

NLP

Lecture 1

- $P(B|A) = \frac{P(A \cap B)}{P(A)} = \frac{\text{likelihood } P(A|B)P(B)}{P(A)} \rightarrow \text{prior}$
- Bernoulli $\rightarrow$ 两分类 $\rightarrow$ Binomial: $\rightarrow$ marginal $\rightarrow$ 一堆 i.i.d 的 Bernoulli variable $\Rightarrow$ 1 个 binomial variable
categorical distribution $\rightarrow$ 多分类 $\rightarrow$ Multi-nomial

Lecture 2

1. n-gram language models

1.1) 句根无率 $P(w_1, \dots, w_n)$

1.2) 怎么拆 $\left\{ \begin{array}{l} \text{Uni-gram: 词相互独立 } P(w_1) \dots P(w_n) \\ \text{Bi-gram: 只和前-词有关 } P(w_1) P(w_2|w_1) \dots P(w_n|w_{n-1}) \end{array} \right.$

1.3) 模型 $\theta = \{P(w) : w \in V\}$ 在 unigram 中. 由 MLE 可得 $P(w_j) = \frac{c(w_j)}{\sum_i c(w_i)}$

Bi-gram $P(w|v) = \frac{\#(vw)}{\sum_{w'} \#(vw')}$

Tri-gram $P(w|uv) = \frac{\#(uvw)}{\sum_{w'} \#(uvw')}$

1.4) Smoothing 解决 missing words

$P_{\text{Laplacian}}(w_j) = \frac{\#w_j + 1}{\sum_i \#w_i + V}$

bigram: $P_{\text{Laplacian}}(w_2|w_1) = \frac{\#(w_1w_2) + 1}{\#(w_1) + V}$

$\leftarrow$ 词表大小

2. 词表示方式 (有 symbolic 与 semantic 之争)

2.1 one-hot 无法用内积 measure similarity

2.2 distributional word-vector 多维. 每维表示一定特征

2.3 Word2Vec 在一个滑动 window 中. 分 center word 和 context word

$P(\text{context} | \text{center}) > P(\text{在窗外} | \text{center})$

一个词 w 对应两个 vector $\left\{ \begin{array}{l} v_w: w \text{ 为 central} \\ u_w: w \text{ 为 context} \end{array} \right.$

$P(w_o | w_c; \theta) = \frac{\exp(u_o^T v_c)}{\sum_{w=1}^V \exp(u_w^T v_c)}$ v_c 是一样的 context word

优化: 令 $L = -\log P \Rightarrow = -u_o^T v_c + \log \sum_{w=1}^V \exp(u_w^T v_c)$ $\frac{\partial L}{\partial v_c} = -u_o + \sum_{w=1}^V \frac{\exp(u_w^T v_c)}{\sum_{w=1}^V \exp(u_w^T v_c)} u_w$

$v_c^{\text{new}} = v_c^{\text{old}} - \eta \frac{\partial L}{\partial v_c}$

当预测完全正确. $\frac{\partial L}{\partial v_c} \rightarrow 0$ (因为 $\sum_w P(w|w_c) u_w = u_o$)

$\Rightarrow = -u_o + \sum_{w=1}^V [6(u_w^T v_c)] u_w$
 $\uparrow$ 正样本分量 $\uparrow$ 远离错误的平均方向

小技巧. 用 negative sampling: 不和所有词比. 随机采 K 个 negative pairs

目标: 正对 $\uparrow$ 负对 $\downarrow$ $O(V) \rightarrow O(K)$ $K \ll V$

2.4 Word2Vec 两个改进

- 2.4.1 glove 利用 word2vec 没利用到的全局信息 $\rightarrow$ 此值能过滤无效信息
用共现 ratio 概率此值比概率更重要 $P(k|x_1) \quad P(k|x_2)$ 用 log 是考虑词频
用 $v_k^T(v_i - v_j)$ 模拟 context word k 偏向 i 还是 j $\log \frac{P(k|i)}{P(k|j)}$ $w_i^T \tilde{w}_j + b_i + b_j \approx \log x_{ij}$
2.4.2 fasttext 把变形词当独立不好 词 = character n -grams 组合 $S(w_c, w_o) = u_o^T (\sum_{g \in w_c} z_g)$
 $= \sum_g u_o^T z_g$

Lecture 3

NLP Lec 3 POS TAG

给定 $O = (o_1, o_2, \dots, o_T)$ 词 $o_t \in V$

预测对应 $Q = (q_1, q_2, \dots, q_T)$ $q_t \in S \subseteq \text{Tag集合}$

核心任务 $Q^* = \arg \max_Q \Pr(Q|O) = \arg \max_Q \Pr(O|Q)P(Q)$ 分母 $P(O)$ 对所有 Q 一样

1. 为什么重要: ① 语法结构 ② 语义信息 (词性不同意思不同)

2. 为什么难: ① Lexical ambiguity 一个词 $\rightarrow$ 多个 Tag 常见词易模糊

3. 一个简单 BS $q_t = \arg \max_s \Pr(q_t|o_t)$ 只看当前词, 不看上下文

A. 只看 emission prob $\Pr(o_t|q_t)$ 因为 O 知道, 其实就是 which tag suits O most

B. 只看 $P(Q)$ 忽略 O 的信息, 通用不相关序列

生成 tag $\Pr(Q) = \Pr(q_1 \dots q_T) = \Pr(q_1) \Pr(q_2|q_1) \Pr(q_3|q_1, q_2) \dots$

emit $\Pr(O|Q) = \Pr(o_1|q_1) \Pr(o_2|q_1, q_2) \dots$

$q_2|q_1$ 都 $|S|^2$

$q_3|q_1, q_2$ 更是 $|S|^3$

prior

Markov assumption ① $\Pr(Q) = \Pr(q_1) \prod_{t=2}^T \Pr(q_t|q_{t-1})$ $\Pr(q_1) = \pi_{q_1}$

定义 Transitional matrix $A_{ij} = \Pr(q_t=j|q_{t-1}=i)$

$\sum_j A_{ij} = 1$

从 POS tag $i \rightarrow j$ 的 prob

likelihood $|S| \times |S|$

$\sum_{j \in S} A_{ij} = 1$

一般不对称

② $\Pr(O|Q) = \prod_{t=1}^T \Pr(o_t|q_t)$ 当前 word 只由当前 tag 生成

emission matrix

$B_{io} = \Pr(o_t=o|q_t=i)$

posterior

$|S| \times |V|$

从 POS tag i 到生成 word o 根元率 $\sum_{o \in V} B_{io} = 1$

$Q^* = \arg \max_Q [\pi_{q_1} \prod_{t=2}^T A_{q_{t-1}, q_t} \prod_{t=1}^T B_{q_t, o_t}]$

Q

出现 prob

NLP Lec 4 Inference 句子出现的 prob

inference

任务: $O = (o_1, o_2, \dots, o_T)$

以及固定的 HMM 参数 $\lambda = (A, B, \pi)$

求 $\Pr(O|\lambda)$ 这个 HMM 生成 O 的 prob

1. 难. $\Pr(O) = \sum_Q \Pr(O, Q)$

hidden tag sequence $|S|$ 有 $|S|^T$ 种组合

2. 为什么可以 DP

重复利用部分最优

3. Forward Algo $\max(ax, ay, bx, by) = \max(a, b) \times \max(x, y)$

Define forward prob $\alpha_t(j) = P(o_1, \dots, o_t, q_t = j)$ [总结了所有可能 path 求 $q_t = j$ 的 prob]

$$\alpha_1(i) = P(o_1, q_1 = i) = \pi_i \times b_i(o_1)$$

$$P(A, B|C) = P(A|C)P(B|A, C)$$

$$N = |S|$$

$$\Pr(q_t = i) \Pr(o_1 | q_t = i)$$

这里 o_t 只和 q_t 有关

$$P(o_t | o_{1:t-1}, q_{t-1} = k, q_t = j) = P(o_t | q_t = j)$$

$$\alpha_t(j) = \sum_{k=1}^N \alpha_{t-1}(k) a_{kj} b_j(o_t)$$

$$\text{Proof: } P(o_{1:t}, q_t = j) = \sum_{k=1}^N P(o_{1:t-1}, q_{t-1} = k, q_t = j)$$

$$= \sum_{k=1}^N P(o_{1:t-1}, q_{t-1} = k) P(q_t = j, o_t | o_{1:t-1}, q_{t-1} = k)$$

$$= \sum_{k=1}^N \alpha_{t-1}(k) P(q_t = j | q_{t-1} = k) \times P(o_t | q_t = j)$$

$$= \sum_{k=1}^N \alpha_{t-1}(k) a_{kj} b_j(o_t)$$

总

↑
T步

算这个 $O(N^2)$

TN^2

$$\Pr(O) = \sum_j P(O, q_T = j) = \sum_j \alpha_T(j)$$

$$N^T \rightarrow TN^2$$

↑ 拆成上一步

4. Backward Algo

N 种可能

转移到下一状态 j 生成预测

$\beta_t(i) = \Pr(o_{t+1}, \dots, o_T | q_t = i)$ 当前 state 为 i , 未来剩余 observation 的 prob

规定

$\beta_T(i) = 1$ 无法生成更多 seq

$$\beta_{t-1}(i) = \sum_{j=1}^N a_{ij} b_j(o_T) \beta_T(j) = P(o_T | q_{t-1} = i)$$

$$\beta_t(i) = \sum_{j=1}^N a_{ij} b_j(o_{t+1}) \times \beta_{t+1}(j)$$

$$P(o_{t+1:T} | q_t = i) = \sum_{j=1}^N P(o_{t+1:T}, q_{t+1} = j | q_t = i) = \sum_{j=1}^N P(q_{t+1} = j | q_t = i) P(o_{t+1:T} | q_{t+1} = j, q_t = i)$$

初始化 $\alpha_1(i)$ $\forall i$ of q_1

for $t = 2, \dots, T$

for $j = 1, \dots, N$

$$\alpha_t(j) = \sum_k \alpha_{t-1}(k) a_{kj} b_j(o_t)$$

$$\Pr(O) = \sum_j \alpha_T(j)$$

$$= \sum_{j=1}^N a_{ij} P(o_{t+1}, o_{t+2:T} | q_{t+1} = j)$$

$$= \sum_{j=1}^N a_{ij} P(o_{t+1} | q_{t+1} = j) P(o_{t+2:T} | o_{t+1}, q_{t+1} = j)$$

$$= \sum_{j=1}^N a_{ij} b_j(o_{t+1}) \beta_{t+1}(j) \quad \text{50th 无关}$$

$$Pr(O) = \sum_{j=1}^N \pi_j b_j(o_1) \beta_1(j)$$

$\uparrow$ $\uparrow$ $\uparrow$
 $Pr(q_1=j)$ $Pr(o_1|q_1=j)$ $P(o_2:T|q_1=j)$

初始化 $\beta_T(i)$ $\forall i$ for q_T

for $t=T-1, \dots, 1$

for $j=1, \dots, N$

$$\beta_t(j) = \sum_k \beta_{t+1}(k) a_{jk} b_j(o_t)$$

$$Pr(O) = \sum_j P(q_1=j) b_j(o_1) \beta_1(j)$$

Lecture 5 Decoding 参数学习 π, A, B

prediction

找最佳 tag

1. Structured Prediction

1.1 独立逐词 fails. 当前依赖邻居. 邻居依赖当前

1.2 Joint prediction

但若穷举 $O(N^T)$

$$Q^* = \arg\max_Q Pr(Q|O; \theta) \quad Q = [q_1, \dots, q_T], q_t \in S$$

与 Viterbi 压缩至 $O(TN^2)$

2. Viterbi Algo 用了 DP

$$V_t(j) = \max_{q_1, \dots, q_{t-1}} Pr(q_1, \dots, q_{t-1}, q_t=j, o_1, \dots, o_t)$$

$$v_1(j) = \max_{\phi} Pr(q_1=j, o_1) = Pr(q_1=j) Pr(o_1|q_1=j)$$

$$v_2(j) = \max_{q_1} Pr(q_1, q_2=j, o_1, o_2) = \max_{q_1} Pr(q_1) Pr(o_1|q_1) Pr(q_2=j|q_1) Pr(o_2|q_2=j)$$

$$= \max_{q_1} v_1(q_1) Pr(q_2=j|q_1) Pr(o_2|q_2=j)$$

$\downarrow$

$$v_t(j) = \max_{q_1, \dots, q_{t-1}} (q_1, \dots, q_{t-1}, q_t=j, o_1, \dots, o_t) = \max_{q_{t-1}} \max_{q_{1:t-2}} Pr(q_t=j|q_{t-1}) Pr(o_t|q_t=j) Pr(q_{1:t-1}, o_{1:t-1})$$

$$(q_{1:t-2}, q_{t-1}=i, o_{1:t-1}) \text{ 交换 max } = \max_{q_{t-1}} Pr(q_t=j|q_{t-1}) Pr(o_t|q_t=j) \max_{q_{1:t-2}} Pr(q_{1:t-1}, o_{1:t-1})$$

$$q_t=j \mid q_{t-1}=i$$

$$o_t \mid q_t=j$$

$$= \max_{q_{t-1}} Pr(q_t=j|q_{t-1}) Pr(o_t|q_t=j) v_{t-1}(q_{t-1})$$

$$= \max_i v_{t-1}(i) a_{ij} b_j(o_t)$$

必然来自使 $v_{t-1}(i) a_{ij}$ 最大的 path

保存最优路

$$\text{令 } p_t(j) = \arg\max_k v_{t-1}(k) a_{kj} b_j(o_t)$$

若当前最优 path 与 j 结尾, 则一定某上一状态 $p_t(j)$

$$\text{从 } q_T^* = \max_k v_T(k) \text{ backtrack } q_{t+1}^* = p_t(q_t^*)$$

算法 ① 初始化 $v_1(i)$ $\forall i \in q_1$

② 从 $t=2, \dots, T$

$$\text{for } j=1, \dots, N \quad v_t(j) = \max_k (v_{t-1}(k) a_{kj} b_j(o_t))$$

$$p_t(j) = \arg\max_k v_{t-1}(k) a_{kj} b_j(o_t)$$

$O(TN^2)$

$O(T)$

③ Backtrack to find Q^*

$$\text{④ 返回 } Pr(Q^*|O) = \max_j v_T(j)$$

3. Expectation Maximization Algo 用了 DP 和 MLE ^{prob}

E-step 在 unlabeled data 上算 estimated prob of labels $\Rightarrow$ pseudo labels

labels: soft labels

0λ

$\lambda = [A \ B \ \pi]$

M-step: estimate model paras using pseudo labels as in MLE

Estimating the soft labels $i \rightarrow j$

$$E: \xi_t(i, j) = \Pr(q_t = i, q_{t+1} = j | O, \lambda) = \Pr(q_t = i, q_{t+1} = j, O | \lambda) / \Pr(O | \lambda)$$

$$\alpha_t(i) = \Pr(O_{1:t}, q_t = i)$$

$$= \Pr(\underline{O_{1:t}}, \underline{O_{t+1}}, \underline{O_{t+2:T}}, q_t = i, q_{t+1} = j | \lambda) / \dots$$

$$\beta_{t+1}(j) = \Pr(O_{t+2:T} | q_{t+1} = j)$$

$$= \Pr(\underline{O_{1:t}}, q_t = i | \lambda) \Pr(\underline{O_{t+2:T}} | \dots)$$

