h_S] + E_S[h_S] - h_S$$

$$\epsilon = E_S E_{(x,y)} [(y - E_S[h_S] + E_S[h_S] - h_S)^2]$$

后面这提出是 0

$$= E_S E_{(x,y)} [(y - E_S[h_S])^2] + E_S E_{(x,y)} [(E_S[h_S] - h_S)^2]$$

$$= E_{(x,y)} [(y - E_S[h_S(x)])^2] + E_S E_x [(E_S[h_S(x)] - h_S(x))^2]$$

$$\text{Expected Generation error} = \text{Bias}^2 + \text{Variance}$$

3. Ridge 与 LASSO.

$$\text{Ridge: } J(w, w_0) = \frac{1}{m} \sum_{i=1}^m [y_i - (w_0 + w^T \phi(x_i))]^2 + \lambda \|w\|_2^2$$

压小权重.

$$\text{Lasso: } J(w, w_0) + \frac{1}{m} \sum_{i=1}^m [y_i - (w_0 + w^T \phi(x_i))]^2 + \lambda \|w\|_1$$

$\phi(x_i)$ 特征映射

诱导稀疏性

$$\Rightarrow \text{统一形式 } \tilde{J}(\theta) = J(\theta) + \alpha \mathcal{L}(\theta) \quad \alpha \in [0, \infty) \text{ 正则强度}$$

4. L2 regularization / weight decay

$$\tilde{J}(w) = J(w) + \frac{\alpha}{2} \|w\|_2^2$$

$$\nabla_w \tilde{J}(w) = \nabla_w J(w) + \alpha w$$

$$w \leftarrow w - \eta (\nabla_w J(w) + \alpha w) = (1 - \eta \alpha) w - \eta \nabla_w J(w)$$

在正则最优点 w^* 附近做二阶近似

$$\tilde{J}(w) \approx J(w^*) + (w - w^*)^T \nabla_w J(w^*) + \frac{1}{2} (w - w^*)^T H (w - w^*)$$

因为 w^* 最优, $\nabla_w J(w^*) = 0$

$$= J(w^*) + \frac{1}{2} (w - w^*)^T H (w - w^*)$$

H 只与 w^* 有关, w^* 是确定. H const

$$\tilde{J}(w) \text{ 在梯度为0时等于最小值. 对 } w \text{ 求导 } \nabla_w \tilde{J}(w) = H(w - w^*) = 0$$

加入 L^2 regu 之后 有

$$\alpha \tilde{w} + H(\tilde{w} - w^*) = 0 \Leftrightarrow \tilde{w} = (H + \alpha I)^{-1} H w^*$$

$$(ABC)^T = C^T A^T B^T$$

当 $\alpha \rightarrow 0$, $\tilde{w} \rightarrow w^*$

$$\alpha \uparrow \text{ 分解 } H = Q \Lambda Q^T$$

Q orthogonal

$$Q^T Q = I \quad Q Q^T = I$$

$$\tilde{w} = (Q \Lambda Q^T + \alpha I)^{-1} Q \Lambda Q^T w^* = (Q (\Lambda + \alpha I) Q^T)^{-1} Q \Lambda Q^T w^* = Q (\Lambda + \alpha I)^{-1} \Lambda Q^T w^*$$

$\tilde{W} = Q(\Lambda + \alpha I)^{-1} \Lambda Q^T W^*$ 在第 i 个特征方向上, 权重被缩放为 $\frac{\lambda_i}{\lambda_i + \alpha}$.

L^2 对低曲率方向惩罚更强, 对高曲率方向影响较弱

动一点对 loss 影响小, 不重要 动一点对 loss 影响大, 重要

$$J(W; X, y) = (XW - y)^T (XW - y) \Rightarrow W^* = (X^T X)^{-1} X^T y.$$

加入 $L^2 \Rightarrow (X^T X + \alpha I)^{-1} X^T y$
防止 $X^T X$ 奇异

Bayesian: L^2 reg 等价于对权重施加 Gaussian

prior 后做 MAP estimation

即一开始假设

5. L_1 -norm

$\nabla_W \tilde{J}(W) = \nabla_W J(W) + \alpha \text{sign}(W)$ ∇ 不可导, 用 subgradient.

对 $W^* = 0$ 附近近似. 设 Hessian 是 diagonal. $H = \text{diag}(H_{1,1}, \dots, H_{n,n})$, $H_{i,i} > 0$.

$$\tilde{J}(W) = J(W^*) + \sum_i \frac{1}{2} H_{i,i} (W_i - W_i^*)^2 + \alpha |W_i|$$

each dim: $W_i = \text{sign}(W_i^*) \max \left\{ |W_i^*| - \frac{\alpha}{H_{i,i}}, 0 \right\}$. soft-thresholding.

$|W_i^*| \leq \frac{\alpha}{H_{i,i}} \Rightarrow W_i = 0$ 否则 $|W_i| = |W_i^*| - \frac{\alpha}{H_{i,i}}$ 向 0 移动.

$$\log p(W) = \sum_i \log \text{Laplace}(W_i; 0, \frac{1}{\alpha}) = -\alpha \|W\|_1 + n \log \alpha - n \log 2$$

6. Norm Penalty 作为 constrained optimization

$\min J(\theta)$, s.t. $\Omega(\theta) - k \leq 0$. 开拉 $L(\theta, \alpha) = J(\theta) + \alpha (\Omega(\theta) - k)$ $\alpha \geq 0$.

