

1. MLE 设 θ 固定, 完全依赖 data 来找参数 $m_h + m_t = m$
 $\hat{\theta}_{MLE} = \frac{m_h}{m_t + m_h}$ $L(\theta|D) = \theta^{m_h} (1-\theta)^{m_t}$ 整合所有信息
但 MLE 无法提供置信度 sufficient statistic e.g. (m_h, m_t)

2. Bayesian: 小数据时先验重要, 把 θ 当 RV, 在有 prior $p(\theta)$ 下求 $P(\theta|D)$
 $P(D_{m+1} = H|D) = \int P(D_{m+1} = H, \theta, D) d\theta = \int P(D_{m+1} = H|\theta, D) p(\theta|D) d\theta$
Bayesian MAP 简化成求最大后验 $\theta^* = \arg\max_{\theta} P(\theta|D) \propto \theta^{m_h} (1-\theta)^{m_t} p(\theta)$
Beta 分布做先验 $p(\theta) = \frac{\Gamma(d_t + d_h)}{\Gamma(d_t)\Gamma(d_h)} \theta^{d_t-1} (1-\theta)^{d_h-1}$
 d_h 和 d_t 是虚构, 认为在 $\alpha = d_t + d_h$ 次里的结果 $E[\theta] = \frac{m_h + d_h}{m + \alpha}$
 $\Rightarrow \theta^{m_h + d_h - 1} (1-\theta)^{m_t + d_t - 1}$ $P(D_{m+1} = H|D) = \int \theta p(\theta|D) d\theta = \frac{m_h + d_h}{m + \alpha}$
 $\alpha \uparrow \Rightarrow$ prior 越自信

3. Jensen's Inequality $\square: \lambda f(x) + (1-\lambda)f(y) \leq f(\lambda x + (1-\lambda)y)$
对 prob $\sum_{i=1}^n p_i f(x_i) \leq f(\sum_{i=1}^n p_i x_i) \Rightarrow E[\log(x)] \leq \log(E[x])$

4. Entropy $H(x) = -\sum_{x \in X} p(x) \log p(x) = -E[\log(p(x))]$ $0 \log 0 = 0$
① $H(x) \geq 0$ ② $H(x) = 0$ iff $\exists x \in X, P(x) = 1$ ③ $H(x) \leq \log |X|$
proof 3 $H(x) = \frac{1}{n} \sum_{i=1}^n \varphi(p_i) \geq \varphi(\frac{1}{n} \sum_{i=1}^n p_i) = \varphi(\frac{1}{n}) = \log \frac{1}{n}$
 $H(p) \leq \log n$ 用了 $\varphi(x) = x \log x (x > 0)$ $\varphi''(x) > 0$ 严格凸 $f(E[X]) \leq E[f(X)]$

5. 条件 entropy $H(Y|X=x) = \sum_y P(y|X=x) \log \frac{1}{P(y|X=x)}$ 把 $P(x)$ 换成 $P(Y|X=x)$
 $H(Y|X=x)$ 可能大于 $H(Y)$, 因为 $X=x$ 可能带来不确定性

平均条件熵 $H(Y|X) = \sum_x p(x) H(Y|X=x)$ 已知 X 后, Y 的剩余平均不确定性

联合熵 $H(X, Y) = \sum_{x,y} p(x,y) \log \frac{1}{p(x,y)}$

6. Divergence $D_{KL}(P||Q) = \sum_x p(x) \log \frac{p(x)}{q(x)}$
① ≥ 0 . $\because \sum_x p(x) \log \frac{p(x)}{q(x)} \geq -\log(\sum_x p(x) \frac{q(x)}{p(x)}) = -\log 1 = 0$
② $p(x) > 0, q(x) = 0$, $KL \rightarrow \infty$

Cross entropy $H(P, Q) = -\sum_x p(x) \log q(x) = -E[\log q(x)]$
 $KL(P||Q) = H(P, Q) - H(P) \geq 0 \Leftrightarrow \sum_x p(x) \log q(x) \leq \sum_x p(x) \log p(x)$
令 $f(x)$ 为非负, 有 $\sum_x f(x) \log q(x) \leq \sum_x f(x) \log p(x)$ $P^*(x) = f(x) / \sum_x f(x)$
 $D_{JS}(P||Q) = \frac{1}{2} KL(P||\frac{P+Q}{2}) + \frac{1}{2} KL(Q||\frac{P+Q}{2})$ 有界 $[0, \log 2]$
 $\log 2$ 为 P, Q 完全不相交, $\frac{P+Q}{2}(x)$ 退化为 $\frac{1}{2} p(x)$ 或 $\frac{1}{2} q(x)$ $\sum_x p(x) \log 2 = \log 2$
最小 $KL(P||Q) = H(P||Q) - H(P)$ 固定 $\approx -\frac{1}{n} \sum \log Q(x_i) = -\frac{1}{n} \log Q(D)$
MC 采样. $P(x, y) \xrightarrow{\text{采}} D \xrightarrow{\text{采}} Q(y|x)$ 最小 $KL \Leftrightarrow \max_{\frac{1}{n}} \log Q(y_i|x_i)$

7. Mutual info $I(X; Y) = H(X) - H(X|Y) = \sum_{x,y} p(x,y) \log \frac{1}{p(x)} - \sum_{x,y} p(x,y) \log \frac{1}{p(x,y)}$
 $= \sum_{x,y} p(x,y) \frac{p(x,y)}{p(x)} = \sum_{x,y} p(x,y) \frac{p(x,y)}{p(x)p(y)} = KL(P(X, Y)||P(X)P(Y))$
 $= H(Y) - H(Y|X)$ 对称 $\Rightarrow I(X, Y) \geq 0$, $= 0$ iff $X \perp Y$ 因为 KL

8. c target concept; C concept Class. S 训练. $H \neq H_y$
泛化 $R(h) = \Pr[h(x) \neq c(x)] = E[\mathbb{1}_{h(x) \neq c(x)}]$ 经验 $\hat{R}_n(h) = \frac{1}{n} \sum \mathbb{1}_{h(x_i) \neq c(x_i)}$

9. Prob Appl Curr Lea (FPM) 对 $\Pr[R(h) > \epsilon] \leq \delta$
 $m \geq \text{poly}(1/\epsilon, 1/\delta, n, \text{size}(c))$: $\Pr[R(h) \leq \epsilon] \geq 1 - \delta$ 存在, 则 PLA-learn
 $\delta \sim D^m$ 有 overlap, 放大 - 个样本没进去 $\delta \sim 1 - \frac{1}{4}$
希望 $\Pr(r_i) \leq \frac{1}{4}$, 这样放完 $\Pr[U_i r_i] \leq \frac{1}{4}$ $\downarrow$ 4个 strip m 个样本
bad case 有些超 $\frac{1}{4}$ $\Pr[U_i r_i] > \frac{1}{4}$ $\leq 4(1 - \frac{1}{4})^m \leq 4e^{-m/4}$
9.2 $H \neq c$, $R(hs) = 0$: consistent, 在训练上能做到 $\hat{R}_n(h) = 0$
Pr sum $[R(h) \leq \epsilon] \geq 1 - \delta$ holds if $m \geq \frac{1}{\epsilon} \log(1/\delta) + \log \frac{1}{\epsilon}$, $R(h) \leq \frac{1}{m} \log(1/\delta) + \log \frac{1}{\epsilon}$
Proof: $\Pr[\exists h \in H: \hat{R}_n(h) = 0 \wedge R(h) > \epsilon] \leq \sum_{h \in H} \Pr[\hat{R}_n(h) = 0 \wedge R(h) > \epsilon] \leq |H| (1 - \epsilon)^m \leq |H| e^{-m\epsilon}$
 $\Leftrightarrow m \geq \frac{1}{\epsilon} \log(1/\delta) \Leftrightarrow R(h) \leq \frac{1}{m} \log(1/\delta) + \log \frac{1}{\epsilon}$ 一个放成所有
 $= \sum \Pr[\hat{R}_n(h) = 0 | R(h) > \epsilon] \Pr[R(h) > \epsilon] \leq |H| (1 - \epsilon)^m \leq |H| e^{-m\epsilon} = \delta$
 $\Rightarrow m \geq \frac{1}{\epsilon} \log(1/\delta) \Leftrightarrow R(h) \leq \frac{1}{m} \log(1/\delta) + \log \frac{1}{\epsilon}$ 0(1/m)
只要 hypothesis set 成立有限, ERM 就是 PAC

9.3 Inconsistent 常见 $H \neq c$ 且 $R(hs) \neq 0$ Hoeff's ineq X_i 独立, $E[X_i] = b_i$
 $\Rightarrow \Pr(\sum X_i - E[\sum X_i] \geq \epsilon) \leq e^{-\frac{\epsilon^2}{2 \sum (b_i - a_i)^2}}$ 和 $\Pr(\sum X_i - E[\sum X_i] \leq -\epsilon) \leq e^{-\frac{\epsilon^2}{2 \sum (b_i - a_i)^2}}$
在 Hypo = $\{0, 1\}$ 下, 代入 $E[\sum X_i] = R(h)$, $\sum X_i$ 是 $R(h)$ 代入 $a=0, b=1$
 $\Pr[|\hat{R}_n(h) - R(h)| \geq \epsilon] \leq 2e^{-\frac{\epsilon^2}{2m}} \Leftrightarrow \frac{\epsilon^2}{2m} \geq \log \frac{2}{\delta} \Leftrightarrow |\hat{R}_n(h) - R(h)| \leq \sqrt{\frac{2 \log \frac{2}{\delta}}{m}}$
 $\forall h \in H, |\hat{R}_n(h) - R(h)| \geq \epsilon \Rightarrow \sqrt{\frac{2 \log |H| + \log \frac{2}{\delta}}{2m}}$, $R(h) \leq \hat{R}_n(h) + O(\sqrt{\frac{\log |H|}{m}})$ $O(\frac{1}{\sqrt{m}})$