对 B $\gamma_t(i) = P(q_t = i | O, \lambda)$ i 和 0 同时出现 λ 和这一项有关

$$= \Pr(q_t = i, O | \lambda) / \Pr(O | \lambda)$$

$$= \alpha_t(i) \beta_{t+1}(i) / \Pr(O | \lambda)$$

还要对 0 求和

Supervised MLE (隐状态序列已知)

$$a_{ij} = \frac{C(i \rightarrow j)}{C(i)}$$

$$b_{i,0} = \frac{C(i \rightarrow 0)}{C(i)}$$

$$\pi_i = \frac{C(q_i = i)}{m}$$

$\leftarrow$ # sentences in corpus

EM 中 Given m sentences 一些有 label

初始化 $\lambda = [A, B, \pi]$ random / POS-TAG if avail

loop until convergence:

E step: run forward, backward. 得到 $d_t(j)$ $\beta_t(i)$

M step: find $\xi_t(i, j)$

$$a_{ij} = \frac{\sum_{k=1}^m \sum_{t=1}^T \xi_t(i, j)}{\sum_{k=1}^m \sum_{t=1}^T \xi_t(i, j)}$$

所有 sentence 应该是 $\sum_j a_{ij} = 1$

$$b_{i,0} = \frac{\sum_{k=1}^m \sum_{t=1}^T \gamma_t(i)}{\sum_{k=1}^m \sum_{t=1}^T \gamma_t(i)}$$

$i \rightarrow 0$ 单个

$i \rightarrow 0$ 所有

每一轮 EM 循环保证 $\log \Pr(O | \lambda)$ 单调不减

但仅收敛至 Local Minimum 不良初始化易 fail

引入少量 supervised data 和多次随机重启效果好

Lec 6 Grammar and Syntax

1.1 n-grams & HMM

long-range dependency x

只有 shallow syntax

1.2 constitute: 一堆能组成 unit 的词. unit 可出现 anywhere 且保持词意. 保持时不能 break down
context: 句中除 constitute 的部分

↓
context-free grammar CFG

- 2.1 4 elements
- ① N : set of non-terminals 如 $\{S, NP, \dots\}$
 - ② Σ : set of terminals 实际词汇表
 - ③ R : set of rules 形式为 $A \rightarrow \beta$, $A \in N, \beta \in (N \cup \Sigma)^*$
 - ④ S : 起始符号且 $S \in N$

2.2 Derivations: ① 单步 $\alpha A \gamma \Rightarrow \alpha \beta \gamma$ 若 $A \rightarrow \beta \in R$
② m 步 $\alpha_1 \Rightarrow \alpha_2 \dots \alpha_{m-1} \Rightarrow \alpha_m$ m 次替换.
 $A \rightarrow \beta$ 的应用仅取决于 non-terminal A 本身

$L_G = \{w | w \text{ is in } \Sigma^* \text{ and } S \xRightarrow{*} w\}$ 只包含 terminal 且从 S 开始
通过 recursive 完美建模 long-range dependency
支持多种句式

问题: ① Ambiguity 一句对应多 parsing tree ② 搜索效率低下

3. Top down vs Bottom up

Top down: 从 S 开始, 用 CFG 展开 non-terminal, 匹配输入. 无法匹配剪枝 / 回溯

Bottom up: 从输入词 (leaves) 开始, 用 CFG 生成 non-terminal 直到汇到 S

Why it fails: ① 重复子问题 ② late pruning

$n^3 \leftarrow \text{rule size}$

4. CYK Algo (DP) $O(n^3)$ 每个子空间 $[i, j]$ 只解析一次, 若 1 个, 需为 term

4.1 CFG $\rightarrow$ CNF 严格限制 - 一条 rule 右侧 $A \rightarrow BC$ 或 $A \rightarrow a$

转换算法

CFG	CNF	至多两个, 且两个得为 non-term
$A \rightarrow Bc$	$A \rightarrow B \quad B \rightarrow c$	

$A \rightarrow B$	将所有 $B \rightarrow \gamma$ 规则赋给 A 变成 $A \rightarrow \gamma$
$A \rightarrow BCD$	$A \rightarrow XD \quad X \rightarrow BC$

多元产生形式 $\Rightarrow$ 严格二叉树形式

4.2 CYK Parsing $(n+1) \times (n+1)$ matrix n words

~~word~~ w_0, w_1, w_2, w_3, w_4

(i, j) 包含所有可能的 non-terminals that can generate the constituent ranging i to j
 $\forall (i, j)$ 有 $j > i+1$ 可以被拆成 $(i, j) = (i, k) + (k, j)$ $(1, 3) = (1, 2) + (2, 3)$

对每个词 w_j $T[j-1, j] = \{A \in N \mid A \rightarrow w_j \in R\}$

递推

$$T[i, j] = \bigcup_{k=i+1}^{j-1} \{A \in N \mid \exists B \in T[i, k], C \in T[k, j] \text{ s.t. } A \rightarrow BC \in R\}$$

终止条件 $S \in T[0, n]$ 则句子合法

for $j \leftarrow$ 从 1 到 句长: n

$\forall \{A \mid A \rightarrow \text{words}[j] \in \text{grammar}\}$

$\text{table}[j-1, j] \leftarrow \text{table}[j-1, j] \cup A$

init

for $i \leftarrow j-2$ 到 0

for $k \leftarrow i+1$ to $j-1$ do

for all $\{A \mid A \rightarrow BC \in \text{grammar and } B \in T[i, k] \text{ } C \in T[k, j]\}$

$\text{table}[i, j] \leftarrow \text{table}[i, j] \cup A$

n^2 所有子串长度

从哪开始有多长

对每个子串还要拆

R

$A \rightarrow BC$ $O(R)$ 检查

Lecture 7 Probabilistic CFG

CFG (CYK) 判断一个句子有哪些合法 parse trees

⇒ 给 tree prob 并找出最优树

但多个, 哪个更有可能

CFG (CYK) 做不到

1. PCFG

1.1 给 production rule + prob $P(A \rightarrow \gamma) \sum_{\gamma} P(A \rightarrow \gamma) = 1$

1.2 Parse Tree prob: $P(t|G) = \prod_{r \in t} P(r)$
 ↑ ↑ ret
 - 一颗 tree PCFG 语法 → r 是 tree 里所有 rule

1.3 Sentence prob: 一个句子 $W_{1:T}$ 不是某 tree 的 prob, 而是所有能生成这个 sentence 的 parsing tree 的总和

$$P(W_{1:m}|G) = \sum_{t \in T(W_1, \dots, W_m)} P(t|G)$$

↑ ↑
输入句子 树

类似 $P(O) = \sum_Q P(O, Q)$
 ↑
 hidden sequence

不能暴力枚举, 可能指数级增 ⇒ DP!
 ↓ tractable

1.4 三个关键假设

- ① Place invariant $P(N_{k,l}^j \Rightarrow \xi)$ 对不同的 k 一样
 同一个 non-term 生成树的 prob 不依赖于句中位置.
- ② Context free $P(N_{k,l}^j \Rightarrow \xi \mid \text{anything outside } [k,l]) = P(N_{k,l}^j \Rightarrow \xi)$
 若当前子树 cover span $[k,l]$, 那么 span 外部上下不影响内部展开
- ③ Ancestor-free 某个节点如何展开不依赖祖先
 $P(N_{k,l}^j \Rightarrow \xi \mid \text{ancestor nodes outside } N_{k,l}^j) = P(N_{k,l}^j \Rightarrow \xi)$
 当前 NP 是从 S/VP 下来, 规则 prob 一样. 很强假设, 也是局限

1.5 Inference and Estimation algos with PCFG

- ① 推句子的 prob $Pr(W_1, \dots, W_m|G) = \sum_t Pr(t, W_{1:m}|G)$
- ② 找最优 parsing tree $t^* = \argmax_t Pr(t|W_{1:m}, G)$
- ③ Learn a PCFG from a training corpus using MLE $G^* = \argmax_G Pr(W_{1:m}|G)$
 假设 p, q 被 N_j^i 代表, 那么外部生成剩余词 prob

1.6 Inside/outside prob

- outside: $\alpha_j(p, q) = P(W_{1:p-1}, N_{pq}^j, W_{q+1:m}|G)$
- inside: $\beta_j(p, q) = P(W_{pq}|N_{pq}^j, G)$

N_{pq}^j 表示 non-terminal j derives word from p to q
 这棵子树内部生成这段 words 的 prob

1.7 Inside

1.7.1 Base case 若 span=1 $p=q=k$ W_k

$\beta_j(k, k) = P(W_k|N_j^i, G)$ N_j^i 展开 W_k 的 prob
 若有 rule $N_j^i \rightarrow W_k$ 则 $\beta_j(k, k) = P(N_j^i \rightarrow W_k)$
 否则为 0

1.7.2 $\beta_j(p, q) = \sum_{r,s} \sum_{d=p}^{q-1} P(N_j^i \rightarrow N_r^i N_s^i) \beta_r(p, d) \beta_s(d, q)$
 枚举规则 ↑ ↑
 更大 切分

$$\begin{aligned}
 \beta_j(p, q) &= \sum_{r,s} \sum_{d=q}^{q-1} P(W_{pd}, N_{pd}^r, W_{(d+1)q}, N_{(d+1)q}^s | \underline{N_{pq}^j}, G) \\
 &= \sum \sum P(\underbrace{N_{pd}^r N_{(d+1)q}^s}_{\text{位置无关}} | \underbrace{N_{pq}^j}_{\text{ancestor 和 DG 无关}}, G) P(W_{pd} | \underbrace{N_{pq}^j}_{\text{ancestor}}, \underbrace{N_{pd}^r, N_{(d+1)q}^s}_{\text{上下文}}, G) P(W_{(d+1)q} | \underbrace{N_{pq}^j}_{\text{context}}, \underbrace{N_{pd}^r, N_{(d+1)q}^s}_{\text{无关}}, G, W_{pd}) \\
 &= \sum \sum P(N^j \rightarrow N^r N^s) \beta_r(p, d) \beta_s(d+1, q)
 \end{aligned}$$

1.8 Outside $\alpha_j(p, q) = P(W_{1:p-1}, N_{pq}^j, W_{q+1:m} | G)$ 和 inside 有关系

Base case $P(W_{1:m} | G) = \sum_j \alpha_j(k, k) P(N^j \rightarrow W_k)$ $P(W_{1:m}, N_{pq}^j | G) = \alpha_j(p, q) \beta_j(p, q)$

若 $N^1 = S$ 覆盖句子 $[1, m]$ 外面什么没有, 那 $\alpha_1(1, m) = 1$ 其余 non-terminal $\alpha_j(1, m) = 0, j \neq 1$.

整句必须从 S 开始
自己在父左侧 更大的父亲 $f \rightarrow j, g$ 用右的里填充.

$$\begin{aligned}
 \alpha_j(p, q) &= \left[\sum_{f, g \neq j} \sum_{e=q+1}^m \alpha_f(p, e) P(N^f \rightarrow N^j N^g) \beta_g(q+1, e) \right] \\
 &\quad + \left[\sum_{f, g \neq j} \sum_{e=1}^{p-1} \alpha_f(e, q) P(N^f \rightarrow N^g N^j) \beta_g(e, p-1) \right]
 \end{aligned}$$

只是单层上一级 用左的 inside prob 填充

1.9 Optimal parsing tree 定义 $\delta_j(p, q)$ 表示生成 span $[p, q]$ 最高 prob parse tree. 在下

若 span 长度为 1 $\delta_j(k, k) = P(N^j \rightarrow W_k)$

prob $\Leftarrow \delta_j(p, q) = \max_{r,s,d} P(N^j \rightarrow N^r N^s) \delta_r(p, d) \delta_s(d+1, q)$ 和 inside 递推类似

一个值 记录 $\varphi_j(p, q) = \argmax_{r,s,d} P(N^j \rightarrow N^r N^s) \delta_r(p, d) \delta_s(d+1, q)$ $\sum \rightarrow \max$

具体方式 $\Leftarrow$ $\downarrow$ 左子树状态 右子树状态

$= (r, s, d)$ 表示左孩是 N^r 右孩 N^s split point 为 d .