$$\theta^* = \arg \min_{\theta} \max_{\alpha, \alpha \geq 0} L(\theta, \alpha)$$

7. Data Augmentation { random crop crop $\rightarrow$ resize
Drop block 遮挡, 用灰色/随机亮灰色

more valid data $\Rightarrow$ less variance

8. Noise Injection 等价于对权重加惩罚

$$\tilde{J}_W = E_{p(x, y, \varepsilon_W)} [(\hat{y}_{\varepsilon_W}(x) - y)^2] = E_{p(x, y, \varepsilon_W)} [\hat{y}$$

$$\hat{y}(W + \varepsilon_W) \approx \hat{y}(W, x) + \varepsilon_W^T \nabla_W \hat{y}(W, x) + \frac{1}{2} \varepsilon_W^T \nabla_W^2 \hat{y}(W, x) \varepsilon_W$$

$$\tilde{J}_W = E_{x, y, \varepsilon_W} [(\hat{y}(x) + \varepsilon_W^T \nabla_W \hat{y}(x) - y)^2] = E_{x, y, \varepsilon_W} [(\hat{y}(x) - y)^2 + 2(\hat{y}(x) - y)(\varepsilon_W^T \nabla_W \hat{y}(x)) + (\varepsilon_W^T \nabla_W \hat{y}(x))^2]$$

$\hat{y}_{\varepsilon_W}(x) \approx \hat{y}(x) + \varepsilon_W^T \nabla_W \hat{y}(x)$
与 ε 无关 $\eta \|\nabla_W \hat{y}(x)\|^2$

Lec 7 Transformer

1. Pos emb.

$$W_i = \frac{1}{10000^{2i/d_m}}$$

$$PE_{pos, 2i} = \sin(W_i pos)$$

$$PE_{pos, 2i+1} = \cos(W_i pos)$$

$$\sin(a+b) = \sin a \cos b + \cos a \sin b$$

$$\cos(a+b) = \cos a \cos b - \sin a \sin b$$

$$\begin{bmatrix} \sin(W_i(pos+k)) \\ \cos(W_i(pos+k)) \end{bmatrix} = \begin{bmatrix} \cos(W_i k) & \sin(W_i k) \\ -\sin(W_i k) & \cos(W_i k) \end{bmatrix} \begin{bmatrix} \sin(W_i pos) \\ \cos(W_i pos) \end{bmatrix}$$

$$PE_{pos+k} = M_k PE_{pos}$$

M_k 只依赖偏移 k , 不依赖绝对位置. self-att 可通过线性变换学习相关位置信息

2. Decoder

2.1 Encoder-Decoder Attention

queries 来自 decoder, k, v 来自 encoder

$$Q_i = X_{\text{decoder}} W_i^Q$$

$$K_i = X_{\text{encoder}} W_i^K$$

$$V_i = X_{\text{encoder}} W_i^V$$

$$Z_i = \text{softmax}\left(\frac{Q_i K_i^T}{\sqrt{d_k}}\right) V_i$$

decoder 在当前位置提出 query, 去 encoder 检索相关信息

2.2 输出层 $P(y_t = v | y_{<t}, X) = \text{softmax}(W h_t + b)_v$

词表中某 token
在位置 t 的 final hidden state

$$\mathcal{L} = - \sum_t \log P(y_t^* | y_{<t}^*, X)$$

3. BERT

3.1 Masked Language Model

$$P(x_i | X_{\setminus i}) = \text{softmax}(W T_i + b) \quad \mathcal{L}_{MLM} = - \sum_{i \in M} \log P(x_i^* | X_{\text{corrupted}})$$

扰乱后全部 input
↓

3.2 Next Sentence Prediction

50% B 在 A 后 50% B 是随机句子.

使用 [CLS] 输出向量 C 预测: 不一定是天然是句子内容好 summary.

$$P(I_s \text{ Next} | A, B) = \text{softmax}(W C + b)$$

下游任务 finetune

$$\text{NSP Loss} \quad \mathcal{L}_{NSP} = -\log P(y_{NSP}^* | C)$$

$$\mathcal{L}_{BERT} = \mathcal{L}_{MLM} + \mathcal{L}_{NSP}$$

3.3 classification → 用 [CLS]

Token-level task

第 i token contextual
↑ representation

$$\hat{y} = \text{softmax}(W C + b)$$

$$\text{每个 token 的输出 } \hat{y}_i = \text{softmax}(W T_i + b)$$

4. GPT

$$N = \frac{HW}{P^2}$$

$$\text{长度为 } \frac{HWC}{HW/P^2} = P^2 C$$

5. ViT

$$\text{一个 patch } R^{1 \times P^2 C}$$

额外加入可学习 class token x_{class} 无 embedding

$Z_0 = [x_{\text{class}}; x_1^E; \dots; 1 + E_{\text{pos}}]$ E 是 embedding matrix x_{class} 不用乘, 但要加 pos emb

标准 encoder

$$y = \text{MLPHead}(Z_L^{\text{class token}}) \quad \text{class token 最终 representation}$$

6. CLIP: Contrastive Language-Image Pre-training.

图像、文本都 embed 进同一个 embedding space

让匹配的 text-image 对 vector 相似. 非 match 不相似.

6.1 Loss 对于一个 batch, 含 N pair $(I_1, T_1), \dots, (I_N, T_N)$

$$L_{\text{CLIP}} = -\frac{1}{2N} \sum_{i=1}^N \log \frac{\exp(I_i \cdot T_i)}{\sum_{j=1}^N \exp(I_i \cdot T_j)} - \frac{1}{2N} \sum_{j=1}^N \log \frac{\exp(I_j \cdot T_j)}{\sum_{i=1}^N \exp(I_i \cdot T_j)}$$

让 image 和 match text
匹配度在 texts 中最大
对 text 求和

让 text 和 match image
在 images 中匹配度最大,
对 image 求和

6.2 zero-shot classification: ① Create a text prompt for each class

一类. $\text{Prob } c = \frac{\exp(I \cdot T_c)}{\sum_{c'=1}^C \exp(I \cdot T_{c'})}$ only one 可以用 argmax 找最匹配的类.

