

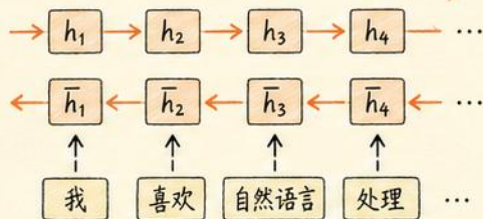
预训练语言模型

Word Embedding → BERT

技术演化图

2. ELMo

基于双向LSTM的上下文表示



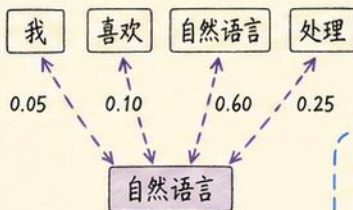
1. Word Embedding

将词映射为低维稠密向量

word	vector
我	[0.21, -0.35, ...]
喜欢	[-0.18, 0.42, ...]
自然语言处理	[0.33, 0.11, ...]
...	...

3. Attention

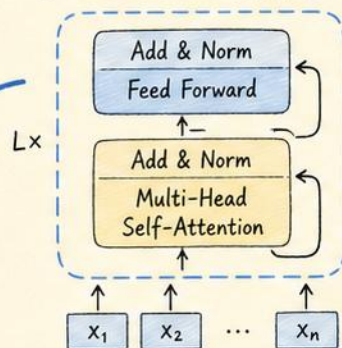
让模型关注更重要的词



加权聚合上下文信息

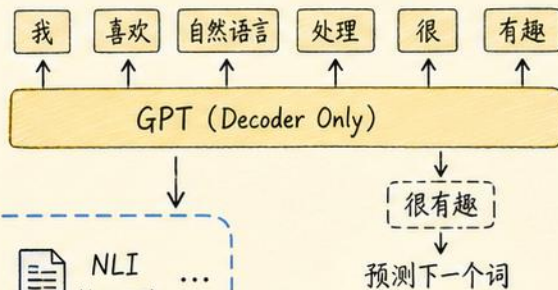
4. Transformer

多头自注意力 + 前馈网络



5. GPT

基于Transformer的自回归语言模型



BERT

Bidirectional Encoder Representations from Transformers

Fine-tune



分类
(如情感分析)



问答
(如SQuAD)



NLI
(如MNLI)

...