然后从 $\delta_1(1, m)$ 开始回溯, 恢复整棵 parse tree

1. n-gram 3问题
- ① 数据稀疏
 - ② 复杂度随数据n指数增长
 - ③ 用 fixed window 太 rigid straightforward

神经网络. 参数更高效. 可用不同设计解决不同问题

2. Logistic Regression

$$a = \sigma(\sum W_i x_i + b), l(a, y) = -[y \log a + (1-y) \log (1-a)]$$

非线性压到 (0,1)

$x \in \mathbb{R}^m$
hidden size n.

3. RNN $h^{(t)} = f(h^{(t-1)}, x^{(t)}; \theta)$

$$\theta = \{U, W, V, b, c\}$$

biases

$h^{(t)} \rightarrow h^{(t+1)}$
 $x^{(t)} \rightarrow h^{(t+1)}$

$W \in \mathbb{R}^{n \times n}$
 $V \in \mathbb{R}^{m \times n}$
 $n \times m$

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

$$h^{(t)} = \tanh(a^{(t)})$$

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

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

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

$$= \tanh(b + W(\tanh(b + W h^{(t-2)} + U x^{(t-1)})) + U x^{(t)})$$

NLL Loss $-\sum_t \log P_{\text{model}}(y^{(t)} | \{x^{(1)}, \dots, x^{(t)}\})$

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

BPTT

gradient vanishing or explosion

Back Propagation through time

4. Why tanh? 引入非线性

$$y = \frac{e^x - e^{-x}}{e^x + e^{-x}}$$

BPTT中梯度含有 $\prod \frac{\partial h^{(k)}}{\partial h^{(k-1)}}$

$\rightarrow 0$ 历史遗忘, $|\lambda| > 1$ $W^t \rightarrow \infty$ 梯度爆炸.
将值域压到 $[-1, 1]$ 对矩阵指数增长起截断. 缓解优化敏感度

梯度衰减问题

tanh 导数 $\in (0, 1]$ 使连乘不易发散, 但长序列仍面临

令 $u = e^x - e^{-x}$ $v = e^x + e^{-x}$

$$(\frac{u}{v})' = \frac{u'v - v'u}{v^2} = \frac{(e^x + e^{-x})(e^x - e^{-x}) - (e^x - e^{-x})^2}{(e^x + e^{-x})^2}$$

$$= \frac{v^2 - u^2}{v^2} = 1 - \frac{u^2}{v^2} = 1 - y^2$$

Lecture 9

翻译难的不是找对应单词, 而是词义消歧, 语序重排, 文化/语境适配

$$\hat{T} = \underset{T}{\operatorname{argmax}} P(T)P(S|T) \quad T \text{ 是 target} \quad S = \text{foreign.}$$

$$\text{fluency } P(E) \quad \text{faithfulness } P(F|E) \quad E = (e_1, \dots, e_I) \quad F = (f_1, \dots, f_J) \quad \text{外语}$$

$$E^* = \underset{E}{\operatorname{argmax}} P(E|F) = \underset{E}{\operatorname{argmax}} \frac{P(F|E)P(E)}{P(F)} = \underset{E}{\operatorname{argmax}} P(F|E)P(E)$$

翻译方向是 $F \rightarrow E$ $\uparrow$ 固定

IBM Model 1 给定 E , 生成 F

1. 生成 F 的长度 J
2. 生成 Alignment A
3. 根据 A , 从 E 一对一生成 F .

$$E \rightarrow A \rightarrow F \quad P(F|E) = \sum_A P(F, A|E)$$

① bag of words. 所有 alignment 位置一样. 不懂语序 所有 alignment 相同 prob.

② independent word trans 只看词对, 不看上下文 $P(f_j | e_{a_j})$ 独立

③ 只能处理简单 one to many 但多个词经常是联合生成

$$P(F, A|E) = \frac{\sum}{(I+1)^J} \prod_{j=1}^J P(f_j | e_{a_j}) \quad I \text{ 和 } J \text{ 能匹配的 prob 是 } \frac{1}{(I+1)^J}$$

$\uparrow$

Lecture 10

1. HMM 对齐联合概率 $\rightarrow$ 先决定长度 I

$$P(F, A|E) = P(J|I) \prod_{j=1}^J P(a_j | a_{j-1}, I) P(f_j | e_{a_j})$$

根据上个 alignment 生成新的

用这个 alignment 生成 f_j NULL

Markov Assumption: ① $P(a_j | f_1^{j-1}, a_1^{j-1}, e_1^I) = P(a_j | a_{j-1}, I)$ 当前 alignment 只与上一个 alignment 有关

② $P(f_j | f_1^{j-1}, a_1^j, e_1^I) = P(f_j | e_{a_j})$ 当前 foreign word 只由它对齐到的目标词生成, 不依赖于别的

$$P(F|E) = P(J|I) \times \sum_A \prod_{j=1}^J P(a_j | a_{j-1}, I) P(f_j | e_{a_j})$$

2. Transition prob 应鼓励 local alignment

对 $P(a_j | a_{j-1}, I)$, 希望 $a_j \approx a_{j-1}$ 看相对距离, 不看绝对位置

trans model language model

3. Decoding $\hat{E} = \underset{E}{\operatorname{argmax}} P(F|E)P(E)$

希望最大化 $P(E|F)$

若 E 已知, 则可用 Viterbi
但翻译时 E 未知 变成在所有 E 中搜

带 bigram language model 的 general coding 是 NP-complete

必须用 heuristic search

3.1 Search-based decoding 把一个状态定义为 partial translation

包含 ① 已翻译的 foreign span ② 对应生成的 target words ③ 当前 accumulated score
searching tree 节点为 partial translation

3.2 Best-first Search: 每次扩展 score 最高的 简单, 但需维护大量候选节点. 只看当前易陷于 local

3.3 A* search: 当前质量 + 未来潜力 $f^*(p) = g(p) + h^*(p) \times \text{未来质量}$ 但 $h^*(p)$ 难估准, 需要 heuristic
综合优先级

3.4 Beam Search 最常见: 每一步只保留 top-k 个当前最好状态

- ① 初始化 Beam ② 扩展 Beam 所有状态 $\uparrow$ beam size
③ 对所有扩展打分 ④ 保留 top-k ⑤ 重复直到生成结束

比 greedy 稳, 比 exhaustive search cheap
实践效果好; 但不保证全局最优, k 太小易错过好翻译, k 太大计算成本高

4. Evaluations BLEU

n-gram precision $prec_n = \frac{\sum_{c \in \text{candidates}} \sum_{n\text{-gram} \in c} \text{Count match}(n\text{-gram})}{\sum_{c \in \text{candidates}} \sum_{n\text{-gram} \in c} \text{Count}(n\text{-gram})}$

BLEU = $(\prod_{n=1}^4 prec_n)^{\frac{1}{4}}$

几何平均

只有 unigram 过于宽松

clip

对每个 n-gram, 限制在参考译文中最多被匹配的次数

不超过参考中该 n-gram 最多出现次数

一个 5 词句, 分母 for unigram 是 5

5. Seq 2 Seq 直接学 $P(y_{1:T} | x_{1:T})$ 输入输出长度可以不同

5.1 encoder $h^{(t)} = f(h^{(t-1)}, x^{(t)}; \theta)$

5.2 decoder $d^{(t)} = f(d^{(t-1)}, y^{(t-1)}; \theta)$

5.3 MLE $P(y_{1:T} | x_{1:T}) = \prod_{t=1}^T P(y_t | v, y_{1:t-1})$ 前一个生成目标词

5.4 问题 ① fix-length hidden vector 会把 source sentence 压到固定长, 无论 sequence 多大 $P(y_t | v, y_1, \dots, y_{t-1}) = \text{softmax}(o^{(t)})$
 $o^{(t)} = Vd^{(t)}$

② 仍然不能 encode long-range dependency: 有 sequential recency bias
较近 token 过矩阵乘法少, 更易保留, 较远被递归, 易被冲淡

③ Gradient vanishing / explosion

④ 无法并行计算, 必须按时间顺序算

5.5 Attention mechanism: 生成一个词可动态看 source sentence 所有 hidden state

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

$$\alpha_{t,j} = \frac{\exp(\text{Score}(t,j))}{\sum_k \exp(\text{Score}(t,k))}$$

$$\text{Score}(t,k) = \langle s_{t-1}, h_t \rangle = s_{t-1}^T W_a h_t$$

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

优: ① 缓解 fixed-length vector bottleneck 不足: 仍是 RNN, 不能并行

② 可直接访问 source sentence 的 hidden state

③ long-range dependency

④ attention weight 有一定解释性

Lecture 11

1. Transformer { PE, MLP, Masked MHA, RC, LN } expressive
更好训

model should be large data

1.1 PE $PE(pos, 2i) = \sin(\frac{pos}{10000^{2i/d}})$ $PE(pos, 2i+1) = \cos(\frac{pos}{10000^{2i/d}})$

pos 是 token 在序列中位置; d - embedding dimension, $2i$ 和 $2i+1$ 是每个 embedded 维度
 $ht = E(W_t) + PE_t$ $i \uparrow \rightarrow T \uparrow$ 拟合了不同距离的 dependency

1.2 MLP

$dim(x) \rightarrow dim(x) * 4 \xrightarrow{GELU} dim(x) * 4 \rightarrow dim(x)$ 进出 dim 同, 可 stack 模块
因为 self-attention layer 只能找 linear combination of value vector

MLP 可引入非线性 $\downarrow$ batch size, token length

1.3 Self-Attention

$Q = XW_Q$ $K = XW_K$ $V = XW_V$

$\Rightarrow Q, K, V \in \mathbb{R}^{B \times n_h \times T \times h_s}$

$(n_embd, 3 * n_embd)$ $self_c_atten = [W_K, W_Q, W_V] \in (n_embd, 3 * n_embd)$
 $q, k, v = self_c_atten(x)$ $\uparrow$ 头数 $\uparrow$ 头大小
 $k = k.view(B, T, self.n_head, C // self.n_head).transpose(1, 2)$
把 C 拆成 $n_head * size_head$

$A = softmax(\frac{QK^T}{\sqrt{d_k}})$

$att = (q @ k.transpose(-2, -1)) \Rightarrow [B, n_h, T, T] / \sqrt{math.sqrt(k.size(-1))}$
 $[B, n_h, T, h_s]$ $[B, n_h, h_s, T]$

$(T \times h_s) * (h_s \times T) \Rightarrow [T \times T]$ 上三角为 0 的矩阵
 $att = att.masked_fill(self.bias[:, :, T, T] = 0, float('-inf'))$ 除以 h_s
 $att = F.softmax(att, dim=-1)$ 不让看 future
 $y = att @ v \Rightarrow [B, n_h, T, h_s]$ $softmax(-\infty) = 0$

$y = y.transpose(1, 2).contiguous().view(B, T, C)$

单个 attention head 只能学一种匹配, 多头让模型在不同子空间学不同关系

1.4 Layer norm $\text{layer-norm}(x) = \frac{x - \text{mean}(x)}{\sqrt{\text{variance}(x) + \epsilon}} \cdot \gamma + \beta$ learnable para

(过完MHA和feed-forward都有)

希望层内, the scale of activations remain the same about the input
 ↓ 值
 多层后仍如此 如果不这样, 若有层扩大/缩小 activation scales, 会出现 saturation of softmax $\Rightarrow$ vanishing gradient

Batch norm: 对每个特征通道C单独算. 对每个C, 收集BXT来做norm

Layer norm: 该样本内C和T, 用CXT来做norm

BN

LN

Batch size 对BC强依赖 不依赖

不同长度 sequence 有问题, padding会干扰 独立处理每个位置, padding不污染整个batch统计量

训/测统一性 训练用batch统计量

测试用全局滑动平均

存在不一致

训/测一样

不同activation之间尺度差异变小, softmax不易saturate

1.5 RC $x = \text{MHA}(x) + x$ $x = \text{MLP}(x) + x$ 只需学 Δ

$\text{MLP}(x) = W_2 \text{GeLU}(W_1 x + b_1) + b_2$

1.6 Training $\text{Raw Text} \xrightarrow{\text{tokenizer/encoder}} \text{Token indices} \rightarrow \text{binary files} \xrightarrow{\text{batch}} \text{输入} x \text{ 输出 } y$
 $\rightarrow$ next token prediction loss

$x = [w_1 : t]$ $y = [w_1 : t+1]$ $p_{\theta} = (w_{t+1} | w_{\leq t})$

$\mathcal{L}(\theta) = - \sum_{t=1}^T \log p_{\theta}(w_{t+1} | w_{\leq t})$

1.7 Inference: 生成慢 $\text{idx}_{\text{cond}} \in \mathbb{R}^{B \times \text{block-size}}$ $\text{logits} \in \mathbb{R}^{|B| \times |V|}$

$\text{logits} \rightarrow \text{softmax} \rightarrow \text{prob} \rightarrow \text{sample next token}$

把新token拼回序列重复生成

低效: ① 自回归得一个token = 一个生成

② 若无KV cache, 则每生成一个新token都要重新生成所有历史token

Lec 12

1. 定义

Pretraining: 大量通用学

Mid training: 特定高质量继续训练

Post training / alignment: base model 转成遵循指令, 符合人类偏好

2. Pretrain

2.1 Why? ① Transferability across task

② Improved task performance: 相比从头训, pretrain 把参数带到更平滑, 有利的 loss landscape 尤其在数据少时

2.2 Impact factors ① Model Architecture 决定表达能力和 scalability 重要但固定

② Training Object

③ Data (数, 质, 范围)

④ Hyper parameters (η , BS) 在大 data 下不那么重要

很重要

2.3 Data filtering deduplication $\Rightarrow$ Acc $\uparrow$

避免反复看到重复, 浪费容量, 增加 memorization 风险

2.4 Objective

2.4.1 MLM: 把输入部分词替换为 Mask, model 完型填空. 没有 causal mask
只用 encoder 部分

$h_{1:T} = \text{Encoder}(w_{1:T}) \quad z_i = Ah_i + b$ loss 只加在有 mask 位置

$$Z_{\text{mask}} = - \sum_{i \in \text{Mask}} \log p_{\theta}(w_i | \tilde{x})$$