10. Infinite Hypothesis Set (之前 $\log |H| \rightarrow \infty$, 失效)
引 a: growth func b. VC dim $\Pi_H: N \rightarrow N$ for a hypothesis H , $\forall m \in N$
 $\Pi_H(m) = \max_{\{x_1, \dots, x_m\} \subseteq X} |\{h(x_1), \dots, h(x_m)\}|$ 最多的用 $h \in H$ 把 m 个点分成不同的数
若 H 为分类 $\Pr[|\hat{R}_n(h) - R(h)| \geq \epsilon] \leq 4 \Pi_H(2m) e^{-m\epsilon^2/8}$

$R(h) \leq \hat{R}_n(h) + \sqrt{\frac{2 \log \Pi_H(m)}{m}} + \log \frac{2}{\delta}$ 复杂 + 置信度 $A \subseteq B$, 则 $\text{VCdim}(A) \leq B$
11. VC Dim ① 一种 label 点的方式 ② 能 shattering 能被打破说明 H 能实现任意 Dichotomy

$\text{VCdim}(H) = \max \{m: \Pi_H(m) = 2^m\}$ $\Pi_H(m) = 2^m$ $\text{VCdim}(R_n^H) = n+1$ 超平面.
 $\text{VCdim}(\sin) = \infty$

12. Sauer's lemma, 令 H with $\text{VCdim}(H) = d$, $\forall m \in N, m < d$
 $\Pi_H(m) \leq \sum_{i=0}^d \binom{m}{i} \leq \sum_{i=0}^d \binom{m}{i} (\frac{d}{i})^{d-i} = (\frac{d}{0}) \sum_{i=0}^d \binom{m}{i} (\frac{d}{i})^{d-i} \leq \dots \leq (\frac{em}{d})^d$
要 $\forall \text{VCdim}(H) = d < \infty$, and $\Pi_H(m) = O(m^d)$ 要 $\forall \text{VCdim}(H) = \infty$, $\Pi_H(m) = 2^m$
令 $H \rightarrow \{-1, +1\}$ $\text{VCdim}(H) = d$, $\forall \delta > 0$. 至少有 $1-\delta$ 的 prob $|R(h) - \hat{R}_n(h)| \leq \sqrt{\frac{8d \log \frac{2}{\delta}}{n} + \log \frac{2}{\delta}}$
要 $|R(h) - \hat{R}_n(h)| \leq \epsilon$ 对 $\forall h \in H$, with prob p 成立, 满足 $m = O(d)$. 模型复杂 $\rightarrow$ 数据多
有限 V 和 H $|\hat{R}_n(h) - R(h)| \leq \sqrt{\frac{2 \log |H| + \log \frac{2}{\delta}}{2m}}$, $\forall h \in H$ $R(h) \leq \hat{R}_n(h) + O(\sqrt{\frac{\log |H|}{m}})$
无 $H \neq c$ consist: $R(h) \leq \frac{1}{m} \log(1/\delta) + \log \frac{1}{\epsilon}$ inconsistent $|\hat{R}_n(h) - R(h)| \leq \sqrt{\frac{2 \log |H| + \log \frac{2}{\delta}}{2m}}$

Lec 3 神经网络信号/GD BP 层间消失爆
1. 初始化 Xavier $w \sim U[-\frac{1}{\sqrt{n}}, \frac{1}{\sqrt{n}}]$ n : input dim 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 负值无, 要更大方差
2. Act ① sigmoid $\sigma(z) = \frac{1}{1 + \exp(-z)}$ $\sigma'(z) = \sigma(z)(1 - \sigma(z))$ vanishing
② ReLU $\max(0, x)$ 快训, 但易 dead: 若某 neuron 对所有输入 output 均为 0.
③ ReLU 变种 $g(z, \alpha) = \max\{0, z\} + \alpha \min\{0, z\}$ $\alpha = 1$ absolute value rectification
 $\alpha = 0.01$ leaky α 可学 parametric
④ tanh 零中心, 1/saturate $g(z) = \frac{\exp(z) - \exp(-z)}{\exp(z) + \exp(-z)} \in [-1, 1]$, u sigmoid
梯度大, 当 $|z| \rightarrow \infty$, saturate $g'(z) = 1 - z^2$
⑤ GeLU: $g(z) = z \phi(z)$ ϕ 是 Gaussian CDF $\uparrow$ softplus $\log(1 + \exp(z))$
⑥ Mish $g(z) = z \tanh(\log(1 + \exp(z))) = z \tanh(\text{softplus}(z))$
⑦ Swish $g(z) = z \sigma(\beta z)$ β fixed / 可学
共同特点: 平滑 负半轴 $\neq 0$, 非单调, 上方无界, 一般 ReLU 不稳

3. 输出层 prob 分布
3.1 linear gaussian $z = W^T h + b$ $p(y|x) = N(y|z, I)$ $L(x, y, \theta) = -\log p(y|x, \theta)$
 $d \pm \|y - z\|_2^2 = \frac{\partial L}{\partial z_k} = z_k - y_k$
3.2 sigmoid = 分类 $p(y|x) = \text{Bernoulli}(y|6(z)) = 6(z)^y (1 - 6(z))^{1-y}$
 $L(x, y, \theta) = -[y \log 6(z) + (1-y) \log (1 - 6(z))]$
 $\frac{\partial L}{\partial z} = \frac{\partial L}{\partial 6(z)} \frac{\partial 6(z)}{\partial z} = (-\frac{y}{6(z)} + \frac{1-y}{1-6(z)}) \cdot 6(z)(1-6(z)) = 6(z) - y$
saturation 只发生在 model 正确且自信情况下 k 是正确标签
3.3 Softmax 多分类 $p(y=k|x) = \frac{\exp(z_k)}{\sum_j \exp(z_j)}$ $L = -\log p(y=k|x)$
对正石角 k $\frac{\partial L}{\partial z_k} = -1 + \frac{1}{\sum_j \exp(z_j)} = -1 + p(y=k|x) = -z_k + \log \sum_j \exp(z_j)$
预测时 $\frac{\partial L}{\partial z_k} = \hat{y}_k - 1(y=k)$ 或 $\frac{\partial L}{\partial z} = -1(y=k) + z$
3.4 BP
 $z_k = \sum_j u_j w_{kj}$ 只和 k 行有关 $u_j = g(z_j)$ $z_j = \sum_i u_i w_{ij}$
 $\frac{\partial L}{\partial w_{ij}} = \frac{\partial L}{\partial z_k} \frac{\partial z_k}{\partial w_{ij}} = u_j \delta_k$ $\delta_k = \frac{\partial L}{\partial z_k}$ 部分
隐藏 $\frac{\partial L}{\partial w_{ji}} = \sum_k \frac{\partial L}{\partial z_k} \frac{\partial z_k}{\partial w_{ji}} = \sum_k \frac{\partial L}{\partial z_k} \frac{\partial z_k}{\partial u_j} \frac{\partial u_j}{\partial w_{ji}} = \sum_k \delta_k w_{kj} g'(z_j) u_i = u_i (g'(z_j) \sum_k \delta_k w_{kj})$ 目标单元误差
定义 $\delta_j = g'(z_j) \sum_k \delta_k w_{kj}$ $\frac{\partial L}{\partial w_{ji}} = u_i \delta_j$ $\frac{\partial L}{\partial b_j} = \delta_j$
 $\uparrow$ 输入激活

4. Optimization
4.1 SGD 只看当前 ∇ , noise 大, 震荡, 狭长 valley, 收敛慢
4.2 Momentum $g \leftarrow \frac{1}{m} \nabla_{\theta} \sum L(f(x^{(i)}; \theta), y^{(i)})$
 $v \leftarrow \alpha v - \epsilon g$ $\theta \leftarrow \theta + v$ 过去平均指数 $\alpha = 0.5, 0.9, 0.99$
4.3 Adagrad $r \leftarrow r + g \circ g$ 对稀疏特征友好, 历史 ∇ 小的
 $\Delta \theta \leftarrow -\frac{\epsilon}{\delta + r} g$ $\theta \leftarrow \theta + \Delta \theta$ r 大, cons: r 单增, 学不到
4.3 RMSProp: 遗忘久远历史
 $r \leftarrow r + (1-\rho) g \circ g$ $\Delta \theta \leftarrow -\frac{\epsilon}{\sqrt{r}} g$ $\theta \leftarrow \theta + \Delta \theta$
指数滑动平均, 不会无限累积, 有 adaptive scaling 但无显示 momentum, noise 仍大