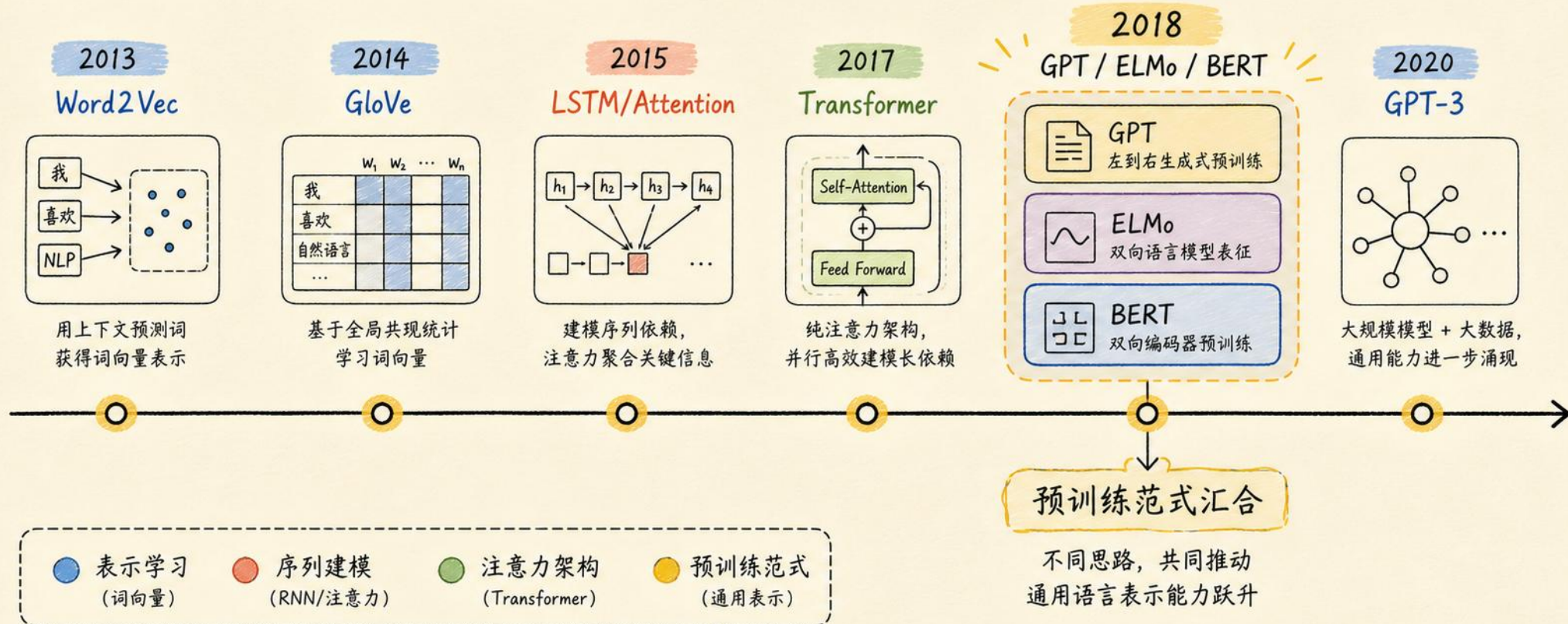


BERT = 双向理解 + 深度上下文 + 大规模预训练

→ 汇聚多种关键思想，成为NLP的里程碑模型

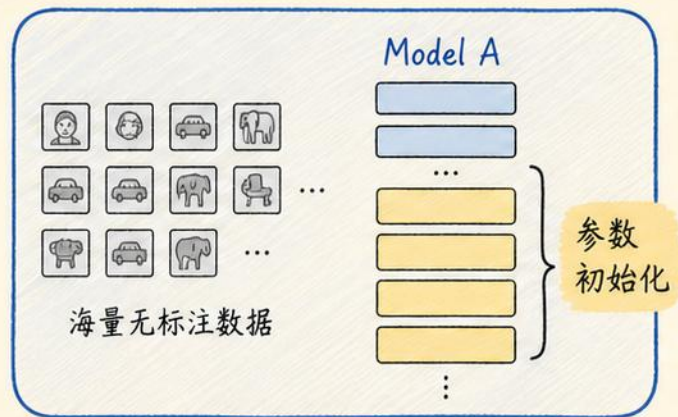
走向 BERT 的时间线

从词向量到预训练：关键技术的演进与汇合



预训练 = 参数迁移

大数据预训练



Frozen

冻结预训练参数，只训练上层

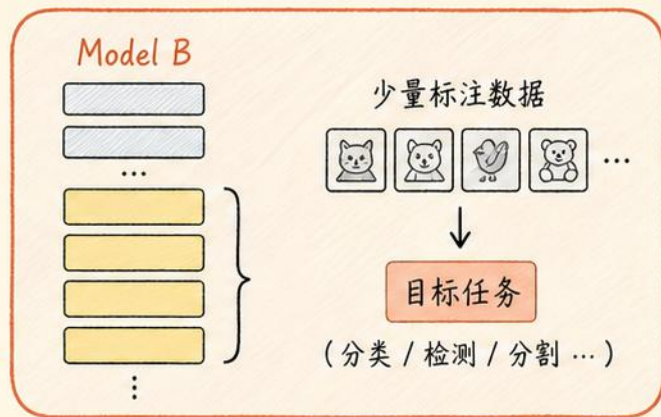
黄色：迁移的预训练参数（可复用）

蓝色：随机初始化的新参数

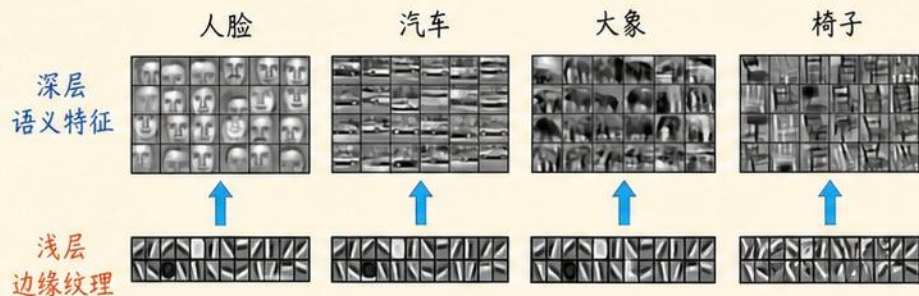
Fine-tuning

微调预训练参数，整体训练

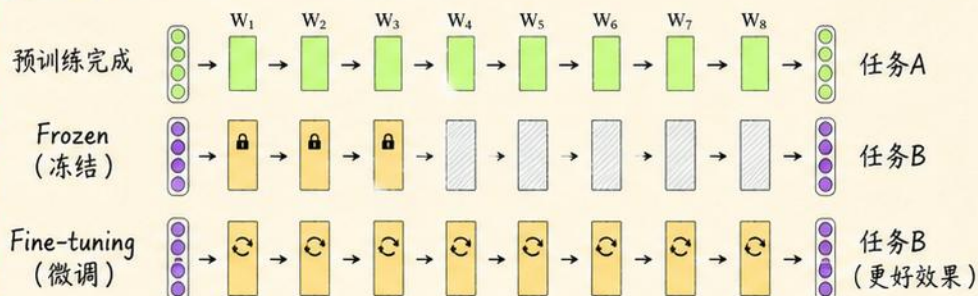
小数据任务



证据1：预训练学习通用特征（以自然图像为例）



证据2：冻结或微调的训练流程



黄色：预训练参数（可复用）

蓝色：新随机初始化参数

🔒：冻结不更新

🔄：参与更新

语言模型：预测下一个词

语言模型的概率基础

统计计数

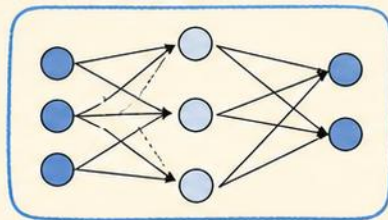
w_1	...	w_{i-1}	count
我	喜欢	自然语言	25
我	喜欢	机器学习	18