随机选 15% 预测, 80% 换成 [MASK], 10% random token 10% 原词不变
SFT 时没有 [MASK], 若 pre-train 老看到 [Mask] 会 mismatch

输入 $e_i = e_i^{\text{token}} + e_i^{\text{segment}} + e_i^{\text{position}}$

[CLS] [SEP] 分句

只用 decoder 部分

sentence A/B

2.4.2 Autoregressive

输入 $[w_{1:T}] \rightarrow$ 得到每个位置 $[h_{1:T}]$ 拿 h_T 来 linear softmax

$$h_{1:T} = f(w_{1:T})$$

$$z_t = Ah_{t-1} + b$$

pros: ① 目标统一 - 简单

② 训练与推理一致

③ Scaling behaviour 强

$$\min_{\theta} \sum_{x \in D_{\text{train}}} \sum_t -\log p_{\theta}(w_t | w_{<t})$$

cons: ① 只能用 left-to-right context

② 自回归生成慢

③ error accumulation

下游任务, 任务转文本,
再输出加 linear layer

Lecture 13

Pretrained LM 能继续补全, 但还不够多 不编
希望 LLM helpful harmless honest
只续写, 不回答问题

1. SFT { single task: 只训练特定任务, 泛化弱
multi task / instruction tuning: 多任务混合, 让 model 学会服从指令

1.1 数据来源 { 早期: 人工标注. pros: ① 高质量 ② 不依赖 teacher model ③ 避免 self-generated error amplification
cons: ① 慢 ② 贵 ③ 多样性有限 ④ 规模有限
LIMA: Less is more for Alignment 不需那么多
后期: AI-generated SFT data: self instruct
少量 seed tasks → LLM 生成新 instructions → 判断 task 是否需 input → filtering
→ task pool 继续迭代

数据格式 { 当 teacher model 足够强, 人工标注 → 机器标注
Alpaca format single-run instruction / input / output
GPT format multi-run
chatML format role-based chat conversations: role + value
system / user / assistant

1.2 Loss $D = \{(x_i, y_i)\}$ $\hat{\theta} \leftarrow \argmin_{\theta} \mathcal{L}(\theta) = \sum_{i \in D} \sum_{j=1}^m -\log(P_{\theta}[y_{i,j} | x_i, y_{i,1:j-1}])$

Do NOT Count x in the loss using a loss Mask

prompt loss weight { 0 → 只学怎么回答
0.1 → 保留少量语言建模信号, 缓解 catastrophic forgetting

1.3 效果 SFT 减少 hallucination most effective 甚至比 PPO 更直接
更 obey 指令 helpful, harmless honest

1.4 Industrial SFT

SWE-lego { 输入: repo + bug fixing task + system prompt
output: bug-fixing trajectory

read files → edit files → run tests → observe execution feedback → refine edits

Synthetic data curation { LLM rewrite
AST reformulation 生成 bug

1.5 Multi-task SFT 不是越混越好 三类能力 general / math / coding

混太多 → 已有能力 ↓

小数量是好的

Dual stage { 1. 学习强化 specific ability
2. 混少量 general data 避免遗忘, 保留泛化

Lec 14 LoRA QLoRA 4bit-base LoRA

Full FT $\xrightarrow{\text{太贵}}$ PEFT $\xrightarrow{\text{低维访问}}$ LoRA $\rightarrow$ QLoRA

1. 参数量 1 byte = 8 bits

对一个 16-bit model Full FT

weight 16 bits fp16
gradient 16 bits

$$65B = 65 \times 10^9 \times 12 \text{ bytes} = 780 \text{ GB}$$

Adam first momentum 32 bits fp32

Adam second momentum 32 bits fp32
96 bits = 12 bytes

$\Rightarrow$ ensure stability of historical data

$$m_t = \beta_1 m_{t-1} + (1 - \beta_1) g_t$$

$$v_t = \beta_2 v_{t-1} + (1 - \beta_2) g_t^2$$

$$\hat{m}_t = \frac{m_t}{1 - \beta_1^t}$$

$$\hat{v}_t = \frac{v_t}{1 - \beta_2^t}$$

$$\theta_t \leftarrow \theta_{t-1} - \alpha \frac{\hat{m}_t}{\sqrt{\hat{v}_t + \epsilon}}$$

2. PEFT: 只更新一小部分参数

① Parameter Overload: 不需改变太多

② Intrinsic dim: 特定任务在低维子空间: pretrain 已把空间组织好
固定随机投影矩阵

$$\theta^D = \theta_0^D + \theta^d M$$

低维子空间

$$M = HG \Pi H^B$$

能达到 Full FT 90% performance
最小 d

① 相同 base model, 不同任务需不同 intrinsic dimension

② pretrain step $\uparrow \Rightarrow$ intrinsic dim $\downarrow$ 本来就厉害

③ #params $\uparrow \Rightarrow$ intrinsic dim $\downarrow$

④ intrinsic dim $\uparrow \Rightarrow$ 多任务表现更差

3. Intrinsic Dim $\Rightarrow$ Low Rank 参数更新 ΔW 形状大, 但 rank 小

$$\text{rank}(\Delta W) = r \ll \min(d, k)$$

$\downarrow$

4. LoRA

$$W = W_0 + \Delta W = W_0 + BA$$

训练 $(d+k)$ 参数 $\ll dk$

$$h = W_0 x + B A x \quad B \in \mathbb{R}^{d \times r} \quad A \in \mathbb{R}^{r \times k}$$

pros: ① 一个 backbone 可以对应多个任务的 Adapter

② No inference latency

$W' = W_0 + BA$ 无延迟, 无新增参数
打破对称, 不同 rank 不能全为 0. 因为对称更新有问题

4.1 Initialization

$$A \sim N(0, 6^{-2})$$

$$B = 0$$

不会因随机 Adapter 破坏原始能力

frozen weight 不算 GD 了, 但仍占显存

一开始 $\Delta W = BA = 0$ 完全等效 base model

5. QLoRA: + Quantization

$$x^{\text{int8}} = \text{round} \left(\frac{127}{\text{absmax}(x^{\text{FP32}})} x^{\text{FP32}} \right) = \text{round} (c^{\text{FP32}} \cdot x^{\text{FP32}})$$

标准 INT4 有 16 个 level. 但均匀分布. LLM 权重集中在 0 附近. 中心高密度, 但标准下误差大

$$E[\text{error}] = \sum_v P(v) \cdot \text{error}(v, \text{quant}(v))$$

NV4 更优. 重新设计 16 level. 每个 bucket 覆盖 1/16 的权重, 减少 E[error]

5.2 Paged Optimizer 长序列训练, activation, gradient, optimizer state 会导致 VRAM $\uparrow$

不够多把部分挪去 CPU RAM. 用时再查回来

Lec 15

1. SFT局限: Loss是最大化reference answer的 token likelihood. 但开放性问题需要比较多个候选回答质量. 可能两词都非 answer, 但一个靠近, 一个远离但 penalty 相同

2. Reward function: measure the quality of multiple open-ended answers

$$R(s) \in \mathbb{R} \quad \max_{\theta} E_{\tilde{s} \sim p_{\theta}(s)} [R(\tilde{s})] \quad \text{高 reward 的回答更应该被频繁生成}$$

3. Policy Gradient: reward不可微

$$\theta_{t+1} := \theta_t + \alpha \nabla_{\theta_t} E_{\tilde{s} \sim p_{\theta_t}(s)} [R(\tilde{s})]$$

$$\nabla f(x) = f'(x) \quad \nabla \log f(x) = \frac{\nabla f(x)}{f(x)} \Rightarrow f(x) \nabla \log f(x)$$

↑
与0无关

$$\begin{aligned} \nabla_{\theta} E_{\tilde{s} \sim p_{\theta}(s)} [R(\tilde{s})] &= \nabla_{\theta} \sum_{\tilde{s}} R(\tilde{s}) p_{\theta}(\tilde{s}) = \sum_{\tilde{s}} R(\tilde{s}) \nabla_{\theta} p_{\theta}(\tilde{s}) = \sum_{\tilde{s}} R(\tilde{s}) p_{\theta}(\tilde{s}) \nabla_{\theta} \log p_{\theta}(\tilde{s}) \\ &= \sum_{\tilde{s}} p_{\theta}(\tilde{s}) \boxed{R(\tilde{s}) \nabla_{\theta} \log p_{\theta}(\tilde{s})} = E_{\tilde{s} \sim p_{\theta}(s)} [R(\tilde{s}) \nabla_{\theta} \log p_{\theta}(\tilde{s})] \end{aligned}$$

用 MC 估计

$$\approx \frac{1}{m} \sum_{i=1}^m R(s_i) \nabla_{\theta} \log p_{\theta}(s_i)$$

↑
采样数 ↑
reward 增加该序列生成概率的方向

只需对 $\log p_{\theta}(s_i)$ 求导 鼓励|抑制

$$\log p_{\theta}(s_i) = \sum_{t=1}^T \log p_{\theta}(w_t^i | W_i^{<t}) \quad \text{本来是 } \pi, \log \text{ 后变累加}$$

4. 为什么用 pairwise comparison?

4.1 R(s)来源: 人类回答 { 贵, 慢
absolute scoring 模糊性, 7和8边界不清
annotator calibration 不一致, 有人严格, 有人宽松

pairwise 更可靠: 两回答, 人判断哪个好. 相对判断比绝对评分稳定, inter-annotator agreement 更高

4.2 Bradley-Terry Model

若每个回答有一潜在强度 $r(s) \in \mathbb{R}$ $r(s) = RM_{\phi}(s)$

$$P(i > j) = \frac{\exp(r_i)}{\exp(r_i) + \exp(r_j)} \Leftrightarrow P(i > j) = \sigma(r_i - r_j) = \frac{1}{1 + \exp(r_j - r_i)} = \frac{\exp(r_i)}{\exp(r_i) + \exp(r_j)}$$

偏好只依赖 $r_i - r_j$, 不依赖绝对分数

4.3 Reward Model Loss Train Reward Model 和要训练的 LLM 无关

$$J_{RM}(\phi) = - E_{(s^w, s^l) \sim D} [\log \sigma(RM_{\phi}(s^w) - RM_{\phi}(s^l))] \Rightarrow \text{训练目标: 让 } RM_{\phi}(s^w) \text{ 尽可能大于 } RM_{\phi}(s^l)$$

s^w : winner s^l : loser 模型认为 winner 胜过 loser 的 prob

4.4 Reward Hacking: 利用 Reward Model 的漏洞

4.5 RLHF Objective with KL Penalty 希望最大化的

$$R(s) = RM_{\phi}(s) - \beta \log \left(\frac{p_{\theta}^{RL}(s)}{p_{\theta}^{PT}(s)} \right) \quad \text{KL penalty}$$

↑ ↑
R Original R (frozen)

PT 是 pretrain (frozen) RL 是正在训练的 policy model

Reward Model 鼓励模型符合人类偏好, KL penalty 让输出接近 $p^{PT}(s)$

Lec 16

动机: 朴素 policy gradient 可以优化 reward, 但 policy network 对参数更新很敏感
小更新可能带来分布巨大变化, 造成 policy collapse

$$E \log X \leq \log E(X)$$

1. KL Divergence $\| \theta_{\text{new}} - \theta_{\text{old}} \|$ 不可靠 相同参数变化, 不同区域可能导致完全不同的输出分布变化

$$D_{KL}(P||Q) = \sum_x P(x) \log \frac{P(x)}{Q(x)} \quad P \text{ 是旧 } Q \text{ 是新} \quad D_{KL}(P||Q) \geq 0 \text{ 当且仅当 } P=Q \text{ 时为}$$

2. Natural Policy Gradient $\sum p \log \frac{p}{q}$ $D_{KL}(P||Q) \neq D_{KL}(Q||P)$

目标是在 trust region 更新参数 $\log x \leq x-1$ $\log(1+x) \approx \frac{1}{1+x}$ $E p(x) - \sum p \log \frac{q}{p} \geq \log E q$

$$\max_d V(\theta+d) \quad \text{s.t.} \quad D_{KL}(\pi_{\theta} || \pi_{\theta+d}) \leq \epsilon$$

精确 KL 约束难求解 对 V 做一阶近似 D_{KL} 做二阶近似

$$V(\theta+d) \approx V(\theta) + \nabla_{\theta} V(\theta)^T d \quad D_{KL}(\pi_{\theta} || \pi_{\theta+d}) \approx \frac{1}{2} d^T F d$$

$$D_{KL} \approx D_{KL}(\pi_{\theta} || \pi_{\theta}) + (\nabla_d D_{KL})^T d + \frac{1}{2} d^T (\nabla_d^2 D_{KL}) d = \frac{1}{2} d^T F d$$

正好取最大值
导数为0

FIM
Fisher Information Matrix

FIM 半正定且对称 $d^T F d \geq 0 \quad \forall d$

↓

变成 $\max_d \nabla_{\theta} V(\theta)^T d \quad \text{s.t.} \quad \frac{1}{2} d^T F d \leq \epsilon$
开拉 $L(d, \lambda) = \nabla_{\theta} V(\theta)^T d - \lambda (\frac{1}{2} d^T F d - \epsilon)$ 令 $\lambda > 0$ 若满足, 则会大. 不满足, 会小.

$$\nabla_d L(d, \lambda) = \nabla_{\theta} V(\theta) - \lambda F d = 0$$

$$\text{令 } d = \alpha F^{-1} g \quad \alpha = \frac{1}{\lambda}$$

不妨令 $d = \alpha F^{-1} g$ 这样

$$d = \frac{1}{\lambda} F^{-1} \nabla_{\theta} V(\theta) = \theta_{\text{new}} - \theta_{\text{old}}$$