DL Lec 8

1. The path to LLMs

1.1 Count-based Models: n-gram

$$P(x_{1:T}) = \prod_{t=1}^T P(x_t | x_{1:t-1}) \approx \prod_{t=1}^T P(x_t | x_{t-n+1}, \dots, x_{t-1})$$

Markov assumption: 当前词只依赖前面有限的 $n-1$ 个词

问题: ① 稀疏性 ② $n \uparrow$ 指数增长 ③ 很多上下文训练集无

X perplexity $PPL = \exp(-\frac{1}{T} \sum_{t=1}^T \log P(x_t | x_{<t}))$
或用以2为底的log $PPL = 2^{-\frac{1}{T} \sum_{t=1}^T \log_2 P(x_t | x_{<t})}$

PPL 越低, 给真实 token 概率越高

2. Scaling Law

2.1 Kaplan { Performance 依赖 scale, 对 model shape 弱相关
大模型 sample-efficient

超参影响小模型预测 Smooth power laws 小模型趋势外推到大模型

2.2 emergent abilities 小模无, 有模突然有

3. Encoder-decoder (T0/T5) Masked span prediction data, 训法, 架构, tokenizer 优化策略都影响

$$L_{span} = - \sum_{t=1}^{T_y} \log P_{\theta}(y_t | y_{<t}, \vec{x})$$

encoder 双向注意力, decoder causal 翻译, 摘要, 条件生成等

decoder-only 强于 encoder-decoder $\vec{x}$ 是被 mask span 替换序列, y 是要恢复的序列

4. Llama

LayerNorm $\rightarrow$ RMSNorm

$$\text{ReLU} \rightarrow \text{SwiGLU}(x) = \text{Swish}(xW) \times V = xW \text{Sigmoid}(AxW) \times V$$

Position: Absolute/Relative

MHA $\rightarrow$ GQA

5. CoT

$$P_{\theta}(a|x) = \sum_r P_{\theta}(a, r|x) = \sum_r P_{\theta}(r|x) P_{\theta}(a|x, r)$$

引入中间变量 r

先推理再回答. 在大模上收益明显. size $\uparrow$ benefit 越明显.

zero-shot CoT 有时甚至超过 few-shot CoT

$$W \in R^{d \times k}$$

6. LoRA $W' = W + \Delta W = W + BA$ $W \in R^{d \times r}$ $A \in R^{r \times k}$

7. Multimodal LLMs

Continuous repre: 把图/audio encode 成 dense embeddings 与文本 embedding 拼进 LLM

$$Z_{\text{image}} = f_{\text{vision}}(I) \quad Z = [Z_{\text{image}}; Z_{\text{text}}]$$

Discrete

信息丰富, 性能强 计算和内存看重
特征量化为 discrete tokens, VQ-VAE 或 clustering

$$z_t = \underset{j}{\arg \min} \|h_t - g_j\|^2 \quad \Rightarrow \text{多模状态输入} \quad \text{图像文本 token 一样被 LLM 处理.}$$

离散 unit id $\uparrow$ 连续特征 $\uparrow$ j codebook center

省存储, 序列长度 $\downarrow$ 去重. subword-like modeling

DL lec 9 VAE & Diffusion

1. Intro

生成模型 $\Rightarrow$ model distribution

- 处理缺失数据
- 测试对高维数据 handle
- 结合 RL

两类

- explicit generative models
GMM log likelihood. max
- implicit generative models
GAIN 没有 max likelihood

2. Autoencoder

2.1 def: 无监督. target 为输入本身. 一般用于降维
希望学 $x \in \mathbb{R}^n \rightarrow h \in \mathbb{R}^d \quad d \ll n$ h 是低维表示, 也叫 encoding

2.2 结构 encoder: $h = f_{\theta}(x) = a(Wx + b)$ $\theta \quad W, b$
decoder: $z = g_{\theta'}(h) = a(W'h + b')$ $\theta' \quad W', b'$

$$\theta^*, \theta'^* = \arg \min_{\theta, \theta'} \frac{1}{n} \sum_{i=1}^n L(x^{(i)}, z^{(i)}) = \arg \min_{\theta, \theta'} \frac{1}{n} \sum_{i=1}^n L(x^{(i)}, g_{\theta'}(f_{\theta}(x^{(i)})))$$

$L(x, z) = \|x - z\|^2$ 若 x 和 z 是 prob 的 vector, L 可以是 CE

$$L_H(x, z) = H(\beta x \| \beta z) = - \sum_{k=1}^d [x_k \log z_k + (1-x_k) \log (1-z_k)]$$

2.3 Under complete: $\dim(h) < \dim(x)$ 压缩信息

Over complete: $\dim(h) > \dim(x)$ 不一定有提取意义, 可能直接学 identity matrix

tied weight: decoder 权重 $W^* = W^T$ encoder 要加 regularization: sparsity, denoising, contractive penalty

3. Linear Autoencoder 与 PCA 关系

在线性 encoder, l_2 reconstruction loss 下, 最优 linear autoencoder 学到的子空间等价于 PCA 主成分子空间.

$$A = U \Sigma V^T \quad B^* = \arg \min_{B: \text{rank}(B)=k} \|A - B\|_F = U_{:, \leq k} \Sigma_{\leq k, \leq k} V_{:, \leq k}^T$$

前 k 个左奇异值向量, 若只能用 rank- r 近似数据, 最优是留前 k 的 principle directions

Autoencoder 平方误差满足 $\min_{\theta} \sum_t \frac{1}{2} (x_i^{(t)} - z_i^{(t)})^2 \geq \min_{W'} \frac{1}{2} \|X - W'h(x)\|_F^2$

最优解为 $W' \leftarrow U_{:, \leq k} \quad h(x) \leftarrow \sum_{\leq k, \leq k} V_{:, \leq k}^T$ 设 $X = U \Sigma V^T$

$$\Rightarrow h(x) = U_{:, \leq k}^T x \quad z = U_{:, \leq k} h(x)$$

这正是 PCA encoder

$$x^{(t)} \leftarrow \frac{1}{T} (x^{(t)} - \frac{1}{T} \sum_{i=1}^T x^{(i)})$$

singular values and vectors = eigenvalues / vectors of covariance matrix

$$X_{\text{centered}} = X - \mu x$$

$$\hat{X} = \frac{1}{T} X_{\text{centered}}$$

$$\hat{X} \hat{X}^T = \frac{1}{T} X_{\text{centered}} X_{\text{centered}}^T = C$$

$$\Sigma = \frac{1}{m-1} B^T B$$

协方差矩阵

4. AE $\rightarrow$ VAE

4.1 为什么 AE 不能直接生成

① 未显式规定 latent space 的概率分布 $\Rightarrow$ 破碎, 连续, 空散

② 训练时 decoder 只见过 encoder 产生的 latent codes, 没见过可任意采样的 latent vector

$$\text{VAE} \quad z \sim p(z) = \mathcal{N}(0, I)$$

4.2 生成假设

假设训练数据 $X = \{x^{(1)}, x^{(2)}, \dots, x^{(N)}\}$ 全由隐变量 z 生成

$$z \sim p(z), \quad x \sim p_\theta(x|z) \quad p(z) \text{ 常取 Gaussian prior } \mathcal{N}(0, I)$$