4.4 Adam: Momentum + RMSProp
 $s \leftarrow \rho s + (1-\rho) g$ $r \leftarrow \rho r + (1-\rho) g \circ g$ $\hat{s} = \frac{s}{1-\rho^t}$ $\hat{r} = \frac{r}{1-\rho^t}$
 $\Delta \theta \leftarrow -\epsilon \frac{\hat{s}}{\sqrt{\hat{r}} + \epsilon}$ $\theta \leftarrow \theta - \epsilon \frac{\hat{s}}{\sqrt{\hat{r}} + \epsilon}$
4.5 LR Scheduler: 固定 / step decay / expo decay / cosine warm start

4.6 Dropout 廉价 bagging $m_j \sim \text{Bernoulli}(p)$ $\hat{y} = m_j u_j$
4.7 convex $O(\frac{1}{\sqrt{t}})$ strongly convex $O(\frac{1}{t})$
 $z^{(l)} = W^{(l)} a^{(l-1)}$ forward $a^{(l)} = f(z^{(l)})$
backward $\delta^{(l)} = (W^{(l+1)})^T \delta^{(l+1)} \odot f'(z^{(l)})$
 $\frac{\partial L}{\partial w^{(l)}} = \delta^{(l)} (a^{(l-1)})^T$
上一步激活值 和 $w^{(l)}$ 的输入

L4 CNN 输入 filter 位置
1. Convolution $(f * g)(t) = \int_{-\infty}^{+\infty} f(t-\tau)g(\tau)d\tau = \int_{-\infty}^{+\infty} f(t-\tau)g(\tau)d\tau$
 $(f * g)(n) = \sum_{m=-\infty}^{+\infty} f[m]g[n-m] = \sum_{m=-\infty}^{+\infty} f[n-m]g[m]$ 卷积核翻转 $k[m,n]$
2. 2D $Z(i,j) = (I * k)(i,j) = \sum_m \sum_n I(m,n)k(i-m,j-n) = \sum_m \sum_n I(i-m,j-n)k(m,n)$
自相关 $Z(i,j) = (I * k)(i,j) = \sum_m \sum_n I(i+m,j+n)k(m,n)$
DL中 kernel 权重可学, 不翻转无所谓

3. 三大优势 { Sparse connectivity: 一个神经元只连 receptive field 里的
parameter sharing: 不同空间相同权重
Equivariance: 先移再卷 = 先卷再移
4. $Z_{j,k} = \sum_{i,m,n} I_{j+m-1,k+n-1} K_{m,n} + b^{(i)}$ 加 stride
 $Z_{j,k} = \sum_{i,m,n} I_{(j-1)s+m,(k-1)s+n} K_{m,n} + b^{(i)}$ 一个是 $Cin F_n F_w$
 $H_{out} = \lfloor \frac{H+2P-F_n}{2} \rfloor + 1$ 参数量 $Cin F_n F_w (C_{out} + C_{out})$
FLOPs = $2 \times K_h \times K_w \times C_{in} \times H_{out} \times W_{out}$ 或再加 $H_{out} \times W_{out}$ 次 bias 相加

5. 反传 $J(k,b)$ 设上一层传回梯度 $G_{j,k}^{(i)} = \frac{\partial J(k,b)}{\partial Z_{j,k}^{(i)}}$ 是第 i 个 output channel
对 kernel $\frac{\partial J(k,b)}{\partial K_{j,k}^{(i)}} = \sum_{m,n} \frac{\partial J}{\partial Z_{mn}^{(i)}} \frac{\partial Z_{mn}^{(i)}}{\partial K_{j,k}^{(i)}} = \sum_{m,n} G_{m,n}^{(i)} I_{(m-1)s+j,(n-1)s+k}^{(i)}$ (是输入 channel)
对 input $\frac{\partial J(k,b)}{\partial I_{p,q}^{(i)}} = \sum_{j,k} K_{j,k}^{(i,l)} G_{j,k}^{(i)}$ 每次看到的输入值都
对 bias $\frac{\partial J(k,b)}{\partial b^{(i)}} = \sum_{j,k} G_{j,k}^{(i)}$ 一个输入像素 $I_{p,q}$ 会被多个滑窗覆盖, 凡覆盖的都要把对应 kernel 的权重乘上游梯度加起来
6. Pooling { max 平均
l2-norm 能量
prob weighted 对平移不敏感, 降计算, 关注特征而非位置

7. BN $\mu_j = \frac{1}{n} \sum x_{ij}$ $\sigma_j^2 = \frac{1}{n} \sum (x_{ij} - \mu_j)^2$ $\hat{x}_{ij} = \frac{x_{ij} - \mu_j}{\sigma_j}$ $\hat{y} = \gamma \hat{x} + \beta$
轻约束正则, 降初始化敏感, 稳深层
8. NIN conv $\rightarrow$ 1x1 conv $\rightarrow$ 1x1 conv GAP, $S_c = \frac{1}{H \times W} \sum A_{c,i,j}$ 作为
9. Dense: 所有前层 output 全拼接进下层类分数, 减参, 不用 MLP, overfit
10. ResNet X1 多路径卷再拼 对位置 robust, 可解系并
11. SENet 通道注意力, Nas Next 自动找

L5 RNN
1. RNN $h^{(t)} = \tanh(b + W h^{(t-1)} + U x^{(t)})$ $o^{(t)} = c + V h^{(t)}$ $\hat{y}^{(t)} = \text{softmax}(o^{(t)})$
Loss $-\sum_{t=1}^T \log p(y^{(t)} | x_{1:t})$ 对所有 timestamp 求和
2. BPTT $\nabla_W L$ $\nabla_U L$ $\nabla_b L$ $\nabla_c L$ $\nabla_{\theta_m} L$ $\hat{y}_i(t) = \frac{\exp(o_i(t))}{\sum_j \exp(o_j(t))}$
 $L^{(t)} = -\log \hat{y}^{(t)}(y(t))$ $\frac{\partial L}{\partial o_i(t)} = \hat{y}_i(t) - 1$ i, y(t) prob 与真实 one-hot 标
 $h^{(t)}$ 到总 loss 通过 $o^{(t)} \rightarrow L^{(t)}$ 或 $\rightarrow h^{(t+1)} \rightarrow L^{(t+1)}$ 等
定义 $\delta^{(t)} \equiv \nabla_{a^{(t)}} L$ $\frac{\partial h^{(t)}}{\partial a^{(t)}} = \text{diag}((1-h^{(t)})^2) \Rightarrow \delta^{(t)} = \text{diag}(1-h^{(t)})^2 \nabla h^{(t)}$
 $\Rightarrow \delta^{(t)} = \text{diag}(1-h^{(t)})^2 [V^T \nabla o^{(t)} L + W^T \delta^{(t+1)}]$ 且 $\delta^{(t+1)} = 0$ 边界
 $\Rightarrow \nabla_b L = \sum_{t=1}^T \delta^{(t)}$ $\nabla_W L = \sum_{t=1}^T \delta^{(t)}(h^{(t+1)})^T$ $\nabla_U L = \sum_{t=1}^T \delta^{(t)}(x^{(t)})^T$
3. Teacher forcing: 把 $y^{(t)}$ 当输入 $\lambda(\log p(y_{1:2} | x_{1:2}) - \log p(y_{2:2} | x_{1:2}))$
+ log p(y, x1:2) 好处时间步间完全 decode $\Rightarrow$ 可以并行

4. 梯度消失/爆炸
线 RNN $h^{(t)} = W h^{(t-1)} + U x^{(t)} + b$ $h^{(t)} = f(W h^{(t-1)} + U x^{(t)} + b)$
 $\frac{\partial h^{(t)}}{\partial h^{(t-1)}} = \frac{\partial f}{\partial \text{net}} \frac{\partial \text{net}}{\partial h^{(t-1)}} = (W^T)^{t-k}$ $\frac{\partial h^{(t)}}{\partial h^{(j)}} = \frac{\partial f}{\partial \text{net}} W^T \text{diag}(f'(h^{(j-1)}))$
若 $W^T = V \text{diag}(\lambda) V^T$ $\leq \|W^T\| \|\text{diag}(\dots)\|$
 $\|\lambda\|_2 \rightarrow \infty$ 爆 非线性激活
 $\|\lambda\|_2 \rightarrow 0$ 消 缓解火爆
② 若 f' 很小, sigmoid/tanh $\rightarrow$ 消失

5. Solu { Gradient clipping $\hat{g} = \bar{v} \otimes g$ $\hat{g} \leftarrow \begin{cases} \hat{g}, & \|\hat{g}\| < r \\ \frac{r}{\|\hat{g}\|} \hat{g}, & \|\hat{g}\| > r \end{cases}$ 防止更新过大
Identity init + ReLU $W = I + \text{ReLU}$
6. LSTM forget $f^{(t)} = \sigma(W^f h^{(t-1)} + U^f x^{(t)} + b^f)$ 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
output gate $o^{(t)} = \sigma(W^o h^{(t-1)} + U^o x^{(t)} + b^o)$
cell state $c^{(t)} = f^{(t)} \odot c^{(t-1)} + i^{(t)} \odot \tanh(W^c h^{(t-1)} + U^c x^{(t)} + b^c)$ candidate values
final hidden state $h^{(t)} = o^{(t)} \odot \tanh(c^{(t)})$