Step $F^T = F$

$$\text{又有 } \frac{1}{2} d^T F d - \epsilon = 0 \Rightarrow \frac{1}{2} \alpha^2 (F^{-1} g)^T F (F^{-1} g) = \frac{1}{2} \alpha^2 g^T F^{-1} g = \epsilon$$

$$\text{所以 } \alpha = \sqrt{\frac{2\epsilon}{g^T F^{-1} g}}$$

$$\Rightarrow d = \sqrt{\frac{2\epsilon}{g^T F^{-1} g}} F^{-1} g$$

步长非普通 learning rate, 而是由 trust region 和 Fisher metric 自动决定

3. TNPG 与 TRPO ($F^{-1} O(d^3)$ d could be millions)

3.1 Truncated natural policy gradient (把 $v = F^{-1} g \Leftrightarrow Fv = g$ 用 conjugate gradient 迭代求解)

Trust Region

3.2 TRPO = TNPG + line search

① 先算 Natural Gradient 更新方向 ② 尝试一大步, 是否改善, 是否小于 KL threshold ③ 不安全缩小步长
避免显示求逆, 仍属二阶优化. 需多次 backward / Hessian-product. 计算仍贵 重试
multiple iterative backward pass

4. PPO Proximal Policy Optimization

不用显式二阶优化, 不用 conjugate gradient. 用 Adam, 但用 clipping 限制新旧策略比近似实现 safe update

$\pi(a|s)$ or π_{old} : old policy $\pi'(a|s)$ or π_{new} : new policy

Probability ratio $r_{sa}(\pi') = \frac{\pi'(a|s)}{\pi(a|s)}$ $r_{sa} > 1$ a is more likely 在新 policy 下

Advantage $A(s, a) = Q(s, A) - V(s)$ MC rollouts
 $A > 0$ action A 比平均好 $A < 0$ 比平均差

Goal $\max r_{sa} \times A$

$$L^{CLIP}(\theta) = \hat{E} [\min(r(\theta) \tilde{A}, \text{clip}(r(\theta), 1-\epsilon, 1+\epsilon) \tilde{A})]$$

$$\text{clip}(r(\theta), 1-\epsilon, 1+\epsilon) = \begin{cases} 1-\epsilon, & r(\theta) < 1-\epsilon \\ r(\theta), & 1-\epsilon \leq r(\theta) \leq 1+\epsilon \\ 1+\epsilon, & r(\theta) > 1+\epsilon \end{cases}$$

NLP Lec 17

1. Recap $R(s) = RM\phi(s) - \beta \log \frac{p_{\theta}^{RL}(s)}{p^{PT}(s)}$

PPO advantage-actor-critic (AC2) 要 rollout

↑ identify true value of a rollout
↑ Policy LM
↑ Value model 相同架构, 但 delayed param 用来算 advantage
需维护 critic, ref model, reward model

GAE 估计 advantage

2. DPO motivation 希望 model 符合 preference data, 能否跳过 reward model + RL 直接优化 model?

closed-form solution to the reward function for simple MLE optimization
把 reward function $\rightarrow$ policy 概率比

2.1 从 RLHF 推 DPO

Objective $E_{y \sim p_{\theta}^{RL}(y|x)} [RM\phi(x, y) - \beta \log \frac{p_{\theta}^{RL}(y|x)}{p^{PT}(y|x)}]$

令 $\pi(y) = p_{\theta}^{RL}(y|x)$

构造拉 $\sum_y \pi(y) = 1$

$J(\pi) = \sum_y \pi(y) [RM(x, y) - \beta \log \frac{\pi(y)}{p^{PT}(y|x)}]$
 $-\beta (\pi(y) \log \frac{\pi(y)}{p^{PT}(y|x)} + \frac{\pi(y)}{p^{PT}(y|x)} \frac{p^{PT}(y|x)}{\pi(y)})$

$L(\pi, \lambda) = \sum_y \pi(y) [RM(x, y) - \beta \log \frac{\pi(y)}{p^{PT}(y|x)}] + \lambda (1 - \sum_y \pi(y))$
对 $\pi(y)$ 求偏导
只对其中一个 $\frac{\partial L}{\partial \pi(y)} = RM(x, y) - \beta (\log \frac{\pi(y)}{p^{PT}(y|x)} + 1) - \lambda = 0$

$\log \frac{\pi^*(y)}{p^{PT}(y)} = \frac{1}{\beta} RM(x, y) - \frac{\beta + \lambda}{\beta}$

$\pi^*(y) = p^{PT}(y) \exp(\frac{1}{\beta} RM(x, y) \exp(-\frac{\beta + \lambda}{\beta}))$ 每一项都有

theoretical optimal $\rightarrow p^*(y|x) = \frac{1}{Z(x)} p^{PT}(y) \exp(\frac{1}{\beta} RM(x, y))$
令 $Z(x) = \sum_y p^{PT}(y|x) \exp(\frac{1}{\beta} RM(x, y))$ $\uparrow$ 按 reward 指数求和

$Z(x)$ 要对所有求和, 无法计算

2.2 用 Policy 反推 Reward Model

$p^*(y|x) = \frac{1}{Z(x)} p^{PT}(y|x) \exp(\frac{1}{\beta} RM(x, y))$

取 log $\log p^*(y|x) = -\log(Z(x)) + \log p^{PT}(y|x) + \frac{1}{\beta} RM(x, y)$

$\Rightarrow RM(x, y) = \beta \log p^*(y|x) + \log(Z(x)) - \beta \log p^{PT}(y|x)$

定义 $RM_{\theta}(x, y) = \beta \log \frac{p^*(y|x)}{p^{PT}(y|x)} + \beta \log(Z(x))$

DPO 定义一个由 policy 隐含表示的 reward model

但仍有 $Z(x)$, 算不了. 用 Bradley-Terry 消掉

Reward Model Loss $P(y_w > y_l | x) = \sigma(r(x, y_w) - r(x, y_l))$

代入 $Loss = -\log \sigma(r(x, y_w) - r(x, y_l))$

$$RM_\theta(x, y_w) - RM_\theta(x, y_l) = \beta \log \frac{P_\theta(y_w | x)}{P^{PT}(y_w | x)} + \beta \log(z(x)) - \beta \log \frac{P_\theta(y_l | x)}{P^{PT}(y_l | x)} - \beta \log(z(x))$$

$$= \beta \log \frac{P_\theta(y_w | x)}{P^{PT}(y_w | x)} - \beta \log \frac{P_\theta(y_l | x)}{P^{PT}(y_l | x)}$$

$$L_{DPO}(\theta) = -E_{(x, y_w, y_l) \sim D} [\log \sigma(\beta \log \frac{P_\theta(y_w | x)}{P^{PT}(y_w | x)} - \beta \log \frac{P_\theta(y_l | x)}{P^{PT}(y_l | x)})]$$

pairwise 只关心 reward difference

对于 y_w , 让 model 相比原来更愿意生成

y_l

不愿意生成

增大两者差值

不用显式训练 Reward Model, 不用 PPO rollout. Reference model 还在, 提供不要离原来太远的基准

开源友好. 轻量. 简洁 便宜, plug-and-play. 对一般 alignment 很有效

PPO 仍有更强全局探索, 性能上限更高. 适合复杂. 长文本

3. GRPO: 去掉 Critic 的 RL, 让同一个 prompt 生成的多个回答互相比, 用 group baseline 算 advantage

对同一问题 q , 生成 G 个回答: $o_1: G$. 对应 reward: $r_1: G$

$$A_i = \frac{r_i - \text{mean}(r)}{\text{std}(r)}$$

同一个 query 的相对分数

i 是第 i 个回答第 t token

3.1 GRPO Loss token-level ratio $r_{i,t}(\theta) = \frac{\pi_\theta(o_{i,t} | q, o_{i,<t})}{\pi_{\theta_{old}}(o_{i,t} | q, o_{i,<t})}$

$$L_{clip} = \min [r_{i,t}(\theta) \hat{A}_{i,t}, \text{clip}(r_{i,t}(\theta), 1-\epsilon, 1+\epsilon) \hat{A}_{i,t}]$$

$$J_{GRPO}(\theta) = E_{q, \{o_i\}} \left[\frac{1}{G} \sum_{i=1}^G \left(\frac{1}{T_i} \sum_{t=1}^{T_i} \min(r_{i,t}(\theta) \hat{A}_i, \text{clip}(r_{i,t}(\theta), 1-\epsilon, 1+\epsilon) \hat{A}_i) - \beta D_{KL} \right) \right]$$

第 i 个长度

4. Extensions

4.1 KTO: 解决 DPO 对 paired preference 依赖. DPO 要一个 prompt 下有 winner, loser

但现实更多是好/不好, 很少有配对 只要 desired / undesired

更依赖 SFT 初始化等超参

4.2 SimPO: 解决 DPO 的长度偏置 PPT 呢? 去掉了 reference model

reward margin

保证两者间有严格差距

DPO 问题: 序列越长, token-level log prob 累积差异越大, 即使每个 token 优势很小

长回答可能因果效应扩大 winner-loser 的 reward gap

输出 length $L_{SimPO}(x, y) = \frac{\beta}{|y|} \log \pi_\theta(y | x)$ $L_{SimPO} = -\log \sigma(r_{SimPO}(x, y_w) - r_{SimPO}(x, y_l) - \gamma)$

4.3 DAPO 解决 GRPO 的 clipping 抑制探索问题

GRPO | PPO 的 clipping 对称 $(1-\epsilon, 1+\epsilon)$

对低概率高 reward 的探索 token 上界太限制 exploration

DAPO: 给低概率但正 advantage 更高 upper bound

$\Rightarrow (1-\epsilon_{low}, 1+\epsilon_{high})$

NLP Lec 18 Synthetic Data

1. 为什么需要

- ① 真实数据增长慢, 但大模型消耗数据越来越快
- ② 重复同一批数据并不等价于新数据
- ③ 合成数据可降 hallucination
- ④ 真实 instruction data 可能质量低
- ⑤ 合成数据可用于 evaluation

2. Synthetic Data 类型

self-instruct: 服务 instruction following; code data: 可通过 execution 自动验证
math data 关注推理/答案 preference data 服务 RLHF / DPO / RM / Alignment

3. Prompting a teacher model

cons: teacher 错, student 错

$x \xrightarrow{\text{Teacher LLM}} (x, y)$

4. Retrieve and Transform

teacher model 生成的缺少 grounding, 易产生不真实样本

Corpus $\xrightarrow{\text{Retrieve}}$ z $\xrightarrow{\text{Transform}}$ (x, y)
已有真实 data 检索到 再把 z 改写成目标格式

pros: grounded in real knowledge sources
hallucination $\downarrow$

cons: quality heavily depends on retrieval quality

5. Extract and Rewrite

Web 上有 good data, but polluted by 广告, 导航, 无关文本.

Raw Web $\rightarrow$ Extract useful Passage $\rightarrow$ Score/Filter $\rightarrow$ Rewrite into IFT Format

Pros: standardize formatting

cons: 依赖 rules/filters. 且需 LLM filter, 很贵

6. Reparsing: 不引入新知识, 改表达形式

$d \xrightarrow{\text{Rephrase Model}} \{d_1', d_2', \dots, d_k'\}$ 一个原始多改写

pros: 保留原知识, hallucination $\downarrow$

cons: LLM 改写引入自己风格, tone 更统一, 降低多样性

7. knowledge graph extraction: data 少/私有/offline 从小规模文档构建知识图谱, 再生成本

subject 关系 客体实体
(s, p, o)

pros: 小 corpus $\rightarrow$ 大 synthetic corpus

cons: relation extraction 错误 / LLM 误解关系 $\rightarrow$ 数据不真实

8. AI Rating / Hybrid Preference: human 准贵, AI 便宜不准

$r(x) \in \{\text{Human}, \text{AI}\}$ 把难给入 r , 简单 AI r

9. Self-Instruct: Teacher LLM贵: $\rightarrow$ 让模型从少量AI seed自动扩展任务

生成instruction和instance

$S_0 \rightarrow G(S_0) \rightarrow S_1 \rightarrow G(S_1) \rightarrow \dots$

$\uparrow$
AI Seed Task

$S_t := t$ 轮后任务池

pros: ① instruction data $\uparrow$

② cons: ① seed决定质量上限 ② 生成任务可能表面多样, 实则单-

10. Self-guide: 在self-instruct上加结构化示例/过滤

input generation $\rightarrow$ input filter $\rightarrow$ output generation $\rightarrow$ input-output pair filter

通过task sample控制数据生成, 减少随机低质输出

cons: 示例本身 incomplete / bias

11. Evol-Instruct: 从instruct出发, 迭代t次

如 In-Breadth Evolving, Deepening, 加限制, 加推理...

生成复杂任务. 复杂 $\stackrel{?}{=}$ 好

12. Multi-agent Methods: 多agent互评, 修订, 过滤.

pros: 减少单模型误差 cons: 互相强化错误, 贵

13. Evaluation

13.1 Correctness: LLM1执行代码, 比GT...

$\Rightarrow \{0, 1\}$ 满意?