Decoder 建模 $p_\theta(x|z) = \mathcal{N}(\mu_x, \sigma_x^2)$ 由对 ± 1 值用 Bernoulli

希望最大化 $p_\theta(x) = \int p_\theta(z) p_\theta(x|z) dz$ 要对所有 z , intractable. 且 posterior $p_\theta(x|z)$ 也不可角

引入 encoder network

$$q_\phi(z|x) \text{ 来近似 } p_\theta(x|z) \quad \text{设 } q_\phi(z|x) = \mathcal{N}(\mu_z(x), \sigma_z(x))$$

4.3 ELBO

$$\begin{aligned} \log p_\theta(x) &= E_{z \sim q_\phi(z|x)} [\log p_\theta(x)] \quad \log p_\theta(x) \text{ 与 } z \text{ 无关} \\ &= E_{z \sim q_\phi(z|x)} [\log \frac{p_\theta(x, z)}{p_\theta(x|z)}] \quad \text{贝叶斯 } p_\theta(x) = \frac{p_\theta(x, z)}{p_\theta(x|z)} \\ &= E_{z \sim q_\phi(z|x)} [\log \frac{p_\theta(x, z)}{q_\phi(z|x)}] + E_{z \sim q_\phi(z|x)} [\log \frac{q_\phi(z|x)}{p_\theta(x|z)}] \quad \text{上下同乘 } q_\phi(z|x) \\ &= E_{z \sim q_\phi(z|x)} [\log \frac{p_\theta(x, z)}{q_\phi(z|x)}] + D_{KL}(q_\phi(z|x) \| p_\theta(x|z)) \quad \text{拆开} \\ &\geq \text{tractable } L(\theta, \phi, x) = E_{z \sim q_\phi(z|x)} [\log \frac{p_\theta(x, z)}{q_\phi(z|x)}] \quad \text{KL} \geq 0 \end{aligned}$$

$z \sim q_\phi(z|x)$ 其实因有随机节点无法 sample

但可采 $z \sim \mathcal{N}(0, I)$

$$z = \mu_\phi(x) + \epsilon \quad \epsilon \sim \mathcal{N}(0, I) \quad \Rightarrow \quad L(\theta, \phi, x) = E_{z \sim q_\phi(z|x)} [\log \frac{p_\theta(x, z)}{q_\phi(z|x)}] + E_{z \sim q_\phi(z|x)} [\log \frac{p_\theta(x|z)}{q_\phi(z|x)}]$$

encoder 输出的 $z|x$ 不要偏离 prior $p(z)$

reconstruction error - regularization term 让 decoder 从 z 重构 x

4.4 Gaussian KL 闭式解

当 prior 为标准高斯 $p(z) = \mathcal{N}(0, I)$

$$q_\phi(z|x) = \mathcal{N}(\mu_z(x; \phi), \sigma_z^2(x; \phi))$$

$$-D_{KL}(q_\phi(z|x) \| p(z)) = \frac{1}{2} (1 + \log \sigma_z^2 - \mu_z^2 - \sigma_z^2) \quad \text{多维时对 latent dimension } j=1, \dots, J \text{ 求和}$$

$$-D_{KL}(q_\phi(z|x) \| p(z)) = \frac{1}{2} \sum_{j=1}^J (1 + \log \sigma_{z,j}^2 - \mu_{z,j}^2 - \sigma_{z,j}^2) \quad \mu_{z,j} \text{ 是 } j\text{-th dim 均值 } \sigma_{z,j}^2 \text{ 是方差}$$

化简向令 $\mu_z \rightarrow 0 \quad \sigma_z^2 \rightarrow 1$

4.5 Representation Trick

$$E_{q_\phi(z|x)} [\log p_\theta(x|z)] \approx \frac{1}{K} \sum_k \log p_\theta(x|z^{(k)}), \quad z^{(k)} \sim q_\phi(z|x)$$

$z^{(k)}$ 是 random sample
BP 不能穿过 sampling node

$$\Rightarrow L(\theta, \phi, x) \approx \frac{1}{2} \sum_{j=1}^J (1 + \log \sigma_z^2 - \mu_{zj}^2 - \sigma_{zj}^2) + \frac{1}{K} \sum_k \log p_\theta(x|z^{(k)})$$

$z^{(k)} \sim N(0, I)$
 $z^{(k)} = \mu_z(x, \phi) + \sigma_z(x, \phi) \cdot \epsilon^{(k)}$
 z 对 μ_x, σ_z 是可微的

4.6 Generation: 采样 $z \sim N(0, I)$

Diagonal prior on $z \Rightarrow$ independent variables

4.7 VQ-VAE: 离散 latent representation

动机: token, visual code 等数据具有天然离散结构. VAE 用连续 Gaussian latent 不一定适合.