我	喜欢	深度学习	12
...	...	...	...

$$P(w_i | w_1 \dots w_{i-1})$$

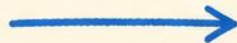
$$P(\text{next} | \text{context})$$

神经估计

稀疏、维度高



学习语义与上下文



Softmax



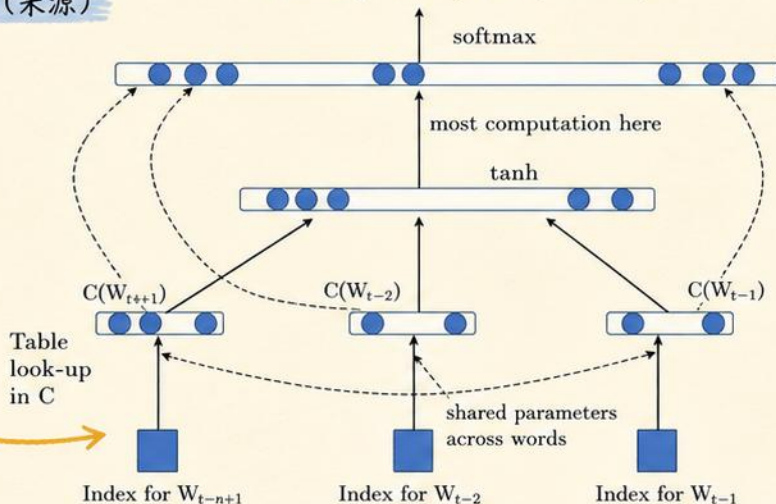
每个词的概率分布
(和为 1)

神经语言模型 (NNLM) 架构 (来源)

$$i\text{-th output} = P(W_t = i | \text{context})$$

1

输入：词的索引
通过共享的查表矩阵 C
得到词向量 (Embedding)



2

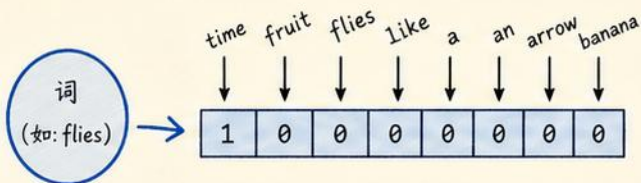
输出：词表上的概率分布
Softmax 将隐藏表示映射
为每个词的概率

W_t ：待预测的当前词
context：前面的上下文词
 C ：共享的词向量矩阵

第一代 NLP 预训练

学习一个共享的 Embedding 矩阵，作为下游任务的初始化参数。但向量是静态的，不懂上下文。⚠

1. One-hot (稀疏表示)

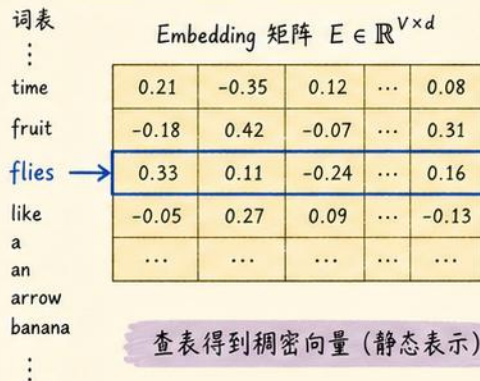


One-hot (稀疏, 高维)