13.2 Complexity $(x) = \alpha s(x) + \beta q(x) + \gamma k(x)$

$\uparrow$ 推理步数 $\uparrow$ 约束 $\uparrow$ 知识深度

简单 $\Rightarrow$ 学不会 难的

难 $\Rightarrow$ 不现实

13.3 Diversity

$D = \{x_1, \dots, x_n\}$ ① self-BLEU(D) = $\frac{1}{n} \sum_{i=1}^n \text{BLEU}(x_i, D \setminus \{x_i\})$

低 $\Rightarrow$ 多样性 $\uparrow$

② Similarity Matrix $K \in \mathbb{R}^{n \times n}$

$K_{ij} = \sin(|x_i, x_j|)$ 对K做谱分解, 用特征值分布的

entropy 衡量 "有效独特样本", 谱集中 $\Rightarrow$ diversity 低

③ Diversity Greedy Selection

$\max_{x_j \in S} \text{sim}(x_i, x_j) < r$ 加入 x_i

13.4 Fidelity: 是否忠于原始事实或 transformation 前句子

-致性

RoBERTa-MNLI $(x, x') < 0.75 \Rightarrow$ 忽略

14. 局限

① Model Collapse 丢失细节, 保留模板表达

② lack of data provenance 合成数据吸收进参数后, 后续犯错难以追踪来源

③ data leakage 复杂 rephrasing 难查

NLP Lec 19 Scaling Law

1. Scale

N : 参数量 D : 训练 token 数 C : 训练计算量 in GFLOPS
希望平衡于 N, D, C

2. FLOPS: FLOP floating-point-operation 加、减、乘、除都算

transformer: 以矩阵乘为主, 忽略 bias, Layer Norm, residual, nonlinearity, softmax

$$\begin{matrix} m \times n & R^n & R^m & \text{output} \\ A \times & \Rightarrow 2mn & \text{1个要算 } n \text{ 乘 } n \text{ 加} \\ & m \text{ 个} & \Rightarrow \underline{\underline{2mn}} \end{matrix}$$

$$\begin{matrix} m \times n & n \times p \\ AB \in R^{m \times p} & m \times p \text{ output} \\ & \text{一个 output 要 } n \text{ 乘 } n \text{ 加} \Rightarrow 2mnp \end{matrix}$$

训练成本 $C \approx 6ND$

每个 token 前传 每个参数乘加一次 $\approx 2N \Rightarrow 2ND$ D token
反传约为前传两倍
(一个算 dW , 一个算 dX) $\approx 4N \Rightarrow 4ND$
 $\downarrow$
 $6ND$

$$N \begin{cases} QKV + \text{投影} & 4 \times D \times D = 4D^2 \\ FFN(\text{中间维度 } 4D) & D \times 4D + 4D \times D = 8D^2 \Rightarrow 12D^2 \\ \text{一层} & N = 12Ld^2 \end{cases}$$

$$N = 12 \times n_{\text{layers}} \times d_{\text{model}}^2 + (n_{\text{vocab}} + n_{\text{pos}}) \times d_{\text{model}}$$

embedding 的. 若 pos embed 可学

$$\text{Duration} = \frac{C}{\text{cluster throughput}}$$

10^3 忽略 softmax, GELU, layernorm 等
GPU 通信占 15% - 30% 时长
只能到 GPU 性能 30% - 60%

3. Kaplan Scaling Law: pretrained cross-entropy loss 随 scale 呈 power-law 下降

$$L(C) = AC^{-\alpha} + L_{\infty}$$

比例常数 L_{∞} 无限扩展仍无法消除的 loss

$\Rightarrow$ 更大模型有更陡 loss curve, 同样 data 下更能有效提取模式, 达到更低最终 loss
局限: 易被误解成 model size 越大越好, 但未考虑固定 compute 下数据量不足的问题

4. Chinchilla Scaling Law: 在固定 compute 下, 很多旧大 model 是 under-trained: 参数太多, 数据太少

$$C \approx 6ND \Rightarrow R \approx \frac{C}{6N}$$

对每个 C , 找 loss 最低的 N^* 和 D^*

$$N_{\text{opt}} \propto C^a \quad D_{\text{opt}} \propto C^b \quad a, b \text{ 接近 } 0.5 \quad \text{而 Kaplan } N_{\text{opt}} \propto C^{0.73} \quad D_{\text{opt}} \propto C^{0.27} \text{ 更偏参数}$$

近似同步增长

5. Variants {
- ① batch size和lr也有scaling law. $C \uparrow \Rightarrow BS^* \uparrow, lr^* \downarrow$. 可用小规模实验预测大规模训练超参
 - ② data 质量影响scaling. 少高质-样能打
 - ③ 多模态也可能存在SL

6. Limitation {
- ① pretrain loss $\neq$ 下游表现
 - ② inverse scaling: 模型越大某些任务越差: 喜欢记而不跟从指令; 模仿数据, 挑简单任务做
 - ③ data filtering {
 - 小compute下 干净小数据更好
 - 大compute下 加一些noisy有利
- 被正石角但不充分示例带偏

NLP Lec 20 KV-cache

1. Naive Decoding 浪费

$$p(x_{1:T}) = \prod_{t=1}^T p(x_t | x_{<t})$$

每一步生成新 tokens 时都要对当前上下文做 self-attention

第 t 步生成时, 历史 token KV 已算 $\sum_{t=1}^n O(t) = O(n^2)$

2. KV cache 用显存换时间: 历史 token KV 不会因新 token 带来而改变, 只存, 不算

At time step t $k_t = h_t W_k, v_t = h_t W_v, q_t = h_t W_q$

然后把新的 k_t, v_t append 到 cache

$$K_{1:t} = [K_{1:t-1}; k_t] \quad V_{1:t} = [V_{1:t-1}; v_t]$$

生成第 t 个 token 时

$$o_t = \text{softmax}\left(\frac{q_t K_{1:t}^T}{\sqrt{d_k}}\right) V_{1:t}$$

$$O(n^2) \rightarrow O(n)$$

3. Prefill vs Decode

3.1 Prefill: 处理输入 Prompt

$$X_{\text{prompt}} \rightarrow K_{\text{prompt}}, V_{\text{prompt}}$$

可以并行, 可以跨 token 运算 瓶颈是 compute-bound 指标 TTFT

但不能跨层并行, H 层依赖 L 层

Time to First Token

3.2 Decode 逐 token 生成 $x_t \sim p(x_t | x_{<t})$ 并行: 低. GPU 利用低

指标 Time Per Output Token

瓶颈: memory-bandwidth-bound

每步都从 HBM 读 KV Cache

4. KV Cache 大小

$$\text{Size} = B \times L \times 2 \times N_{\text{layers}} \times N_{\text{heads}} \times d_{\text{head}} \times 2 \text{ bytes}$$

$\uparrow$ K, V $\uparrow$ FP16

size $\propto B L N_{\text{layers}} N_{\text{heads}} d_{\text{head}}$ 长文档任务 OOM 元凶是 KV Cache

5. Flash Attention 减少 attention 中间矩阵写回

$$A = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)$$

A 很大, 读写慢.

不存完整 A, 而是 block by block 算. 在线维护 softmax 归一化项

对于 logits $m_i = q^T k_i$

softmax 权重为

$$w_i = \frac{\exp(m_i - m^*)}{\sum_j \exp(m_j - m^*)}$$

$$m^* = \max_j m_j$$

输出 $s = \sum_i w_i v_i$

在线计算时维护

$$d_i = \sum_{j=1}^J \exp(m_j - m^*)$$

$$s_i = \frac{s_{i-1} d_{i-1} + \exp(m_i - m^*) v_i}{d_i}$$

d_i

6. MHA $\rightarrow$ MQA $\rightarrow$ GQA: 减少 KV head

6.1 MHA KV cache 随头数线性增长

6.2 MQA 所有 query head 共享同一组 K, V, 极大压缩 KV Cache. 但表达能力 $\downarrow$

6.3 GQA 把 query head 分成若干组, 每组共享 K, V.
压缩比 $\frac{H_{kv}}{H_q}$ 每个不同

7. Page Attention { ① 每个请求预留最大长度, 但有些没用完
② 不同请求物理显存不连续, 没用的无法复用
logical block id $\xrightarrow{\text{block table}}$ physical block id

8. Continuous Batching & chunked prefill

Prefill 大矩阵. compute-bound
decode 小步读 KV-cache memory-bound

- 一个在 prefill, 其他 decode 要等

Pre fill (L) $\rightarrow$ 分块 prefill 把 decode 插 prefill 中

9. Streaming LLM 与 Attention Sink

* 上下文问题 { 全 KV-cache VRAM OOM
最近 window: 丢掉最初 token, 性能崩溃
每次重算 $O(L^2)$

初始几个 token 虽语义信息不强, 但稳定吸收大量 attention. 形成 attention sink
训练中 初始位置被最多看见. softmax 归一化也要某些位置接受冗余注意力

Streaming LLM KV Cache = Sink Tokens + Recent Window Tokens

$$K_{\text{cache}} = [K_{\text{sink}}; K_{t-w:t}]$$

$$V_{\text{cache}} = [V_{\text{sink}}; V_{t-w:t}]$$

NLP Lec 2 | Quantization, Pruning, Distillation

1. Format

若 cache 容量为 C bytes, 每个参数用 P_{bit} bits 能存 $N = \frac{C \times 8}{P_{bit}}$

FP16 (2 Bytes) $[-65504, 65504]$

INT8 (1 Bytes) $[-128, 127]$

INT4 (0.5 Bytes) $[-8, 7]$

2. Quantization

$$q = \text{clamp}(\text{round}(\frac{x}{s} + z), q_{min}, q_{max})$$

反量化 $\hat{x} = s(q - z)$

Scale: s zero-point: z

2.1 Absmax symmetric quan $s = \frac{\max(|x|)}{127}$ INT8

$[-\max|x|, \max|x|] \rightarrow [-127, 127]$

2.2 Zero Point 量 $s = \frac{\max(x) - \min(x)}{2 \times 127}$ 或 $s = \frac{x_{max} - x_{min}}{q_{max} - q_{min}}$

$z = q_{min} - \frac{x_{min}}{s}$

2.3 Precision Cliff: 4-bit 好 3-bit 太差

Performance $\approx f(\# \text{ params}) - g(\text{quantization noise})$

3. PTQ 与 QAT

PTQ: 预训练完直接量化. 快, 便宜, 无需额外训练. 低 bit 易掉点

QAT: Quantization-Aware Training 训练时模拟量化误差 精度恢复好, 需数据训练

3.1 QAT round(.) 不可导 $\Rightarrow$ 梯度消失

STE Straight-Through Estimator

FP32

forward 用 round backward 假装 round 是 identity

4. Activation-outlier: 某些 hidden vector 维度 ≥ 60 , 出现在至少 25% 层, 至少 6% token 中

问题来自 $s = \frac{\max(|x|)}{127}$ 若某 outlier 比普通值大 100 倍 $\Rightarrow$ 大量非 outlier 信息被压成 0

outlier 虽少, 但对 model 重要

$$L = \sum_i p_i l_i \text{ 最小化 } L \Rightarrow \text{罕见用长的 } l_i, \text{ 常见用短}$$

5. LLM.int8 (mixed precision quantization) outlier 单独高精度

$$xW = x_{normal} W_{normal} + x_{outlier} W_{outliers} \Rightarrow \text{INT8} \quad \text{FP16} \Rightarrow \text{FP16} + \text{FP16} \Rightarrow \text{FP16}$$

6. SmoothQuant Optimization 把 activation outlier 搬去平滑的 W

$$Y = (X \text{diag}(s)^{-1}) (\text{diag}(s)W)$$

smoothing factor

7. Pruning 删参. unstructured 如删1

7.1 Activation Pruning 与 WANDA

$$S_{ij} = |W_{ij}| \cdot \|X_j\|_2 \quad \text{weight 本身大不一定重要, 它连的 activation 经常大才重要}$$

7.2 Iterative Pruning $N \times N$

剪低分 blocks $\rightarrow$ finetune / retrain 恢复性能 $\rightarrow$ 重复

$$R_i = \|W_{\text{block}, i}\|_2 \quad \text{小的 masked zero}$$

8. Distillation

$\begin{cases} \text{hard label: 正确类别} \\ \text{soft: 概率分布} \end{cases}$

带 temperature $p_i^{(T)} = \frac{\exp(z_i/T)}{\sum_j \exp(z_j/T)}$

$T=1$ low entropy
 $T>1$ scales down logit 分布平滑 entropy $\uparrow$

pros: ① 完全不同架构

cons: ② 需训 student 且 teacher 得 available
比 PT 贵

9. 2:4

~~2:4~~ 2 bit

连续 4 个 变成 0

horizontal elements

2 bit 存 spatial index

2x throughput 吞吐

NLP Lec 22 MoE

一. 两核心组件 (Experts 多个子网络, FFN; Gating network/Router 谁干)

$$\hat{y} = \sum_{i \in E_k} G(x)_i E_i(x)$$

2. MoE in LSTM

⇒ 两专家 prob 归一化

$$G = \text{softmax}(W_g X)$$

⇒ 选 top-k 个 experts

$$I = \{i_0, i_1\} = \text{Topk}(G, k=2)$$

$$S_0 = \frac{G_{i_0}}{G_{i_0} + G_{i_1}}$$

$$S_1 = \frac{G_{i_1}}{G_{i_0} + G_{i_1}}$$

$$Y = S_0 \text{FFN}_{i_0}(X) + S_1 \text{FFN}_{i_1}(X)$$

3. MoE in Transformer

对于 token hidden state h , router 计算

Self-Attention → MoE FNN

$$P = \text{softmax}(W_r h)$$

router 学习参数

当前 token 分配给不同 experts prob

若选中 k 个 expert

$$\text{MoE}(h) = \sum_{k=1}^K P_k E_k(h)$$

对 token representation 的 FFN 变换

4. Routing Algo

Maximum: 选 1 个. pros: 低参数 cons: routing collapse 少数常用

Top-2/Top-k: Top 2 选 1 = $\text{Topk}(P, k=2)$ MoE(h) = $S_1 E_{i_1}(h) + S_2 E_{i_2}(h)$ 多数闲置

pros: 表达更强. 缓解 routing collapse 训练信号分配给多专家

cons: ① FLOPS ↑

② token 发到不同 device, communication overhead 更大

5. GLaM: 每 token 相似 & FLOPS

MoE 优于 Dense

6. Deep Seek-MoE: 区分

每个 token 必过 shared experts

通用知识

由 router 选

专门知识

每个 MoE 含 2 shared + 64 routed. 每次选 2 shared + 6 routed

token-level

7. MoE Training: 选择 Topk 不可微

$$\mathcal{L} = \mathcal{L}_{\text{task}} + \lambda \mathcal{L}_{\text{aux}}$$

主任务

辅助负载均衡损失, 鼓励 experts 均衡用

$$\mathcal{L}_{\text{task}} = -\frac{1}{T} \sum_{t=1}^T \log P_0(x_t | x_{<t})$$

$$h_t = u_t + \sum_{i=1}^N g_{t,i} \text{FFN}_i(u_t)$$

残差

$$g_{t,i} = \begin{cases} S_{t,i}, & S_{t,i} \in \text{TOPK}(\{S_{t,j} | 1 \leq j \leq N\}, K) \\ 0, & \text{otherwise} \end{cases}$$

$$S_{t,i} = G(u_t; \theta)_i$$

$$\mathcal{L}_{\text{aux}} = \alpha N \sum_{i=1}^N f_i P_i$$

router 给 expert 的 prob

$$P_i = \frac{1}{T} \sum_{t=1}^T S_{t,i}$$

token fraction

KT 个选择

8. Build MoE LLMs

8.1 Sparse Upcycling

copy FFNs to form experts { 可复用 dense model 的 pretrained knowledge
比从 0 开始训练经济高效
参数增加 DN 倍