7. GRU Update $z^{(t)} = \sigma(W^z h^{(t-1)} + U^z x^{(t)} + b^z)$ 留多少旧 hidden state
Reset $r^{(t)} = \sigma(W^r h^{(t-1)} + U^r x^{(t)} + b^r)$ 出成候选是否忘过去
New memory content: $\tilde{h}_t^* = \tanh(W^h r^{(t)} \odot h^{(t-1)} + U^h x^{(t)} + b^h)$ 候选记忆
Final memory: $h^{(t)} = z^{(t)} \odot h^{(t-1)} + (1-z^{(t)}) \odot \tilde{h}_t^*$ 让 gradient 能沿 cell state 加法路径传播
GRU 参少简单训练 LSTM 大任务强, 长程依赖
8. 双向 $\hat{g} = g(Uh + c) = g(U[\tilde{h}_t^* \parallel \tilde{h}_t] + c)$, 可堆多层
9. Attention $\text{Ply}(y_{1:t-1}, x) = g(y_{1:t-1}, s_t, c_t)$ $s_t = f(s_{t-1}, y_{t-1}, c_t)$
 $c_t = \sum_{j=1}^T \alpha_{t,j} h_j$ $\alpha_{t,j} = \frac{\exp(e_{t,j})}{\sum_k \exp(e_{t,k})}$ $e_{t,j} = a(s_{t-1}, h_j)$
10. 难并行, BP through seq. 局部信息与全局信息通过 bottleneck
11. 单层 $O(nd^2)$ vs $O(nd)$ 顺序操作 $O(n)$ vs $O(1)$ 最短 $O(n)$ vs $O(1)$
12. Mamba $h_t = A(x_t)h_{t-1} + B(x_t)x_t$ $y_t = C(x_t)h_t$
重要 token 强更 memory 不重要不改, inside a structured state-space evolution

Lec 6 代 $y - h_s = y - E_s[h_s] + E_s[h_s] - h_s$
1. 分解 $\hat{z} = E_s[E(x,y) | (y - E_s[h_s])^2] = E_s[E(x,y) | (y - E_s[h_s] + E_s[h_s] - h_s)^2]$
 $= E_s[E(x,y) | (y - E_s[h_s])^2] + E_s[E(x,y) | (E_s[h_s] - h_s)^2]$
 $= \text{Bias}^2 + \text{Variance}$ 交叉项为 0, 因为 $E_s[E(x,y) | (y - E_s[h_s]) (E_s[h_s] - h_s)] = 0$

2. Ridge $J(w, w_0) = \frac{1}{n} \sum [y_i - (w_0 + w^T \phi(x_i))]^2 + \lambda \|w\|_2^2$ 固定 x, y , 前面是常数
Lasso $J(w, w_0) = \frac{1}{n} \sum [y_i - (w_0 + w^T \phi(x_i))]^2 + \lambda \|w\|_1$ $\hat{J}(\theta) = J(\theta) + \alpha \|\theta\|_1$
3. L2 正则/weight decay 等价于对权重加 Gaussian prior 后做 MAP
 $\hat{y}(w) = J(w) + \frac{\alpha}{2} \|w\|_2^2$ $w \leftarrow w - \eta (\nabla_w J(w) + \alpha w) = (1 - \eta \alpha) w - \eta \nabla_w J(w)$ 在未正则化 $\rightarrow$ 近似
 $f(x) \approx f(x_0) + (x - x_0)^T f'(x_0) + \frac{1}{2} (x - x_0)^T f''(x_0) (x - x_0)$ $\hat{y}(w) = J(w^*) + (w - w^*)^T \nabla_w J(w^*) + \frac{1}{2} (w - w^*)^T H(w^*) (w - w^*)$
 $= J(w^*) + \frac{1}{2} (w - w^*)^T H(w - w^*)$ H 与 w^* 有关, w^* 确定, 再求 $\nabla_w J(w) = H(w - w^*) = 0$
加 λI 后, $\alpha \tilde{w} + H(\tilde{w} - w^*) = 0 \Leftrightarrow \tilde{w} = (H + \alpha I)^{-1} H w^*$ 当 $\alpha \rightarrow 0$ $\tilde{w} \rightarrow w^*$ 对分解 $H = Q \Lambda Q^T$ 取
 $\tilde{w} = (Q \Lambda Q^T + \alpha I)^{-1} Q \Lambda Q^T w^* = Q (\Lambda + \alpha I)^{-1} \Lambda Q^T w^*$ 用了 $(\Lambda B C)^{-1} = C^{-1} A^{-1} B^{-1}$ w
在每一个特征方向上, 权重被缩放为 $\frac{\lambda_i}{\lambda_i + \alpha}$ α 对低曲率方向 (动力一点, 对 loss 影响小, 不重要) 惩罚强, 在高曲率 (重要, 影响 loss) 方向影响弱, $w^* = (X^T X + \alpha I)^{-1} X^T y$ 防止 $X^T X$ 奇异

4. L2-norm $\nabla_w J(w) = \nabla_w J(w) + \lambda \text{sign}(w)$ 不可导, 用 subgradient. 对 $w^* = 0$ 附近似, 设 H 对用
 $\hat{y}(w) = J(w^*) + \sum \pm H_{ii}(w_i - w_i^*)^2 + \alpha |w_i|$ 对每个 dim $w_i = \text{sign}(w_i^*) \max\{|w_i|, \frac{\alpha}{2H_{ii}}\}$ 且元素 > 0
soft-thresholding $|w_i| \leq \frac{\alpha}{2H_{ii}} \rightarrow w_i = 0$ 否则 $|w_i| = |w_i^*| - \frac{\alpha}{2H_{ii}}$ 向 0 移动
 $\log p(w) = \sum \log \text{Laplacian}(w_i; 0, \alpha) = -\alpha \|w\|_1 + n \log \alpha - n \log 2$
5. Norm Penalty 作为 constrained optimization min $J(\theta)$, s.t. $\|\theta\|_1 \leq K$ 开拉
 $L(\theta, \alpha) = J(\theta) + \alpha (\|\theta\|_1 - K)$ $\alpha \geq 0 \Rightarrow \theta^* = \argmin_{\theta} \max_{\alpha \geq 0} L(\theta, \alpha)$
6. Data Aug: random crop: crop $\rightarrow$ resize Prop block 遮挡, 用灰色随机颜色
7. Noise Injection 等效对权重加惩罚. 对 input 加为 Data aug
对 weight 加 $\tilde{w} = w + \epsilon w$ $\epsilon w \sim N(0, \sigma^2)$ $\hat{y}(x) = f_w(x)$ 变成 $\hat{y}_{\tilde{w}}(x) = f_{w+\epsilon w}(x)$
 $\tilde{J}_w = E_{\epsilon} [L(x, y, \tilde{w})] = E_{\epsilon} [L(x, y, w) + 2\epsilon y^T w + \epsilon^2 L(x, y, w)] = E_{\epsilon} [L(x, y, w) + 2\epsilon y^T w + \epsilon^2 L(x, y, w)]$
 $\nabla_w \tilde{J}_w(x) = \nabla_w L(x, y, w) + 2y + 2\epsilon \nabla_w L(x, y, w)$ 当 η 很小, 对 $y^T w$ $\approx y^T w(x) + \nabla_w L(x, y, w)^T \epsilon w$ 值近似
交叉项依赖 ϵw , $E[\epsilon w] = 0$ 消去, 后面平方项变成 $\eta E_{\epsilon} [L(x, y, w) + 2\epsilon y^T w + \epsilon^2 L(x, y, w)]$ $\tilde{J}_w \approx J_w + \eta E_{\epsilon} [L(x, y, w) + 2\epsilon y^T w + \epsilon^2 L(x, y, w)]$
后面项 measure 输出对权重敏感度 (6. 惩罚 $\|w\|_2^2$ 权重稍变, output 变 $\| \nabla_w \hat{y}(x) \|^2$)
可以看作是 Bayesian inference over weights 的 stochastic implementation
 w 应是一个分布 $p(w)$, 精确求解很难, 近似做法训练时随机采样 $w + \epsilon w$, 不让 model
建立在一个极端点上, 不要 one-hot label 过度自信

8. Output Noise: Label Smoothing $y_c = 1 - \frac{\epsilon}{K}$ $y_i = \frac{\epsilon}{K}$, $i \neq c$ 错误类 $L_{LS} = -\sum_{j=1}^K y_j \log p(j|x)$
Lec 7 Transformer Score = $\text{softmax}(\frac{QK^T}{d_k}) V$ pos emb $w = \frac{1}{10000} \frac{2\pi i d_m}{d_m}$
PE pos, $z_i = \sin(W_{pos} i)$ PE pos, $z_i = \cos(W_{pos} i)$ $\begin{bmatrix} \sin(w_i(\text{pos}+k)) \\ \cos(w_i(\text{pos}+k)) \end{bmatrix} = \begin{bmatrix} \cos(w_i k) \sin(w_i \text{pos}) \\ -\sin(w_i k) \cos(w_i \text{pos}) \end{bmatrix} = \begin{bmatrix} \sin(w_i \text{pos}) \\ \cos(w_i \text{pos}) \end{bmatrix} \begin{bmatrix} \sin(w_i k) \\ -\cos(w_i k) \end{bmatrix}$
BERT $P(x_i | x_{-i}) = \text{softmax}(W_i^T b)$ LM = $-\sum_{i=1}^M \log P(x_i^* | x_{-i}^*)$ Next Sentence prediction: 50% 在 A 后, 50% 随机, 使用 [CLS] 输出向量 c 预测
 $P(\text{Is Next} | A, B) = \text{softmax}(W_c b)$ NSP Loss $L_{NSP} = -\log P(y_{NSP} | c)$
LBERT = LM + LNSP classification 用 [CLS] token-level task
每个 token 的输出 $\hat{y}_i = \text{softmax}(W_i^T b)$ T_i 第 i-th token 表示, 还有 segment embedding 区分两句子
2. ViT, 一个 patch $R \times (P^2 c)$ $N = HW/P^2$ 需加可学习 token x_{class} 无 embedding
CNN 有 locality, 2D neighbourhood structure, translation equivariance
ViT 里只有 MLP 是 local 且 translationally equivalent, self-attention 是 global 的