用 **discrete latent embeddings** - posterior 和 prior 都用 categorical distribution

采样结果是整数 index. 用 index 查 embedding table

$$q(z=k|x) = \begin{cases} 1, & k = \arg \min_j \|z_e(x) - e_j\|_2 \\ 0, & \text{otherwise} \end{cases}$$

$$L = \log p(x|z_q(x)) + \|sg[z_e(x)] - e\|_2^2 + \beta \|z_e(x) - sg[ze]\|_2^2$$

codebook embedding

让靠近 encoder 输出

encoder 输出会被量化到最近 codebook vector

commitment loss

让 encoder 输出承诺靠近某 embedding

$sg[\cdot]$ 是 stop-gradient operator
前传为 identity, 后传 $\nabla=0$

注意是一个格子选一个代码, 而不是整张图选一个

5. Diffusion

5.1 forward: $x_0 \sim q(x)$

$$x_t = \sqrt{1-\beta_t} x_{t-1} + \sqrt{\beta_t} \epsilon_t, \quad \epsilon_t \sim N(0, I)$$

$$q(x_t|x_{t-1}) = N(x_t; \sqrt{1-\beta_t} x_{t-1}, \beta_t I), \quad q(x_{1:T}|x_0) = \prod_{t=1}^T q(x_t|x_{t-1})$$

$$q(x_{0:T}) = q(x_0) q(x_{1:T}|x_0) = q(x_0) \prod_{t=1}^T q(x_t|x_{t-1})$$

$z \sim N(\mu_x + \mu_T, \sigma_x^2 + \sigma_T^2)$

$$\text{定义 } \alpha_t = 1 - \beta_t, \quad \bar{\alpha}_t = \prod_{s=1}^t \alpha_s$$

$$x_t = \sqrt{\alpha_t} x_{t-1} + \sqrt{\beta_t} \epsilon_t = \sqrt{\alpha_t} (\sqrt{\alpha_{t-1}} x_{t-2} + \sqrt{1-\alpha_{t-1}} \epsilon_{t-1}) + \sqrt{1-\alpha_t} \epsilon_t$$

$$\Rightarrow x_t = \sqrt{\bar{\alpha}_t} x_0 + \sqrt{1-\bar{\alpha}_t} \epsilon, \quad \epsilon \sim N(0, I)$$

$$t \uparrow \Rightarrow \bar{\alpha}_t \rightarrow 0, \quad q(x_T|x_0) \approx N(0, I)$$

5.2 Noise Scheduler

linear: $\beta_t = \beta_1 + \frac{t-1}{T-1} (\beta_T - \beta_1)$

cosine: $\bar{\alpha}_t = \frac{f(t)}{f(0)}$
 $f(t) = \cos^2\left(\frac{t}{T} \cdot \frac{\pi}{2}\right)$

quadratic: $\beta_t = \left[\sqrt{\beta_1} + \frac{t-1}{T-1} (\sqrt{\beta_T} - \sqrt{\beta_1}) \right]^2$

sigmoid: $\beta_t = \beta_1 + \frac{\beta_T - \beta_1}{1 + e^{\tau(1 - \frac{t-1}{T-1})}}$

5.3 反向 $X_T \sim N(0, I)$ $X_{t+1} = \mu_\theta(X_t, t) + \sigma_t Z$, $Z \sim N(0, I)$ $p(X_T) = N(X_T; 0, I)$

$$p_\theta(X_{t+1}|X_t) = N(X_{t+1}; \mu_\theta(X_t, t), \sigma_t^2 I) \quad p_\theta(X_{0:T}) = p(X_T) \prod_{t=1}^T p_\theta(X_t|X_t)$$

5.4 Noise Predictor: 不直接预测 μ_θ , 而是训练 network 预测加到 x_0 的 compound noise

$$\varepsilon_\theta(X_t, t) \approx \tilde{\varepsilon}_t$$

noisy image

U-Net

$$L = E_{x_0 \sim q(x_0), t \sim U(1, T), \tilde{\varepsilon}_t \sim N(0, I)} [\|\tilde{\varepsilon}_t - \varepsilon_\theta(\sqrt{\alpha_t}x_0 + \sqrt{1-\alpha_t}\tilde{\varepsilon}_t, t)\|^2]$$

随机取 - 真实图. 时间步. 加噪. 预测

5.5 为什么 noise prediction 等价 reverse mean

$$q(X_{t+1}|X_t, x_0) = N(X_{t+1}; \tilde{\mu}_t(X_t, x_0), \tilde{\beta}_t I) \quad \tilde{\mu}_t(X_t, x_0) = \frac{1}{\sqrt{\alpha_t}} (X_t - \frac{\beta_t}{1-\alpha_t} \varepsilon)$$

$$\mu_\theta(X_t, t) = \frac{1}{\sqrt{\alpha_t}} (X_t - \frac{1-\alpha_t}{\sqrt{1-\alpha_t}} \varepsilon_\theta(X_t, t))$$

$$D_{KL} \propto \|\tilde{\mu}_t - \mu_\theta\|^2 \Rightarrow E_{x_0, \varepsilon, t} [\|\varepsilon - \varepsilon_\theta(X_t, t)\|^2]$$

5.6 ELBO

$$E_{x_0 \sim q(x_0)} [-\log p_\theta(x_0)] \leq E_{x_0, x_{1:T} \sim q(x_{0:T})} [-\log \frac{p_\theta(x_0, x_{1:T})}{q(x_{1:T}|x_0)}]$$

$$\begin{aligned} \mathbb{P}_\theta &= [-\log p(x_T) - \sum_{t=1}^T \log p_\theta(X_{t+1}|X_t) + \sum_{t=1}^T \log q(X_t|X_{t-1})] \\ &= \end{aligned}$$

telescoping

$$= \frac{q(X_t|x_0) q(X_{t+1}|X_t, x_0)}{q(X_{t+1}|x_0)} \quad \log \text{完} = \sum_{t=1}^T \log q(X_t|x_0)$$

$$\begin{aligned} &= D_{KL}(q(X_T|x_0) \| p(X_T)) + \sum_{t=2}^T E_q [D_{KL}(q(X_{t+1}|X_t, x_0) \| p_\theta(X_{t+1}|X_t))] + \sum_{t=1}^T \log q(X_{t+1}|X_t) \\ &\quad \text{reverse transition matching} \\ &\quad \text{让 final noisy distribution} \quad \text{匹西 Gaussian prior} \quad \text{reconstruction term} = \log q(X_T|x_0) \\ &\quad \text{telescoping} \end{aligned}$$