8.2 Sparse Splitting

split FFNs to form { 激活 $\downarrow$
expert 容量 $\frac{1}{N}$ 单个 expert 表达有限
in

$$\begin{aligned} \max \sum_t h_t \cdot W_{\alpha t} \\ \text{s.t. } \forall e \sum_{t=0}^T 1_{\alpha t=e} = \frac{T}{E} \end{aligned}$$

1. Overview

1.1 GAN

$$\min_G \max_D v(D, G) = E_{x \sim P_{data}(x)} [\log(D(x))] + E_{z \sim p_z(z)} [\log(1 - D(G(z)))]$$

对D来说都希望最大化 G 希望小 D 希望大 假图

1.2 VAE

$$E_{q_\phi(z|x)} [\log p_\theta(x|z)] - D_{KL}(q_\phi(z|x) || p(z))$$

从 latent z 重构 x $z \sim$ latent noise prior

latent prior $N(0, I)$
希望 $q_\phi(z|x)$ 贴近 $N(0, I)$
encoder: $x \rightarrow z$ (latent distribution)

1.3 Diffusion: reverse

$$p_\theta(x_{0:T}) = p(x_T) \prod_{t=1}^T p_\theta(x_{t-1} | x_t)$$

其中 $p_\theta(x_{t-1} | x_t) = N(x_{t-1}; \mu_\theta(x_t, t), \Sigma_\theta(x_t, t))$

GAN 快. High quality

VAE 快. Mode Coverage / Diversity

Diffusion: High quality + Mode Coverage

2. Forward Process

$$q(x_t | x_{t-1}) = N(x_t; \sqrt{1-\beta_t} x_{t-1}, \beta_t I) \text{ 或 } q(x_t | x_{t-1}) = N(x_t; \sqrt{\alpha_t} x_{t-1}, (1-\alpha_t) I)$$

其中 $\alpha_t = 1 - \beta_t$

$$x_t = \sqrt{\alpha_t} x_{t-1} + \sqrt{1-\alpha_t} \epsilon_{t-1}, \epsilon_{t-1} \sim N(0, I)$$

若不用 $\sqrt{\alpha_t}$ 缩放, 连续加噪会导致 variance 不可控

$$x_{t-1} = \sqrt{\alpha_{t-1}} x_{t-2} + \sqrt{1-\alpha_{t-1}} \epsilon_{t-2}, \epsilon_{t-2} \sim N(0, I)$$

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

两个独立 Gaussian 相加仍是 Gaussian

$$= \sqrt{\alpha_t \alpha_{t-1}} x_{t-2} + \sqrt{1-\alpha_t \alpha_{t-1}} \epsilon$$

推广到任意 t , 定义 $\bar{\alpha}_t = \prod_{i=1}^t \alpha_i$ $\epsilon \sim N(0, I)$

$$q(x_t | x_0) = N(x_t; \sqrt{\bar{\alpha}_t} x_0, (1-\bar{\alpha}_t) I)$$

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

$$\text{定义 } \alpha_t = 1 - \beta_t, \bar{\alpha}_t = \prod_{i=1}^t \alpha_i$$

$\bar{\alpha}_t$ 控制 x_0 在 x_t 中还有多少信号.

$$t \uparrow \quad \bar{\alpha}_t \rightarrow 0 \Rightarrow x_t \approx \epsilon \sim N(0, I)$$

scheduler 控制 $\bar{\alpha}_t$. linear 信号可能会过早崩掉. cosine scheduler 会信号留存更久, 避免 collapse

3. Reverse Process : 从 noise 回 data

目标是学 $p_\theta(x_{t-1}|x_t)$ 因为加噪是 infinite small $\Rightarrow$ a reverse transition 建模成高斯

$$\Rightarrow p_\theta(x_{t-1}|x_t) = \mathcal{N}(x_{t-1}; \mu_\theta(x_t, t), \Sigma_\theta(x_t, t))$$

但真实 reverse posterior $q(x_{t-1}|x_t)$ 不可角平

$$q(x_{t-1}|x_t) = \frac{q(x_t|x_{t-1})q(x_{t-1})}{q(x_t)}$$

对整个数据空间积分 intractable
 $q(x_t) = \int q(x_t|x_{t-1})q(x_{t-1})dx_{t-1}$

但如果客额外 condition on x_0 $q(x_{t-1}|x_t, x_0) = \frac{q(x_t|x_{t-1}, x_0)q(x_{t-1}|x_0)}{q(x_t|x_0)}$
 tractable 高斯

$$q(x_{t-1}|x_t, x_0) = \mathcal{N}(x_{t-1}; \tilde{\mu}_t(x_t, x_0), \tilde{\beta}_t I)$$

其中 $\tilde{\mu}_t(x_t, x_0) = \frac{\sqrt{\alpha_t}(1-\alpha_{t-1})}{1-\alpha_t} x_t + \frac{\sqrt{\alpha_{t-1}}\beta_t}{1-\alpha_t} x_0$

4. 为什么预测 noise $\varepsilon \Leftrightarrow$ 预测 reverse mean

在 inference 中, 由于 x_0 未知, 但由于 $x_t = \sqrt{\alpha_t}x_0 + \sqrt{1-\alpha_t}\varepsilon$

可解出 $x_0 = \frac{1}{\sqrt{\alpha_t}}(x_t - \sqrt{1-\alpha_t}\varepsilon)$

用 neural network 预测 noise: $\varepsilon_\theta(x_t, t) \approx \varepsilon$

代入 posterior mean

直接

$$\mu_\theta(x_t, t) = \frac{1}{\sqrt{\alpha_t}} \left(x_t - \frac{\beta_t}{\sqrt{1-\alpha_t}} \varepsilon_\theta(x_t, t) \right)$$

model 不用预测 $x_{t-1}|x_0$ 只要 predict 加进去的 noise

$$\beta_t = 1 - \alpha_t$$

5. ELBO

目标最大化 $\max_\theta \log p_\theta(x_0) \Leftrightarrow \min_\theta -\log p_\theta(x_0)$

$p_\theta(x_0) = \int p_\theta(x_{0:T}) dx_{1:T}$ 不可算, 引入 $q(x_{1:T}|x_0)$

$$-\log p_\theta(x_0) = -\log \int q(x_{1:T}|x_0) \frac{p_\theta(x_{0:T})}{q(x_{1:T}|x_0)} dx_{1:T} \leq E_{q(x_{1:T}|x_0)} \left[\frac{q(x_{1:T}|x_0)}{p_\theta(x_{0:T})} \right]$$

LVLB

$$L_{VLB} = E_q \left[\log \frac{\prod_{t=1}^T q(x_t|x_{t-1})}{p(x_T) \prod_{t=1}^T p_\theta(x_{t-1}|x_t)} \right]$$

Jensen

已经假设给定真实数据 x_0

$$= E_q \left[\log \frac{q(x_T|x_0)}{p(x_T)} + \sum_{t=2}^T \log \frac{q(x_{t-1}|x_t, x_0)}{p_\theta(x_{t-1}|x_t)} - \log p_\theta(x_0|x_1) \right]$$

怎么推? $q(x_{t-1}|x_t, x_0) = \frac{q(x_t|x_{t-1}, x_0)q(x_{t-1}|x_0)}{q(x_t|x_0)} \Leftrightarrow q(x_t|x_{t-1}) = \frac{q(x_{t-1}|x_t, x_0)q(x_t|x_0)}{q(x_{t-1}|x_0)}$

由马尔可夫

$$q(x_t|x_{t-1}) = q(x_t|x_{t-1}, x_0)$$

$$L_{VLB} = E_q \left[\log \frac{q(x_T|x_0)q(x_{T-1}|x_T, x_0)}{p(x_T) \dots q(x_{T-1}|x_0)} \dots \right] = E_q \left[\log \frac{q(x_T|x_0)}{p(x_T)} \prod_{t=1}^T \frac{q(x_{t-1}|x_t, x_0)}{p_\theta(x_{t-1}|x_t)} \times \frac{q(x_0|x_1, x_0)=1}{p(x_0|x_0)=1} \right]$$

上下相消

$$= E_q \left[\log \frac{q(x_T|x_0)}{p(x_T)} + \sum_{t=2}^T \log \frac{q(x_{t-1}|x_t, x_0)}{p_\theta(x_{t-1}|x_t)} - \log p_\theta(x_0|x_1) \right]$$

$$= E_q \left[D_{KL}(q(x_T|x_0) || p(x_T)) + \sum_{t=2}^T D_{KL}(q(x_{t-1}|x_t, x_0) || p_\theta(x_{t-1}|x_t)) - \log p_\theta(x_0|x_1) \right]$$

塔性质 $q(x_{1:T}|x_0) = q(x_{1:T \setminus \{t, t-1\}} | x_t | x_0) q(x_{t-1} | x_t, x_0)$

两个 Gaussian KL $E_{q(x_{1:T}|x_0)} = E_{q(x_{1:T \setminus \{t, t-1\}} | x_t | x_0)} E_{q(x_{t-1} | x_t, x_0)} \Rightarrow E_{x_{t-1} \sim q(x_{t-1} | x_t, x_0)} [-D_{KL}(\cdot)]$
 $D_{KL}(N(\mu_1, \sigma_1^2) || N(\mu_2, \sigma_2^2)) = \frac{1}{2\sigma_2^2} \|\mu_1 - \mu_2\|_2^2$

$$L_{simple} = E_{t, x_0, \varepsilon} [\|\varepsilon - \varepsilon_\theta(x_t, t)\|_2^2]$$

$$x_t = \sqrt{\alpha_t} x_0 + \sqrt{1 - \alpha_t} \varepsilon$$

6. Training Algo

1. Sample $x_0 \sim \text{data}$
2. Sample $t \sim \text{Uniform}(\{1, \dots, T\})$
3. Sample $\varepsilon \sim N(0, I)$
4. Construct $x_t = \sqrt{\alpha_t} x_0 + \sqrt{1 - \alpha_t} \varepsilon$
5. Predict $\varepsilon_\theta(x_t, t)$
6. Minimize $\|\varepsilon - \varepsilon_\theta(x_t, t)\|$

7. Sample Algo $\beta_t = \sigma_t^2$

1. Sample $x_T \sim N(0, I)$ $\beta_t = \tilde{\beta}_t = \frac{1 - \alpha_{t-1}}{1 - \alpha_t} \beta_t$

2. For $t = T, T-1, \dots, 1$:

$$x_{t-1} = \mu_\theta(x_t, t) + \beta_t z$$

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

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

↑ 其实是从 ~~x_0~~ 没有

$$x_t = \sqrt{\alpha_t} x_0 + \sqrt{1 - \alpha_t} \varepsilon$$

解 $x_0 = \frac{1}{\sqrt{\alpha_t}} (x_t - \sqrt{1 - \alpha_t} \varepsilon)$

↑ 用 $\varepsilon_\theta(x_t, t)$ 来代替

7. Diffusion 通常用 UNET

① noisy x_t

{ downsample: 提 global structure
 bottleneck: 全局语义
 upsampling: 恢复 spatial resolution
 skip-connection: 保留细节

time embedding $\text{Emb}(t) \quad \sin$

$$= q(x_t | x_{t-1}) = N(x_t; \sqrt{\alpha_t} x_{t-1}, \beta_t I), \beta_t = 1 - \alpha_t$$

$$q(x_{t+1} | x_t, x_0) = \frac{q(x_t | x_{t+1}, x_0) q(x_{t+1} | x_0)}{q(x_t | x_0)}$$

$$\log q(x_{t+1} | x_t, x_0) = \log q(x_t | x_{t+1}) + \log q(x_{t+1} | x_0) - \log q(x_t | x_0) + \text{const}$$

代入高斯密度

$$\log q(x_t | x_{t+1}) = -\frac{1}{2\beta_t} \|x_t - \sqrt{\alpha_t} x_{t+1}\|^2 + \text{const}$$

$$\log q(x_{t+1} | x_0) = -\frac{1}{2(1 - \alpha_{t+1})} \|x_{t+1} - \sqrt{\alpha_{t+1}} x_0\|^2 + \text{const}$$

$$\log q(x_t | x_0) = -\frac{1}{2(1 - \alpha_t)} \|x_t - \sqrt{\alpha_t} x_0\|^2$$

$$\mu: \sqrt{\alpha_t} \quad \sigma^2 = 1 - \alpha_t$$

$$\log q(\underline{x}_{t-1} | x_t, x_0) = \log q(x_t | x_{t-1}) + \log q(x_{t-1} | x_0) - \log q(x_t | x_0)$$