3. CLIP: Contrastive Language-Image Pretraining 图文都 embed 进同一个 embedding space
让匹配的相似, 不匹配的最远 Loss. 对一个 patch, 含 N pair Z_i, T_i
 $L_{CLIP} = -\frac{1}{2N} \sum_{i,j} \log \frac{\exp(Z_i \cdot T_j)}{\sum_k \exp(Z_i \cdot T_k)} - \frac{1}{2N} \sum_{i,j} \log \frac{\exp(T_i \cdot Z_j)}{\sum_k \exp(T_i \cdot Z_k)}$ image 对多 text
zero shot: 先 create prompt Prob = $\frac{\exp(I \cdot T_c)}{\sum_k \exp(I \cdot T_k)}$ text 对多 image
4. encoder-decoder masked span prediction $\sum_{c=1}^C \exp(I \cdot T_c)$ argmax 找最匹配
 $L_{span} = \sum_{c=1}^C \log p(y_{c+1:t} | y_{1:c}, x)$ $\hat{x}$ 是被 mask span 替换的序列, y 是要恢复的序列
5. LLaMa: Layer Norm $\rightarrow$ RMSnorm ReLU $\rightarrow$ SwiGLU MHA $\rightarrow$ GQA
6. 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 r, size benefit 明显
7. LoRA $W' = W + \Delta W = W + B A$ $W \in \mathbb{R}^{d \times k}$ $B \in \mathbb{R}^{d \times r}$ $A \in \mathbb{R}^{r \times k}$ A 是 Gaussian

Lec 9 VAE & Diffusion
1. Intro. 两类 { explicit, 显式建模 GMM 生成 $\rightarrow$ 处理缺失 data, 对高维数据 handle
implicit GAN 生成 $\rightarrow$ 结合 RL 生成 env
2. Autoencoder 无监督, target 作为输入本身 $x \in \mathbb{R}^n \rightarrow h \in \mathbb{R}^d \ll n$ 用于降维
encoder $h = f_{\phi}(x) = a(Wx + b)$ decoder $z = g_{\theta}(h) = a(W_h h + b')$
 $\theta^*, \phi^* = \argmin_{\theta, \phi} \frac{1}{n} \sum L(x_i^*, z_i^*)$ $g_{\theta}(f_{\phi}(x^{(i)}))$ $L(x, \hat{x}) = \|x - \hat{x}\|_2^2$ 或 CE
undercomplete 中间维度比输入小, overcomplete 大, tied weight decoder $W' = W^T$
要加 regularization: sparsity, denoising, contractive penalty
3. Linear AE 与 PCA 关系. 在线性 encoder, l2 重建 error 下, 最优 linear autoencoder
学到的子空间等价于 PCA 生成的子空间. $A = U \Sigma V^T$ $\theta^* = \argmin \|A - B\|_F^2 = U_{:k} \Sigma_{:k} V_{:k}^T$
AE 平方误差满足 $\min_w \sum (x_i^{(t)} - z_i^{(t)})^2 = \min_w \frac{1}{2} \|x - W'h\|_2^2$ rank(B) = k $\Sigma = \frac{1}{n} V^T B^T B V$
最优解 $W' \leftarrow U_{:k} \Sigma_{:k}^{-1} V_{:k}^T$ $W' \leftarrow \frac{1}{n} \sum_{i=1}^n x_i x_i^T$ $\hat{x} = \frac{1}{n} X^T x$ centered
正是 PCA $x^{(t)} \leftarrow \frac{1}{n} (x^{(t)} - \frac{1}{n} \sum_{i=1}^n x^{(i)})$ $X_{centered} = X - \mu X$ $\mu = \frac{1}{n} \sum_{i=1}^n x^{(i)}$

L9 4. AE → VAE
4.1 形式规定 latent space 的概率分布为破碎连续、离散
训练时 decoder 只见过 encoder 出来的, 没见过任意条件的
4.2 $z \sim p(z)$ $x \sim p(x|z)$ $p(z) \sim N(0, I)$ Decoder 建模 $p(x|z) \sim N(\mu, \sigma^2)$
希望最大化 $p(x) = \int p(x|z) p(z) dz$ 对所有 z , intractable. 且 $p(x|z)$ 不可解
引入 encoder network $q(z|x)$ 来近似 $p(x|z)$. 设 $q(z|x) = N(\mu_z(x), \sigma_z(x))$
4.3 ELBO

$\log p(x) = E_{z \sim q(z|x)} [\log p(x|z)] = E_{z \sim q(z|x)} [\log \frac{p(x, z)}{p(z|x)}]$ 贝叶斯规则
 $= E_{z \sim q(z|x)} [\log \frac{p(x, z)}{q(z|x)} \frac{q(z|x)}{p(z|x)}]$ 上下同乘 $q(z|x)$
 $= E_{z \sim q(z|x)} [\log \frac{p(x, z)}{q(z|x)}] + D_{KL}(q(z|x) || p(z|x))$ 拆开
 $\geq E_{z \sim q(z|x)} [\log \frac{p(x, z)}{q(z|x)}] = L(\theta, \phi, x)$ 最大化 ELBO 最大
 $= E_{z \sim q(z|x)} [\log p(x|z) p(z)]$ 间接上 approximate posterior
 $= E_{z \sim q(z|x)} [\log \frac{p(z)}{q(z|x)}] + E_{z \sim q(z|x)} [\log p(x|z)]$

$= -D_{KL}(q(z|x) || p(z|x)) + E_{z \sim q(z|x)} [\log p(x|z)]$
encoder 输出不要背向 prior 重建损失
 $= \frac{1}{2} \sum_j (1 + \log \hat{\sigma}_j^2 - \mu_j^2 - \sigma_j^2)$ μ_j 是 j -th dim 均值 σ 方差
1. 正向令 $\mu_j \rightarrow 0$ $\sigma_j^2 \rightarrow 1$; $z \sim q(z|x)$ 有随机节点, 化为
sample $z \sim N(0, I)$ $z = \mu(x) + \epsilon$
 $\Rightarrow L(\theta, \phi, x) \approx \frac{1}{2} \sum_j (1 + \log \hat{\sigma}_j^2 - \mu_j^2 - \sigma_j^2) + \frac{1}{2} \sum_j \log p(x|z^{(k)})$
前一项对 latent dim 求和, 后一项对采样次数求和, 用 SMC

4.4 posterior collapse: 训练 encoder 学到的 posterior 退化为 prior
几乎不含信息. encoder 虽然存在, 但被忽略. KL 变成 0, 正好喜欢
解决 KL annealing. 初期减小 KL 权重, 学会重构后慢慢加大. 限制
decoder capacity 也行

4.5 VQ-VAE: 离散 latent 的 representation. 一个格子一个代码.
 $q(z = k|x) = \begin{cases} 1, & k = \arg \min_l \|z(x) - e_l\|_2 \\ 0, & \text{otherwise} \end{cases}$ encoder 输入量化到
最近格子
 $L = \log p(x|z) + \log q(z|x) - \text{ell}_2^2 + \beta \|z(x) - e_l\|_2^2$
sg. L 前传 identity, 后传为 0 (code book commitment)
让 z 靠近 encoder 输出 令其输出靠近 emb

5. Diffusion
5.1 forward $x_0 \sim q(x)$ $x_t = \sqrt{1 - \beta_t} x_{t-1} + \sqrt{\beta_t} \epsilon_t$ $\epsilon_t \sim N(0, I)$
 $q(x_t|x_{t-1}) = N(x_t; \sqrt{1 - \beta_t} x_{t-1}, \beta_t I)$ $q(x_{1:T}|x_0) = \prod_{t=1}^T q(x_t|x_{t-1})$
 $q(x_0; x_T) = q(x_0) \prod_{t=1}^T q(x_t|x_{t-1})$ $z \sim N(\mu_x + \mu_y, \sigma_x^2 + \sigma_y^2)$
定义 $x_t = \sqrt{1 - \beta_t} x_0 + \sqrt{1 - \beta_t} \epsilon_t$ $x_t \sim N(0, I)$
 $q(x_t|x_0) = N(x_t; \sqrt{1 - \beta_t} x_0, (1 - \beta_t) I)$ t.f. $\beta_t \rightarrow 0$