5.7 Sample

$$\textcircled{1} X_T \sim N(0, I)$$

② 对 $t = T, T-1, \dots, 1$ denosing direction

$$X_{t+1} = \mu_\theta(X_t, t) + \sigma_t Z \quad \leftarrow \text{stochastic sampling term} \quad Z \sim N(0, I)$$

$$\text{其中 } \mu_\theta(X_t, t) = \frac{1}{\sqrt{\alpha_t}} (X_t - \frac{1-\alpha_t}{\sqrt{1-\alpha_t}} \varepsilon_\theta(X_t, t))$$

$$\text{理论 } \sigma_t^2 = \tilde{\beta}_t = \frac{1-\alpha_{t+1}}{1-\alpha_t} \beta_t$$

但实践 $\sigma_t^2 = \beta_t$

若 $\beta_t \neq 0$ 则 stochastic sampling
若 $\beta_t = 0$, 则 DDIM

deterministic

DL Lec 10 Diffusion models

1. ODE SDE

ODE

$$\frac{dx}{dt} = f(x, t) \text{ or } dx = f(x, t) dt$$

Analytical solu $x(t) = x(0) + \int_0^t f(x, \tau) d\tau$

Iterative Num solu $x(t+\Delta t) \approx x(t) + f(x(t), t) \Delta t$

确定性轨迹, 给定 $x(0)$ 后未来状态完全确定
无 randomness

SDE

$$\frac{dx}{dt} = \underbrace{f(x, t)}_{\text{drift coefficient}} + \underbrace{g(x, t)}_{\text{diffusion coefficient}} W_t$$

$$dx = f(x, t) dt + g(x, t) dW_t \text{ random}$$

$$x(t+\Delta t) \approx x(t) + f(x(t), t) \Delta t + g(x(t), t) \sqrt{\Delta t} \epsilon$$

$$\epsilon \sim N(0, I)$$

每次运行都不同, 引入概率云

2. 从离散 diffusion $\rightarrow$ 连续时间 SDE

2.1 forward diffusion 连续极限

无限多小步的极限

Taylor

$$x_t = \sqrt{1 - \beta_t} x_{t-1} + \sqrt{\beta_t} N(0, I) \quad \text{令 } \beta_t := \beta(t) \Delta t \quad \text{当 } \Delta t \text{ 很小时: } \sqrt{1 - \beta(t) \Delta t} \approx 1 - \frac{1}{2} \beta(t) \Delta t$$

$$x_t \approx x_{t-1} - \frac{\beta(t) \Delta t}{2} x_{t-1} + \sqrt{\beta_t} \Delta t N(0, I)$$

令 $\Delta t \rightarrow 0$ 得到前向 SDE

$$dx_t = \underbrace{-\frac{1}{2} \beta(t) x_t dt}_{\text{drift term}} + \underbrace{\sqrt{\beta(t)} dW_t}_{\text{diffusion term}}$$

Wiener process 布朗

$\rightarrow$ 连续时间噪声日常

把 $x_t \rightarrow 0$ 拉, 一方面加入新噪声

2.2 Reverse-Time SDE and score function

SDE 前向: $dx_t = -\frac{1}{2} \beta(t) x_t dt + \sqrt{\beta_t} dW_t$

反向: $dx_t = \underbrace{[-\frac{1}{2} \beta(t) x_t - \beta(t) \nabla_{x_t} \log q_t(x_t)]}_{\text{score}} dt + \sqrt{\beta(t)} d\bar{W}_t$

时间边缘分布

反向

$\rightarrow$ 布朗

$\nabla_{x_t} \log q_t(x_t)$ 表示当前 x_t 应往哪里走, 才能更接近高频数据区

Score function $\log$ -density 对 input 的 ∇ $s(x) = \nabla_x \log q(x)$

若某点在低频区, score 指引移动至高频区域方向

在 diffusion 中, 反向生成的本质是学每个噪声等级下的 score

$$s_\theta(x_t, t) \approx \nabla_{x_t} \log q_t(x_t)$$

反向 SDE 推导: 前向 $dx_t = f(x_t, t) dt + g(t) dW_t$

小时间步转移

$$x_{t+\Delta t} = x_t + f(x_t, t) \Delta t + g(t) \sqrt{\Delta t} \epsilon, \quad \epsilon \sim N(0, I)$$

$$p(x' | x) = N(x'; x + f(x, t) \Delta t, g(t)^2 \Delta t I)$$

- 泰勒展开

$$\log q_t(x) \approx \log q_t(x') + (x - x')^T \nabla \log q_t(x')$$

Gaussian Completion

$$E[x_t | x_{t+\Delta t}] = x_{t+\Delta t} - f(x_{t+\Delta t}, t) \Delta t + g(t)^2 \nabla \log q_t(x_{t+\Delta t}) \Delta t$$

$$\Rightarrow dx_t = [f(x_t, t) - g(t)^2 \nabla_{x_t} \log q_t(x_t)] dt + g(t) d\bar{W}_t, \text{ 代入 VP-SDE}$$

$$g(t) = \sqrt{\beta(t)} \quad f(x_t, t) = -\frac{1}{2} \beta(t) x_t$$

$$dx_t = [-\frac{1}{2} \beta(t) x_t - \beta(t) \nabla_{x_t} \log q_t(x_t)] dt + \sqrt{\beta(t)} d\bar{W}_t$$

3. Score Matching & Denoising Score Matching

希望 $\min_{\theta} E_{t \sim U(0,T)} E_{x_t \sim q_t(x_t)} [\|S_{\theta}(x_t, t) - \nabla_{x_t} \log q_t(x_t)\|_2^2]$ 直接训, 但 $\nabla_{x_t} \log q_t(x_t)$ 不可算。
 条件得分在后验分布下期望 = 边缘得分
 但有 $\nabla_{x_t} \log q_t(x_t) = E_{x_0 \sim q(x_0|x_t)} [\nabla_{x_t} \log q_t(x_t|x_0)]$ 所有数据点扩散后边缘混合分布。
 且 $q_0(x_0)$ 未知