只关注 x_{t-1} 的
- 次项和二次项

$$= -\frac{1}{2\beta_t} \|x_t - \sqrt{\alpha_t} x_{t-1}\|^2 - \frac{1}{2(1-\alpha_{t-1})} \|x_{t-1} - \sqrt{\alpha_{t-1}} x_0\|^2$$

$$+ \frac{1}{2(1-\alpha_t)} \|x_t - \sqrt{\alpha_t} x_0\|^2$$

$$= -\frac{\alpha_t}{2\beta_t} x_{t-1}^2 - \frac{1}{2(1-\alpha_{t-1})} x_{t-1}^2 + \frac{\sqrt{\alpha_t}}{\beta_t} x_t x_{t-1} + \frac{\sqrt{\alpha_{t-1}}}{1-\alpha_{t-1}} x_{t-1} x_0$$

- $\frac{1}{2} A x_{t-1}^2 + B \cdot x_{t-1}$ 型
条件方差为 A^{-1} , 均值为 $A^{-1}B$

$$= -\frac{1}{2} \left(\frac{\alpha_t}{\beta_t} + \frac{1}{1-\alpha_{t-1}} \right) x_{t-1}^2 + \left(\frac{\sqrt{\alpha_t}}{\beta_t} x_t + \frac{\sqrt{\alpha_{t-1}}}{1-\alpha_{t-1}} x_0 \right) x_{t-1}$$

$$\text{Var} = \frac{1}{A} = \frac{1}{\frac{\alpha_t}{\beta_t} + \frac{1}{1-\alpha_{t-1}}} = \frac{\beta_t (1-\alpha_{t-1})}{\alpha_t - \alpha_t \alpha_{t-1} + \beta_t} = \frac{\beta_t (1-\alpha_{t-1})}{1 - \alpha_t \alpha_{t-1}} = \frac{\beta_t (1-\alpha_{t-1})}{1 - \alpha_t}$$

$$\text{均值 } \mu = \frac{\beta_t (1-\alpha_{t-1})}{1 - \alpha_t \alpha_{t-1}} \cdot \left(\frac{\sqrt{\alpha_t}}{\beta_t} + \frac{\sqrt{\alpha_{t-1}}}{1-\alpha_{t-1}} \right)$$

$$= \frac{\sqrt{\alpha_t} (1-\alpha_{t-1})}{1 - \alpha_t \alpha_{t-1}} x_t + \frac{\beta_t (1-\alpha_{t-1}) \sqrt{\alpha_{t-1}}}{(1-\alpha_t \alpha_{t-1}) (1-\alpha_{t-1})} x_0$$

$$= \frac{\sqrt{\alpha_t} (1-\alpha_{t-1})}{1 - \alpha_t} x_t + \frac{\sqrt{\alpha_{t-1}} \beta_t}{1 - \alpha_t} x_0 \quad -\frac{B}{\sqrt{\alpha_t}} = \frac{\sqrt{\alpha_{t-1}} \beta_t}{(1-\alpha_t) \sqrt{\alpha_t}}$$

将 $x_0 = \frac{x_t - \sqrt{1-\alpha_t} \varepsilon}{\sqrt{\alpha_t}}$ 代入

$$\mu_0(x_t, t) = \frac{\sqrt{\alpha_t} (1-\alpha_{t-1})}{1 - \alpha_t} x_t + \frac{\sqrt{\alpha_{t-1}} \beta_t}{1 - \alpha_t} \cdot \frac{x_t - \sqrt{1-\alpha_t} \varepsilon}{\sqrt{\alpha_t}}$$

$$\tilde{\mu} = A x_t + B \cdot \frac{x_t - \sqrt{1-\alpha_t} \varepsilon}{\sqrt{\alpha_t}}$$

$$= \left(A + \frac{B}{\sqrt{\alpha_t}} \right) x_t - \frac{B \sqrt{1-\alpha_t}}{\sqrt{\alpha_t}} \varepsilon$$

$$= \frac{\sqrt{\alpha_t} \sqrt{\alpha_t} (1-\alpha_{t-1}) + \sqrt{\alpha_{t-1}} \beta_t}{\sqrt{\alpha_t} (1-\alpha_t)} x_t - \frac{\sqrt{1-\alpha_t}}{1-\alpha_t} \varepsilon$$

$$= \frac{\sqrt{\alpha_t} (1-\alpha_{t-1})}{1-\alpha_t} + \frac{\sqrt{\alpha_{t-1}} \beta_t}{(1-\alpha_t) \sqrt{\alpha_t}}$$

$$= \frac{\sqrt{\alpha_t} (1-\alpha_{t-1})}{1-\alpha_t} + \frac{\beta_t}{(1-\alpha_t) \sqrt{\alpha_t}} = \frac{1}{(1-\alpha_t) \sqrt{\alpha_t}} \frac{\alpha_t (1-\alpha_{t-1}) + \beta_t}{1-\alpha_t}$$

$$= \frac{1}{1-\alpha_t} (1 - \alpha_t \alpha_{t-1}) = \frac{1-\alpha_t}{1-\alpha_t} \cdot \frac{1}{\sqrt{\alpha_t}} = \frac{1}{\sqrt{\alpha_t}}$$

$$\Rightarrow \frac{1}{\sqrt{\alpha_t}} \left(x_t - \frac{\beta_t}{\sqrt{1-\alpha_t}} \varepsilon \right)$$

NLP Lec 24 Score-based View & Text Diffusion

DDPM: predict noise $\varepsilon \rightarrow$ Score-based $\nabla_x \log p_t(x)$

1. Score function $s(x) = \nabla_x \log p(x)$ 指向 proba density 上升方向最快的向量场
 prob density

2. Langevin Dynamics

用 score 来采样 $\rightarrow$ Step Size

$$x_{k+1} = x_k + \frac{\eta}{2} \nabla_x \log p(x_k) + \sqrt{\eta} z_k, z_k \sim N(0, I)$$

物理直学 令 potential energy $U(x) = -\log p(x)$ 随机扰动力, 防止陷入局部 mode

带热 随机扰动的 GD

3. Score unknown

真实数据分布 $p(x)$ 未知 $\nabla_x \log p(x)$ 不能算

训练神经网络 $S_\theta(x) \approx \nabla_x \log p(x)$

理想 loss
$$L_{\text{score}} = E_{p(x)} [\|S_\theta(x) - \nabla_x \log p(x)\|^2]$$

target score unknown

$\Downarrow$
denoising score matching

4. Denoising Score Matching

对 clean data 加 noise $x_t = x_0 + \sqrt{t} \varepsilon, \varepsilon \sim N(0, I)$

$$\Rightarrow q(x_t | x_0) = \frac{1}{\sqrt{(2\pi t^2)^d}} \exp\left(-\frac{\|x_t - x_0\|^2}{2t^2}\right)$$

取 log

$$\log q(x_t | x_0) = -\frac{1}{2t^2} \|x_t - x_0\|^2 - \log((2\pi t^2)^d) + \text{const}$$

对 x_t 求梯度

$$\nabla_{x_t} \log q(x_t | x_0) = \frac{x_0 - x_t}{t^2}$$

unbiased estimate of the gradient of the marginal density

$$= \frac{-\varepsilon}{t} = -\frac{\varepsilon}{t}$$

$$E_{q(x_0 | x_t)} [\nabla_{x_t} \log q(x_t | x_0)] = \nabla_{x_t} \log p(x_t)$$

Denoising Score Matching

说明预测 score 与 noise 等价

$$S_\theta(x, t) \propto -\frac{1}{t} \varepsilon_\theta(x, t)$$

$$E_{q(x_0, x_t)} [\|S_\theta(x_t) - \nabla_{x_t} \log q(x_t | x_0)\|^2]$$

$$= E[\|S_\theta(x_t) + \frac{\varepsilon}{t}\|^2]$$

5. 为什么 NLP 需要 diffusion

自回归问题 { 左→右生成
前边错只能继续生成
global constraint 难以提前规划

Diffusion 优势 { whole sequence refinement
所有 token 同时被反复修改
早期确定 global semantics
后期修复 local syntax

5.1 核心困难 image 虽离散存储, 但可视为连续变量 $x \in \mathbb{R}^{H \times W \times C}$

token 不连续 $w_i \in V$ { token id 间无自然距离
cat dog 的 id 差距无语义意义
对 token id 加 Gaussian Noise 会落到无效 token

不能直接 $w_t = w_0 + \epsilon$

5.2 Discrete Diffusion: 用 transition matrix 替代 Gaussian Token

$$Q_t \in \mathbb{R}^{|V| \times |V|}$$

forward pass: $q(x_t | x_{t-1}) = (a_t(x_t; Q_t x_{t-1}))$

三种 corruption

Keep 保持原 token, 加噪

$$Q_t = (1 - \beta_t) I + \beta_t U$$

Mask / Absorbing diffusion

[mask] token 的 one-hot

$$Q_t = (1 - \beta_t) I + \beta_t e_{\text{mask}} e_{\text{mask}}^T$$

token 有概率替换为 mask

Swap / Uniform diffusion

$$Q_t = (1 - \beta_t) I + \frac{\beta_t}{|V| - 1} (I I^T - I)$$

被替换成 V 中其他 token

5.3 Discrete forward posterior

定义累积 transition $\bar{Q}_t = Q_1 Q_2 \dots Q_t$

前乘行向量 x

forward $q(x_t | x_{t-1}) = Q_t(x_{t-1}, x_t)$

$$q(x_t | x_0) = \bar{Q}_t(x_0, x_t)$$

若后接列向量 x

$$q(x_{t+1} | x_t, x_0) = \frac{Q_t(x_{t-1}, x_t) \bar{Q}_{t-1}(x_0, x_{t-1})}{\bar{Q}_t(x_0, x_t)}$$

reverse model

$$p_\theta(x_{t+1} | x_t, t) \approx q(x_{t+1} | x_t, x_0)$$

$$L_{KL} = E_{q(x_t, x_0)} [D_{KL}(q(x_{t+1} | x_t, x_0) \| p_\theta(x_{t+1} | x_t))]$$

对于 categorical distribution KL-minimization 等价于 cross entropy

$$L_{CE} = -E_{q(x_{t+1} | x_t, x_0)} [\log p_\theta(\hat{x}_{t+1} = x_{t+1} | x_t)]$$

但 $Q_t \in \mathbb{R}^{|V| \times |V|}$

$|V|$ 很大 $\Rightarrow$ continuous latent diffusion / Diffusion-LM

5.4 Diffusion LM

把 token $\rightarrow$ embedding space

$$x_0 = \text{Emb}(w) \in \mathbb{R}^{n \times d}$$

$$x_t = \sqrt{\alpha_t} x_0 + \sqrt{1 - \alpha_t} \epsilon$$

$$L_{\text{simple}}(x_0, \theta) = \sum_t E_{q(x_t | x_0)} [\| \mu_\theta(x_t, t) - \tilde{\mu}(x_t, x_0) \|^2]$$

rounding 到 token $p_\theta(w | x_0) = \prod_{i=1}^n p_\theta(w_i | x_i)$

$$L(x_0, w, \theta) = L_{\text{simple}}(x_0, \theta) - \log p_\theta(w | x_0)$$

clean embedding $\rightarrow$ 离散词

5.5 Controlled Generation

目标是生成满足控制条件 c 的文本 $p(x_{0:T} | c)$ c 是

在每一步 reverse diffusion 中 $p(x_{t+1} | x_t, c) \propto p(x_{t+1} | x_t) p(c | x_{t+1}, x_t)$

" 条件独立
 $p(c | x_{t+1})$

$$\text{取 log } \log p(x_{t+1} | x_t, c) = \log p(x_{t+1} | x_t) + \log p(c | x_{t+1}) + \text{const}$$

对 x_{t+1} 求梯度 $\nabla_{x_{t+1}} \log p(x_{t+1} | x_t, c) = \nabla_{x_{t+1}} \log p(x_{t+1} | x_t) + \nabla_{x_{t+1}} \log p(c | x_{t+1})$ ↓ classifier

5.6 Classifier Training

给定 labeled dataset $D = \{(w, c)\}$

- ① ~~训练~~ Sample (w, c)
- ② $w \mapsto x_0$
- ③ Sample t
- ④ $x_t = \sqrt{\alpha_t} x_0 + \sqrt{1 - \alpha_t} \epsilon$
- ⑤ classifier predict c from x_t and t
- ⑥ 最小化 cross entropy

$$L_{cls} = -\log p(c | x_t, t)$$

5.7 Block diffusion

普通 discrete diffusion { ① fixed-length gen
② 无法复用 KV-cache

在 AR 和 diffusion 之间折中

$$p_\theta(x) = \prod_{b=1}^B p_\theta(x_b | x_{<b})$$

每个 block 内部用 diffusion 并行去噪

Lec 25 Agentic System + planning | memory | tools | feedback

多步执行, 记忆历史, 调用工具, 反思错误, 与环境交互

4部分 LLM (brain) Memory Tools Environment

1. Context Engineering 优化送入 LLM 的上下文

limited context window

lost in the middle

不相关冲突的 retrieved info

其他形式图等不能序列化成 token

⇒ 检索, 压缩, 结构化

工具调用, 分阶段 reasoning

2. Memory

long term: 事实, 规则, 用户偏好; 过去事件, 交互轨迹, 技能, 习惯, 操作流程, 形成的行为模式

short term: 当前输入和对话, 中间状态, 计划, 变量

3. RAG 解决 outdated knowledge, LLM information source not transparent hallucination parametric memory 不可靠

① Indexing: build indexed database for documents

② Retrieval: retrieve top-k relevant chunks

③ Generation: insert chunks into LM context and generate answer

4. Reflection: CoT, LLM-as-a-judge

Prompt → Intermediate Reasoning

→ Answer

main LLM-output → judge LLM evaluates quality

→ score | critique | retry

变成 agentic system 的 loop 的一个模块

5. Tools

6. Collaboration / MCP 统一模型和工具的接口

model
↓
传统 $N \times M$ service

MCP ⇒ $N+M$