5.2 Noise Schedule: linear $\beta_t = \beta_1 + \frac{t-1}{T-1}(\beta_T - \beta_1)$ $q(x_t|x_0) \approx N(0, I)$
quadratic $\beta_t = [\frac{T+1}{T-1} + \frac{t-1}{T-1}(\frac{\beta_T - \beta_1}{T-1})]^2 \beta_T - \beta_1$
cosine $\beta_t = \frac{f(t)}{f(T)}$ $f(t) = \cos^2(\frac{\pi t}{T})$ sigmoid $\beta_t = \beta_1 + \frac{\beta_T - \beta_1}{1 + \exp(-\frac{t-1}{T-1})}$
5.3 反向: 一步步从小 Gaussian $x_T \sim N(0, I)$ $x_{t-1} = \mu_\theta(x_t, t) + \sigma_t \epsilon_t$
 $p_\theta(x_{t+1}|x_t) = N(x_{t+1}; \mu_\theta(x_t, t), \sigma_t^2 I)$
5.4 为什么 noise prediction 等价 reverse mean
 $q(x_{t+1}, x_t|x_0)$ 是高斯分布, 有其均值 $\mu_\theta(x_t, x_0) = \frac{\sqrt{1 - \beta_t}}{1 - \beta_t} x_0 + \frac{\sqrt{1 - \beta_t}}{1 - \beta_t} \epsilon_t$
底 $(x_t, x_0) = \frac{\sqrt{1 - \beta_t}}{1 - \beta_t} (x_t - \frac{\beta_t}{1 - \beta_t} \epsilon_t)$ 不妨让模型预测 $\mu_\theta(x_t, t) = \frac{1}{\sqrt{1 - \beta_t}} (x_t - \frac{\beta_t}{1 - \beta_t} \epsilon_t)$

最小化 $\| \hat{\mu}_t - \mu \|$ 等价于让 $\hat{\mu}_t$ 接近 μ . 等价于让 $\hat{z}_0(x_t, t) = z$
 $D_{KL} \propto \| \hat{\mu}_t - \mu \|^2 \Rightarrow E_{x_0, \epsilon_t} [\| \hat{z}_0(x_t, t) - z \|^2]$ $q(x_t|x_{t+1}) = \frac{q(x_t|x_0)q(x_0|x_t)}{q(x_{t+1}|x_0)}$
怎么得到? $q(x_{t+1}|x_t, x_0) = \frac{q(x_t|x_{t+1})q(x_{t+1}|x_0)}{q(x_t|x_0)}$
固定 x_t 取条件期望 $E_{x_0 \sim q(x_0|x_t)} [\nabla_{x_t} \log q(x_0|x_t)] = \nabla_{x_t} [q(x_0|x_t) dx_0] = 0$
 $S_0(x_t, t) \approx \nabla_{x_t} \log q(x_t|x_0)$ 或 $\nabla_{x_t} \log q(x_t|x_0) = E_{x_0} [\nabla_{x_t} \log q(x_t|x_0)]$
代入 $x_t = y_t x_0 + b_t \epsilon_t$ $\nabla_{x_t} \log q(x_t|x_0) = -\nabla_{x_t} [\frac{1}{2} (x_t - y_t x_0)^2]$
 $\Rightarrow S_0(x_t, t) = -\frac{\epsilon_t}{\sigma_t^2} (x_t - y_t x_0) = -\frac{\epsilon_t}{\sigma_t^2} x_t$
本质仍是 noise, 差值 $\frac{\epsilon_t}{\sigma_t^2}$ $\Rightarrow \min_{\theta} E_{x_0, \epsilon_t} [\| \hat{z}_0(x_t, t) - z \|^2]$

4. Probability Flow ODE 反向生成 SDE 在分布意义等价于 TODE
 $dx_t = -\frac{1}{2} \beta(t) [x_t + \nabla_{x_t} \log q(x_t|x_0)] dt$ 解这个 ODE 等价于 $q(x_t|x_0)$
当初初始 $q_T(x_T) \approx N(x_T; 0, I)$. 可用 ODE solver 求解
5. Flow Matching
Flow Model $x_0 \sim p_{init}$, ODE: $dx_t = U_t(x_t) dt$ 到 $t=1$ return
Diffusion $x_0 \sim p_{init}$ SDE $dx_t = U_t^0(x_t) dt + \sigma_t dW_t$
5.1 Time dependent vector field
 $U: \mathbb{R}^d \times [0, 1] \rightarrow \mathbb{R}^d$ $(x_t) \mapsto U_t(x_t)$ 轨迹 $x_t = U_t(x_t) \Leftrightarrow \frac{dx_t}{dt} = U_t(x_t)$
瞬时速度场, 输入 x 和 t , 输出一个运动向量和速度. 沿场把分布过去
给定 point $z \in \mathbb{R}^d$, 定义 cond prob path $p_t(\cdot|z)$ 满足从 noise - 跑到 目标 z
 $p_0(\cdot|z) = p_{init}$ $p_1(\cdot|z) = \delta_z$ 集中在 z 的 Dirac delta distribution
若用 Gaussian $p_t(\cdot|z) = N(\mu_t, \Sigma_t)$ 控制 μ_t 从 0 到 z , 方差从 1 到 0
边界 $\mu_0 = 0, \mu_1 = z \Rightarrow p_0(\cdot|z) = N(0, I)$ $p_1(\cdot|z) = N(z, 0) = \delta_z$
5.2 conditional probability path
5.3 conditional vector field 对任意 cond prob path 都有在等价 vector field
s.t. $x_0 \sim p_{init}$ $\frac{dx_t}{dt} = U_t(x_t|z) \Rightarrow x_t \sim p_t(\cdot|z) = N(\mu_t, \Sigma_t)$
 $U_t(x|z) = (\dot{\mu}_t - \frac{\beta_t}{\sigma_t^2} \dot{\Sigma}_t \mu_t) + \frac{\beta_t}{\sigma_t^2} \dot{\Sigma}_t x$ $\dot{\mu}_t = \frac{d\mu_t}{dt}$ $\dot{\Sigma}_t = \frac{d\Sigma_t}{dt}$
推导: 若有 $x_t = \mu_t + \beta_t \epsilon_t, \epsilon_t \sim N(0, I)$. 则 $\frac{dx_t}{dt} = \frac{d\mu_t}{dt} + \beta_t \frac{d\epsilon_t}{dt}$
整理 $U_t(x|z) = \frac{dx_t}{dt} = (\dot{\mu}_t - \frac{\beta_t}{\sigma_t^2} \dot{\Sigma}_t \mu_t) + \frac{\beta_t}{\sigma_t^2} \dot{\Sigma}_t x$

3.6 Train: 取圈 $t, \epsilon, \epsilon_0 \| \epsilon - \epsilon_0 (\sqrt{1 - \beta_t} x_0 + \sqrt{1 - \beta_t} \epsilon_t) \|^2$
Sample: $x_T \sim N(0, I)$ 对每一步 $z \sim N(0, I)$ $x_{t-1} = \frac{1}{\sqrt{1 - \beta_t}} (x_t - \frac{\beta_t}{1 - \beta_t} \epsilon_0(x_t, t))$
返回 x_0 . $\epsilon_0(x_t, t) \mapsto \text{embedding} \rightarrow \text{fully-connected}$ 注入
Loss $L = E_{x_0 \sim q(x_0), t \sim U(0, 1), \epsilon_t \sim N(0, I)} [\| \epsilon_t - \epsilon_0(\sqrt{1 - \beta_t} x_0 + \sqrt{1 - \beta_t} \epsilon_t, t) \|^2]$
 ϵ_t is the compounded noise added to x_0 to obtain x_t DDIM
理论 $\beta_t^2 = \hat{\beta}_t^2 = \frac{1 - \alpha_{t-1}}{1 - \alpha_t} \beta_t$ 但实践 $\beta_t^2 = \beta_t$. 若 $\beta_t \rightarrow 0$ DDPM. $\beta_t \rightarrow 0$

Lec 10
1. ODE SDE
ODE: $\frac{dx}{dt} = f(x, t)$ or $dx = f(x, t) dt$ 解析解 $x(t) = x(0) + \int_0^t f(x, \tau) d\tau$
迭代 $x(t+\Delta t) \approx x(t) + f(x(t), t) \Delta t$ 确定性轨迹. 给定 $x(0)$ 后未来状态完全确定
SDE $\frac{dx}{dt} = f(x, t) + g(x, t) dW_t$ $dx = f(x, t) dt + g(x, t) dW_t$ 若每次运行都
 $x(t+\Delta t) \approx x(t) + f(x(t), t) \Delta t + g(x(t), t) \sqrt{\Delta t} N(0, I)$ 不同, 引随机性
2. 离散 diffusion $\rightarrow$ 连续时间 ODE
2.1 forward diffusion 连续极限 $x_t = \sqrt{1 - \beta_t} x_{t-1} + \sqrt{\beta_t} \epsilon_t$ $\epsilon_t \sim N(0, I)$ 令 $\beta_t = \beta(t) \Delta t$
 $x_t \approx x_{t-1} - \frac{\beta(t) \epsilon_t}{\sigma_t} x_{t-1} + \sqrt{\beta(t) \Delta t} \epsilon_t$ 当 Δt 很小时, $\sqrt{\beta(t) \Delta t} \approx \frac{1}{2} \beta(t) \Delta t$
令 $\Delta t \rightarrow 0$ 得到前向 SDE $dx_t = -\frac{1}{2} \beta(t) x_t dt + \sqrt{\beta(t)} dW_t$
把 x_t 往 0 拉. 注入 noise
2.2 Reverse time SDE and score function
前向: $dx_t = -\frac{1}{2} \beta(t) x_t dt + \sqrt{\beta(t)} dW_t$ score = $\nabla_x \log q(x)$
反向: $dx_t = [-\frac{1}{2} \beta(t) x_t - \beta(t) \nabla_x \log q(x_t)] dt + \sqrt{\beta(t)} dW_t$
Proof: $dx_t = f(x_t, t) dt + g(t) dW_t$ $x_t + \Delta t = x_t + f(x_t, t) \Delta t + g(t) \sqrt{\Delta t} \epsilon_t$
 $p(x'|\Delta t) = N(x'; x_t + f(x_t, t) \Delta t, g(t)^2 \Delta t)$
 $p(x_t|x_{t+\Delta t}) p(x_{t+\Delta t}|x_t) = p(x_t|x_{t+\Delta t}) p(x_{t+\Delta t}|x_t)$
 $\log q(x) \approx \log q(x') + (x - x') \nabla \log q(x)$
 $E[x_t|x_{t+\Delta t}] = x_{t+\Delta t} - f(x_{t+\Delta t}, t) \Delta t + g(t) \nabla \log q(x_{t+\Delta t}) \Delta t$
 $\approx x_t + f(x_t, t) \Delta t - g(t)^2 \nabla_x \log q(x_t) \Delta t$ Gaussian completion
 $\approx x_t + f(x_t, t) \Delta t - g(t)^2 \nabla_x \log q(x_t) \Delta t$ 在 diffusion 中, 反向生成成本顺序每个 noise 下 $g = \sqrt{\beta(t)}$
的 score $S_0(x_t, t) \approx \nabla_x \log q(x_t)$