Proof: $\log q(x_0|x_t) = \log q_0(x_0) + \log q_t(x_t|x_0) - \log q_t(x_t)$ 贝叶斯. 对 x_t 求 ∇

$$\nabla_{x_t} \log q(x_0|x_t) = \nabla_{x_t} \log q_0(x_0) + \nabla_{x_t} \log q_t(x_t|x_0) - \nabla_{x_t} \log q_t(x_t)$$

$$\text{移项 } \nabla_{x_t} \log q_t(x_t|x_0) = \nabla_{x_t} \log q(x_0|x_t) + \nabla_{x_t} \log q_t(x_t) = \int q(x_0|x_t) \nabla_{x_t} \log q(x_0|x_t) dx_0$$

固定 x_t 取条件期望 $E_{x_0 \sim q(x_0|x_t)} [\nabla_{x_t} \log q(x_0|x_t)] = \nabla_{x_t} \log q_t(x_t)$ 在 x_t 固定时与 x_0 无关

$$\text{用 } \nabla_{x_t} \log q(x_0|x_t) = \frac{\nabla_{x_t} q(x_0|x_t)}{q(x_0|x_t)} \text{ log trick } \Rightarrow \int \nabla_{x_t} q(x_0|x_t) dx_0 = \nabla_{x_t} \int q(x_0|x_t) dx_0 = \nabla_{x_t} 1 = 0$$

可得 $\min_{\theta} E_{t \sim U(0,T)} E_{x_0 \sim q_0(x_0)} E_{x_t \sim q_t(x_t|x_0)} [\|S_{\theta}(x_t, t) - \nabla_{x_t} \log q_t(x_t|x_0)\|_2^2]$

PPT: 在期望后 $S_{\theta}(x_t, t) \approx \nabla_{x_t} \log q_t(x_t)$ 其实应为 $\nabla_{x_t} \log q_t(x_t) = E_{x_0 \sim q(x_0|x_t)} [\nabla_{x_t} \log q_t(x_t|x_0)]$

$$\text{代入 } x_t = \gamma x_0 + \beta t \epsilon \quad \epsilon \sim N(0, I) \quad \nabla_{x_t} \log q_t(x_t|x_0) = -\nabla_{x_t} \frac{(x_t - \gamma x_0)^2}{2\beta t^2} = -\frac{\epsilon}{\beta t} \Rightarrow S_{\theta}(x_t, t) = -\frac{S_{\theta}(x_t)}{\beta t}$$

$$\Rightarrow \min_{\theta} E_{t, x_0, \epsilon} \left[\frac{1}{6t^2} \|\epsilon - S_{\theta}(x_t, t)\|_2^2 \right] \text{ 说明DDPM中预测噪声本质上在学 score}$$

4. Probability Flow ODE

反向生成 ODE 在分布意义上等价于一个石角定性 ODE: $dx_t = -\frac{1}{2} \beta(t) [x_t + \nabla_{x_t} \log q_t(x_t)] dt$ probability flow ODE
 随 二者在每个时间 t 的边缘分布 $q_t(x_t)$ 相同, 可用 ODE solver 采样
 当用 $q_T(x_T) \approx N(x_T; 0, I)$ 初始化

5. Flow Matching

Flow Model $\frac{dx}{dt} = u_t^{\theta}(x_t)$ 不随机, 学 vector field, 解 ODE

Diffusion $dx_t = u_t^{\theta}(x_t) + \beta_t dw_t$ 随机, 学 drift/score/noise 解 SDE/ODE

Flow Model 生成: ① $x_0 \sim p_{init}$, $dx_t = u_t^{\theta}(x_t) dt$ ② 从 $t=0$ 模拟到 $t=1$, 得到 $x_t \sim p_{data}$

5.1 Time dependent vector field

瞬时速度 time = t loc = x

$u: \mathbb{R}^d \times [0, 1] \rightarrow \mathbb{R}^d \quad (x, t) \mapsto u_t(x)$ 给每点一个运动方向和速度. 沿速度场就能把初始分布运过去

$$\text{轨迹 } \dot{x}_t = u_t(x_t) \Leftrightarrow \frac{dx_t}{dt} = u_t(x_t)$$

5.2 Conditional Prob Path

给定目标数据点 $z \in \mathbb{R}^d$, 定义 cond prob path $p_t(\cdot|z)$ 满足 从 noise 收缩到目标 z 的路径.
 如 Gaussian $p_t(\cdot|z) = N(x_t; z, \beta_t^2 Id)$ 满足 控制从 0 到 z 控制从 1 到 z

$p_0(\cdot|z) = p_{init}$, $p_1(\cdot|z) = \delta_z$ 集中在 z 的 Dirac delta

$$\text{边界 } x_0=0, \beta_0=1 \Rightarrow p_0(\cdot|z) = N(0, Id) = p_{init} \quad p_1(\cdot|z) = N(z|0) = \delta_z$$

5.3 Cond Vector Field 对任意 cond prob path 都存在等价 vector field, 使得

$x_0 \sim p_{init}$, $\frac{d}{dt} x_t = u_t(x_t|z) \Rightarrow x_t \sim p_t(\cdot|z)$ 对高斯路径 $p_t(\cdot|z) = N(x_t; z, \beta_t^2 Id)$

$$\text{Cond vector field 为 } u_t(x_t|z) = \left(\dot{z} - \frac{\beta_t}{\beta_t^2} \dot{\beta}_t z + \frac{\beta_t}{\beta_t^2} x_t \right) \quad \dot{z} = \frac{dz}{dt} \quad \beta_t = \frac{d\beta_t}{dt}$$

给定 $u_t(x_t|z)$ 在目标 z 处, 把 x 推到 z 的速度

推导. 若有 $x_t = \alpha_t z + \beta_t \epsilon$, ϵ fix, 则 $\frac{dx_t}{dt} = \dot{\alpha}_t z + \beta_t \dot{\epsilon}$, 又因为 $\epsilon = \frac{x_t - \alpha_t z}{\beta_t}$ 整理 $u_t(x_t|z) = \frac{dx_t}{dt} - \left(\dot{\alpha}_t - \frac{\dot{\beta}_t}{\beta_t} \alpha_t \right) z + \frac{\beta_t}{\beta_t^2} x_t$