3. Score matching 但 score 不可算, 且 $q(x_0)$ 未知.
希望 $\min_{\theta} E_{x_t} [\| S_0(x_t, t) - \nabla_x \log q(x_t) \|^2]$
希望 $\min_{\theta} E_{x_t} [\| S_0(x_t, t) - \nabla_x \log q(x_t) \|^2]$
但有 $\nabla_x \log q(x_t) = E_{x_0 \sim q(x_0|x_t)} [\nabla_{x_0} \log q(x_t|x_0)]$

1.2 classifier guidance
 $q(x_0|x_t) = \log q(x_0) + \log q(x_t|x_0) - \log q(x_t)$
 $\nabla_{x_0} \log q(x_0|x_t) = \nabla_{x_0} \log q(x_0) + \nabla_{x_0} \log q(x_t|x_0)$
 $\nabla_{x_0} \log q(x_0|x_t) = \frac{q(x_t|x_0) q(x_0|x_t)}{q(x_0|x_t)^2} \log \text{trick}$
固定 x_t 取条件期望 $E_{x_0 \sim q(x_0|x_t)} [\nabla_{x_0} \log q(x_0|x_t)] = \nabla_{x_t} [q(x_0|x_t) dx_0] = 0$
 $S_0(x_t, t) \approx \nabla_{x_t} \log q(x_t|x_0)$ 或 $\nabla_{x_t} \log q(x_t|x_0) = E_{x_0} [\nabla_{x_t} \log q(x_t|x_0)]$
代入 $x_t = y_t x_0 + b_t \epsilon_t$ $\nabla_{x_t} \log q(x_t|x_0) = -\nabla_{x_t} [\frac{1}{2} (x_t - y_t x_0)^2]$
 $\Rightarrow S_0(x_t, t) = -\frac{\epsilon_t}{\sigma_t^2} (x_t - y_t x_0) = -\frac{\epsilon_t}{\sigma_t^2} x_t$
本质仍是 noise, 差值 $\frac{\epsilon_t}{\sigma_t^2}$ $\Rightarrow \min_{\theta} E_{x_0, \epsilon_t} [\| \hat{z}_0(x_t, t) - z \|^2]$

4. Probability Flow ODE 反向生成 SDE 在分布意义等价于 TODE
 $dx_t = -\frac{1}{2} \beta(t) [x_t + \nabla_{x_t} \log q(x_t|x_0)] dt$ 解这个 ODE 等价于 $q(x_t|x_0)$
当初初始 $q_T(x_T) \approx N(x_T; 0, I)$. 可用 ODE solver 求解
5. Flow Matching
Flow Model $x_0 \sim p_{init}$, ODE: $dx_t = U_t(x_t) dt$ 到 $t=1$ return
Diffusion $x_0 \sim p_{init}$ SDE $dx_t = U_t^0(x_t) dt + \sigma_t dW_t$
5.1 Time dependent vector field
 $U: \mathbb{R}^d \times [0, 1] \rightarrow \mathbb{R}^d$ $(x_t) \mapsto U_t(x_t)$ 轨迹 $x_t = U_t(x_t) \Leftrightarrow \frac{dx_t}{dt} = U_t(x_t)$
瞬时速度场, 输入 x 和 t , 输出一个运动向量和速度. 沿场把分布过去
给定 point $z \in \mathbb{R}^d$, 定义 cond prob path $p_t(\cdot|z)$ 满足从 noise - 跑到 目标 z
 $p_0(\cdot|z) = p_{init}$ $p_1(\cdot|z) = \delta_z$ 集中在 z 的 Dirac delta distribution
若用 Gaussian $p_t(\cdot|z) = N(\mu_t, \Sigma_t)$ 控制 μ_t 从 0 到 z , 方差从 1 到 0
边界 $\mu_0 = 0, \mu_1 = z \Rightarrow p_0(\cdot|z) = N(0, I)$ $p_1(\cdot|z) = N(z, 0) = \delta_z$
5.2 conditional probability path
5.3 conditional vector field 对任意 cond prob path 都有在等价 vector field
s.t. $x_0 \sim p_{init}$ $\frac{dx_t}{dt} = U_t(x_t|z) \Rightarrow x_t \sim p_t(\cdot|z) = N(\mu_t, \Sigma_t)$
 $U_t(x|z) = (\dot{\mu}_t - \frac{\beta_t}{\sigma_t^2} \dot{\Sigma}_t \mu_t) + \frac{\beta_t}{\sigma_t^2} \dot{\Sigma}_t x$ $\dot{\mu}_t = \frac{d\mu_t}{dt}$ $\dot{\Sigma}_t = \frac{d\Sigma_t}{dt}$
推导: 若有 $x_t = \mu_t + \beta_t \epsilon_t, \epsilon_t \sim N(0, I)$. 则 $\frac{dx_t}{dt} = \frac{d\mu_t}{dt} + \beta_t \frac{d\epsilon_t}{dt}$
整理 $U_t(x|z) = \frac{dx_t}{dt} = (\dot{\mu}_t - \frac{\beta_t}{\sigma_t^2} \dot{\Sigma}_t \mu_t) + \frac{\beta_t}{\sigma_t^2} \dot{\Sigma}_t x$

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 $p_t(x) = \int p_t(x|z) p_{data}(z) dz$
对应的 marginal vector field $U_t(x) = \int U_t(x|z) p_{data}(z) dz$
 $= E_{z \sim p_{data}} [U_t(x|z)]$ 在位置 x 上, 有多目标, 数据能解释
中间态, 但 p_{data} 未知, $U_t(x)$ 无法显式写出 $\Rightarrow \hat{U}_t(x)$ 来
5.5 FM Loss $L_{CFM} = E_{x \sim p_{data}, x' \sim p_t(\cdot|z)} [\| U_t(x) - U_t(x') \|^2]$ 直接对 q 做 value iteration
最小化 LCFM $\Leftrightarrow$ 最小化 LFM. 训练用条件, 期望学的 $U_t(x)$
有 $U_t(x|z) = (\dot{\mu}_t - \frac{\beta_t}{\sigma_t^2} \dot{\Sigma}_t \mu_t) + \frac{\beta_t}{\sigma_t^2} \dot{\Sigma}_t x$ 代入 $x_t = \mu_t + \beta_t \epsilon_t$
 $= \dot{\mu}_t + \beta_t \dot{\epsilon}_t$ $L_{CFM} = E_{t, z, \epsilon} [\| \dot{\mu}_t(\epsilon_t + \beta_t \epsilon_t) - \dot{\mu}_t(\epsilon_t + \beta_t \epsilon_t) \|^2]$
5.6 最简单的 optimal transport path 令 $x_t = t$ $\beta_t = 1 - t$ $\dot{\beta}_t = -1$
 $p_t(x|z) = N(tz, (1-t)^2 I) \Rightarrow x = tz + (1-t) \epsilon_t, \epsilon_t \sim N(0, I)$
 $U_t(x|z) = \dot{x} = z - \epsilon_t$ $L_{CFM} = E [\| U_t(x|z) - U_t(x|z') \|^2]$
Training ① 从数据 $z \sim p_{data}$ ② $t \sim U(0, 1)$ ③ $\epsilon_t \sim N(0, I)$
④ $x = tz + (1-t) \epsilon_t$ ⑤ 用 network 预测 $\hat{U}_t(x)$ ⑥ MSE 监督 $z - \hat{z}$

6. DDIM: 少步采样, 允许在确定性采样 $x_{t-1} = \sqrt{1 - \alpha_{t-1}} \hat{x}_0 + \sqrt{1 - \alpha_{t-1} - \beta_t^2} \hat{\epsilon}_t + \beta_t \epsilon_t$
 $\hat{x}_0 = \frac{x_t - \sqrt{1 - \alpha_t} \epsilon_t}{\sqrt{1 - \alpha_t}}$ $\hat{\epsilon}_t = \frac{x_t - \sqrt{1 - \alpha_t} \hat{x}_0}{\sqrt{1 - \alpha_t}}$ 从 $x_t = \sqrt{1 - \alpha_t} \hat{x}_0 + \sqrt{1 - \alpha_t} \epsilon_t$ 来
7. CG, CFG
7.1 Vanilla. 给定条件 y , 直接训 $u_\theta^0(x|y)$ ① 选 prompt y ② $x_0 \sim p_{init}$
③ 模拟 $dx_t = U_t^0(x_t) dt$ 从 $t=0$ 到 1. 不够多样, 不对齐

1.2 classifier guidance
 $q(x_0|x_t) = \log q(x_0) + \log q(x_t|x_0) - \log q(x_t)$
 $\nabla_{x_0} \log q(x_0|x_t) = \nabla_{x_0} \log q(x_0) + \nabla_{x_0} \log q(x_t|x_0)$
 $\nabla_{x_0} \log q(x_0|x_t) = \frac{q(x_t|x_0) q(x_0|x_t)}{q(x_0|x_t)^2} \log \text{trick}$
固定 x_t 取条件期望 $E_{x_0 \sim q(x_0|x_t)} [\nabla_{x_0} \log q(x_0|x_t)] = \nabla_{x_t} [q(x_0|x_t) dx_0] = 0$
 $S_0(x_t, t) \approx \nabla_{x_t} \log q(x_t|x_0)$ 或 $\nabla_{x_t} \log q(x_t|x_0) = E_{x_0} [\nabla_{x_t} \log q(x_t|x_0)]$
代入 $x_t = y_t x_0 + b_t \epsilon_t$ $\nabla_{x_t} \log q(x_t|x_0) = -\nabla_{x_t} [\frac{1}{2} (x_t - y_t x_0)^2]$
 $\Rightarrow S_0(x_t, t) = -\frac{\epsilon_t}{\sigma_t^2} (x_t - y_t x_0) = -\frac{\epsilon_t}{\sigma_t^2} x_t$
本质仍是 noise, 差值 $\frac{\epsilon_t}{\sigma_t^2}$ $\Rightarrow \min_{\theta} E_{x_0, \epsilon_t} [\| \hat{z}_0(x_t, t) - z \|^2]$

1. $\{(s, a, r, s')\}$ r : 即时 reward s' : 下一状态
2. MDP 含 (S, A, P, r) S 有限状态空间, A 动作空间 $P(s'|s, a)$
 $r(s, a) = \sum_{s'} r(s, a, s') P(s'|s, a)$ expected reward
2.1 Markov property $P(s_{t+1}|s_t, a_t, \dots) = P(s_{t+1}|s_t, a_t)$
2.2 在 Π_T , agent 交互形成 trajectory $s_0, a_0, \dots$
discounted total reward $R^\pi(s_0) = r_0 + \gamma r_1 + \dots$
 $R^\pi(s_0) = \sum_{t=0}^{\infty} \gamma^t r_t$ $\gamma \in [0, 1]$

3. Value function and Optimal Policy
return $V^\pi(s) = E[R^\pi(s)]$
最优 π^* 满足 $V^{\pi^*}(s) \geq V^\pi(s), \forall s, \forall \pi$
optimal state value function $V^*(s) = V^{\pi^*}(s) \Rightarrow \pi^*$
4. Planning: 知道 $P(s'|s, a), r(s, a)$ 可直接求 $P(s'|s, a), r(s, a)$
RL 不环境. 收集 $\{(s, a, r, s')\} \Rightarrow \pi^*$
5. Value Iteration
planning 两步 $P(s'|s, a), r(s, a) \Rightarrow V^*$ $V^* \Rightarrow \pi^*$ Bellman
value iter ① 初始化 $V_0(s)$
② 重复更新 $V_{k+1}(s) = \max_a \{ r(s, a) + \gamma \sum_{s'} P(s'|s, a) V_k(s') \}$
直到 $\max_s |V_{k+1}(s) - V_k(s)| \leq \epsilon$ 选最好动作

5.1 Bellman operator 是 contraction $|\max_a - \max_b| \leq \max |a - b|$
 $\max_s |V_{k+1}(s) - V_k(s)| \leq \gamma \max_s |V_k(s) - V_{k-1}(s)|$
当 $\gamma < 1$ 时误差 $\downarrow$ 最终收敛到 V^* optimality equation
当 $V_k = V^*$ 时, 再迭代不变 $V^*(s) = \max_a \{ r(s, a) + \gamma \sum_{s'} P(s'|s, a) V^*(s') \}$
5.2 Optimal Q-function $Q^*(s, a) = r(s, a) + \gamma \sum_{s'} P(s'|s, a) V^*(s')$
由 Q^* 得到策略 $\pi^*(s) = \arg \max_a Q^*(s, a)$
 $V^*(s) = \arg \max_a Q^*(s, a)$

6. Q-learning: Value iteration 需知 $P(s'|s, a), r(s, a)$
在 value iteration 上通过多次采样, 近似为
 $Q_{k+1}(s, a) \approx r(s, a) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q_k(s', a')$
收敛到 $Q^*(s, a) = r(s, a) + \gamma \sum_{s'} P(s'|s, a) \max_{a'} Q^*(s', a')$ 但现实未知
6. Q-learning: Value iteration 需知 $P(s'|s, a), r(s, a)$
若只有一样本则 $YTD = r + \gamma \max_{a'} Q(s, a')$ unbiased
用 $Q(s, a) \leftarrow (1 - \alpha) Q(s, a) + \alpha [r + \gamma \max_{a'} Q(s, a')]$ 方差大
定义 $\delta = r + \gamma \max_{a'} Q(s, a') - Q(s, a)$ TD error
TD target $Q(s, a) \leftarrow Q(s, a) + \alpha \delta$ 提高估计

7. CG, CFG
7.1 Vanilla. 给定条件 y , 直接训 $u_\theta^0(x|y)$ ① 选 prompt y ② $x_0 \sim p_{init}$
③ 模拟 $dx_t = U_t^0(x_t) dt$ 从 $t=0$ 到 1. 不够多样, 不对齐

sigmoid + binary $\hat{y} = \sigma(z)$ z 是一个数

$$L = -y \log \hat{y} - (1-y) \log(1-\hat{y}) \quad \frac{\partial L}{\partial z} = \hat{y} - y$$

$$\frac{\partial L}{\partial z} = \frac{\partial L}{\partial \hat{y}} \frac{\partial \hat{y}}{\partial z} = -\frac{y}{\hat{y}} + \frac{1-y}{1-\hat{y}} \quad \frac{\partial \hat{y}}{\partial z} = \hat{y}(1-\hat{y})$$

$\Rightarrow \hat{y} - y$

$$\text{softmax} + \text{CE} \quad \hat{y}_i = \frac{e^{z_i}}{\sum_j e^{z_j}} \quad \frac{\partial L}{\partial z_i} = \hat{y}_i - y_i$$

$\frac{\partial L}{\partial W^{(2)}} = \delta^{(2)T} \cdot \frac{\partial L}{\partial W^{(2)}}$
 四个值相加于 $\delta^{(1)}$ 和相加
 $CNN \quad X = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \quad k = \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix} \quad S_{ij} = \sum_{a,b} k_{ab} x_{i+a, j+b}$
 $A = \text{ReLU}(S) \quad z = [1 \ 1 \ 2] \text{vec}(A)$
 $L = \frac{1}{2} \|z - y\|^2, y = 2 \quad S = \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix} \quad A = \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix}$
 $\text{vec}(A) = \begin{bmatrix} 1 \\ 0 \\ 0 \\ 2 \end{bmatrix} \quad z = \begin{bmatrix} 1 \\ 1 \\ 2 \end{bmatrix}$

$$\frac{\partial L}{\partial W} = \delta^{(2)} \text{vec}(A)^T \quad \& \quad \frac{\partial L}{\partial \text{vec}(A)} = W^T \delta^{(2)}$$
$$\begin{aligned} \delta^{(2)} &= \frac{\partial L}{\partial z} = z - y \quad \frac{\partial L}{\partial V} = \delta^{(2)} h_2 \quad \frac{\partial L}{\partial h_2} = W \delta \\ \delta^{(2)} &= \frac{\partial L}{\partial a_2} = \frac{\partial L}{\partial h_2} (1 - h_2^2) \quad \frac{\partial L}{\partial h_1} = W \delta_2 \\ \delta^{(1)} &= \frac{\partial L}{\partial a_1} = \frac{\partial L}{\partial h_1} (1 - h_1^2) \quad \frac{\partial L}{\partial U} = \sum_i \delta^t x^t \\ \frac{\partial L}{\partial W} &= \sum_t \delta^t h^t + 1 \end{aligned}$$
$$D_{KL}(p||q) = \log \frac{b_1}{b_2} + \frac{b_1 + (\mu_1 - \mu_2)}{2b_2} - \frac{1}{2} \text{tr} \left(\frac{\Sigma_1}{\Sigma_2} \right) + \frac{1}{2} \log \frac{|\Sigma_2|}{|\Sigma_1|} + \frac{1}{2} (\mu_2 - \mu_1)^T \Sigma_2^{-1} (\mu_2 - \mu_1)$$

G 在 I 上卷积