5.4 Marginal Prob Path and Vector field 生成整个数据分布

令 $z \sim p_{data}$ 定义 marginal prob path $p_t(x) = \int p_t(x|z) p_{data}(z) dz$

对应 marginal vector field $u_t(x) = \int u_t(x|z) \frac{p_t(x|z) p_{data}(z)}{p_t(x)} dz = E_{z \sim p(z|x_t=x)} [u_t(x|z)]$ 在位置 x 上, 有多目标数据点 z 能解释中间态. marginal 是后验平均. 但 p_{data} 未知, $u_t(x)$ 无法显式写出 用 $\hat{u}_t(x)$ 逼近

5.5 FM Loss $L_{CFM} = E_{t \sim U(0,1), z \sim p_{data}, x \sim p(\cdot|z)} [\|\hat{u}_t(x) - u_t(x|z)\|^2]$ 最小化 L_{CFM} 等价最小化 LFM. 训练用条件, 但期望给的 $u_t(x)$

loss推导: 有 $u_t(x|z) = (\alpha_t - \frac{\beta_t}{\alpha_t} \alpha_t)z + \frac{\beta_t}{\alpha_t} (\alpha_t z + \beta_t \varepsilon) = \cancel{\alpha_t z} - \cancel{\frac{\beta_t}{\alpha_t} \alpha_t z} + \beta_t \varepsilon$ 代入 $x = \alpha_t z + \beta_t \varepsilon$
 $L_{CFM} = E_{t,z,\varepsilon} [\| \tilde{u}_t(\alpha_t z + \beta_t \varepsilon) - (\alpha_t z + \beta_t \varepsilon) \|^2]$

5.6 最简单 optimal transport path 令 $\alpha_t = t$ $\beta_t = 1-t$ $\dot{\alpha}_t = 1$ $\dot{\beta}_t = -1$ 采样公式 $x = tz + (1-t)\varepsilon$ $\varepsilon \sim N(0, Id)$
 $L_{CFM} = E_{t \sim U(0,1), z \sim p_{data}, \varepsilon \sim N(0, Id)} [\| \tilde{u}_t(tz + (1-t)\varepsilon) - (tz - \varepsilon) \|^2]$ 预测 x 真 x $u_t(\hat{x}) = \dot{\alpha}_t z + \dot{\beta}_t \varepsilon = z - \varepsilon$
 ① $z \sim p_{data}$ ② $t \sim U(0,1)$ ③ $\varepsilon \sim N(0, I)$ ④ 构造中间点 $x = tz + (1-t)\varepsilon$ ⑤ 用 network 预测 $\tilde{u}_t(x)$ ⑥ MSE 监督

6. 加速采样 DDIM: 少步采样. 允许石角定性采样

传统: $x_{t-1} = \sqrt{\alpha_{t-1}} x_0 + \sqrt{1-\alpha_{t-1}} \varepsilon_{t-1}$
 DDIM: $x_{t-1} = \sqrt{\alpha_{t-1}} \hat{x}_0 + \sqrt{1-\alpha_{t-1} + \beta_t^2} \varepsilon_t + \beta_t z$ 从 $x_t = \sqrt{\alpha_t} x_0 + \sqrt{1-\alpha_t} \varepsilon_t$ 推出
 其中 $\hat{x}_0 = \frac{x_t - \sqrt{1-\alpha_t} \varepsilon_t}{\sqrt{\alpha_t}}$ 其中 $\hat{x}_0 = \frac{x_t - \sqrt{1-\alpha_t} \varepsilon_t}{\sqrt{\alpha_t}}$
 若 $\beta_t = 0$, 则 DDIM 石角定. 可只在子序列时间点采样 走 K 步 $\varepsilon_t = \frac{x_t - \sqrt{\alpha_t} \hat{x}_0}{\sqrt{1-\alpha_t}}$
 $t_1 = 1 < t_2 < \dots < t_K \leq T$ 用来估 ε_{t-1}

7. DPM Solver

$$\left(\frac{f \circ \log x}{\log f} \right)' = \frac{f'}{f}$$

8. Stable Diffusion

9. Classifier Guidance 与 Classifier-free guidance

9.1 Vanilla: 给定条件 y . 直接训 $u_t^\theta(x|y)$ ① 选 prompt y ② $x_0 \sim p_{init}$ ③ 模拟 $dx_t = u_t^\theta(x_t|y)dt$ t 从 0 到 1
 prob: 不够多样. image 和 prompt 未对齐 不能正常预测 y 分类器

9.2 classifier guidance: $\nabla_x \log p_t(x|y) = \nabla_x \log p_t(x) + \nabla_x \log p_t(y|x)$ 分解为 C
 $u_t^{target}(x|y) = u_t^{target}(x) + \alpha \nabla_x \log p_t(y|x)$ uncondition dir + prompt dependent
 进一步增强 $\tilde{u}_t^w(x|y) = u_t^{target}(x) + w \alpha \nabla_x \log p_t(y|x)$ w 是 guidance scale. $w \uparrow \rightarrow$ 约束 $\uparrow \rightarrow$ 饱和, 失真, 多样性 $\downarrow$

9.3 CFG: 无需外部 classifier $u_t^{target}(x|y) - u_t^{target}(x)$ 可近似看作 prompt-dependent direction
 于是放大差值 $\tilde{u}_t^w(x|y) = u_t^{target}(x) + w (u_t^{target}(x|y) - u_t^{target}(x)) = (1-w) u_t^{target}(x) + w u_t^{target}(x|y)$
 CFG 采样向量差项为 $u_t^{\theta,w}(x) = (1-w) u_t^\theta(x|\phi) + w u_t^\theta(x|y)$ 空文本
 非严格数据分布建模 模拟无条件
 $w=1$ 近似正常条件生成
 $w>1$ $u_t^{\theta,w}(x) = u_t^\theta(x|y) + (w-1) [u_t^\theta(x|y) - u_t^\theta(x|\phi)]$ 沿 prompt-dependent direction 走得比真实分布更远