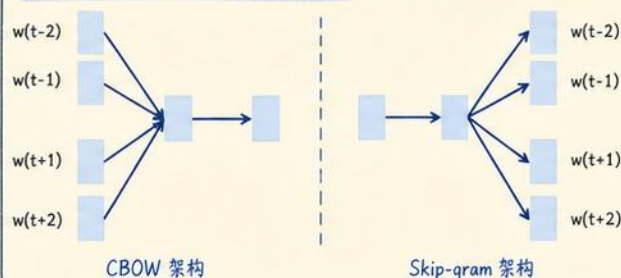
参考图: One-hot 表示

	time	fruit	flies	like	a	an	arrow	banana
1 time	0	0	0	0	0	0	0	0
1 fruit	0	1	0	0	0	0	0	0
1 flies	0	0	0	0	0	0	0	0
1 like	0	0	0	1	0	0	0	0
1 a	0	0	0	1	1	0	0	0
1 an	0	0	0	0	0	0	0	0
1 arrow	0	0	0	0	0	0	1	0
1 banana	0	0	0	0	0	0	1	1

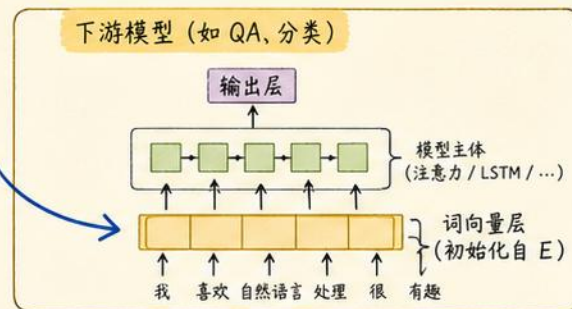
2. Embedding 矩阵 (查表)



参考图: Word2Vec (查表思想)

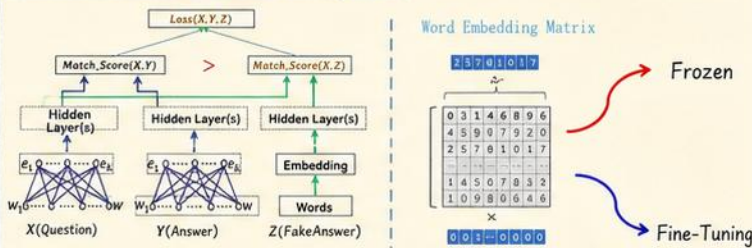


3. 下游初始化 (静态词向量)



静态词向量: 同一词一成不变

参考图: 下游任务中的 Embedding 复用



Word Embedding 是第一代 NLP 预训练: 通过在大规模语料上学习 Embedding 矩阵, 作为通用的初始化参数迁移到下游任务。但它是静态的, 无法根据上下文动态变化。

词义随上下文变化

context matters

一词多义

...very useful to protect **banks** or slopes from being washed away by river...

...the location because it was high, about 100 feet above the **bank** of river...

...The **bank** has plan to branch throughout the country...

...They throttled the watchman and robbed the **bank**...

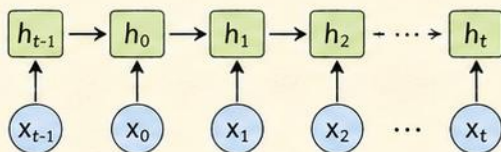
静态Word Embedding (同一向量)

bank

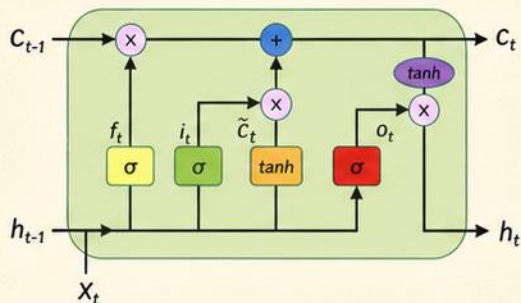
0.25 0.57 -0.13 ... 0.17

无法区分不同含义 ☹️

RNN / LSTM



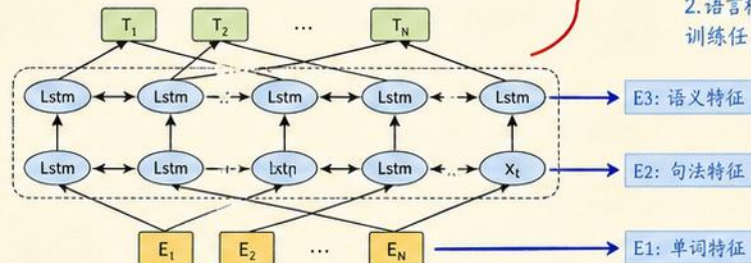
LSTM



顺序建模, 捕获上下文信息

ELMo

ELMo ELMo: Embeddings from Language Models



1. 双层双向RNN预训练
2. 语言模型作为训练任务

结合多层双向上下文, 得到动态表示

... washed away by river ...	0.82	-0.31	0.47	...	0.21
... above the bank of river ...	-0.63	0.71	-0.22	...	0.08
... bank has plan to branch ...	0.11	0.46	0.93	...	-0.35
... robbed the bank ...	-0.27	-0.89	0.31	...	0.66

同一词在不同上下文 → 不同表示 ✓



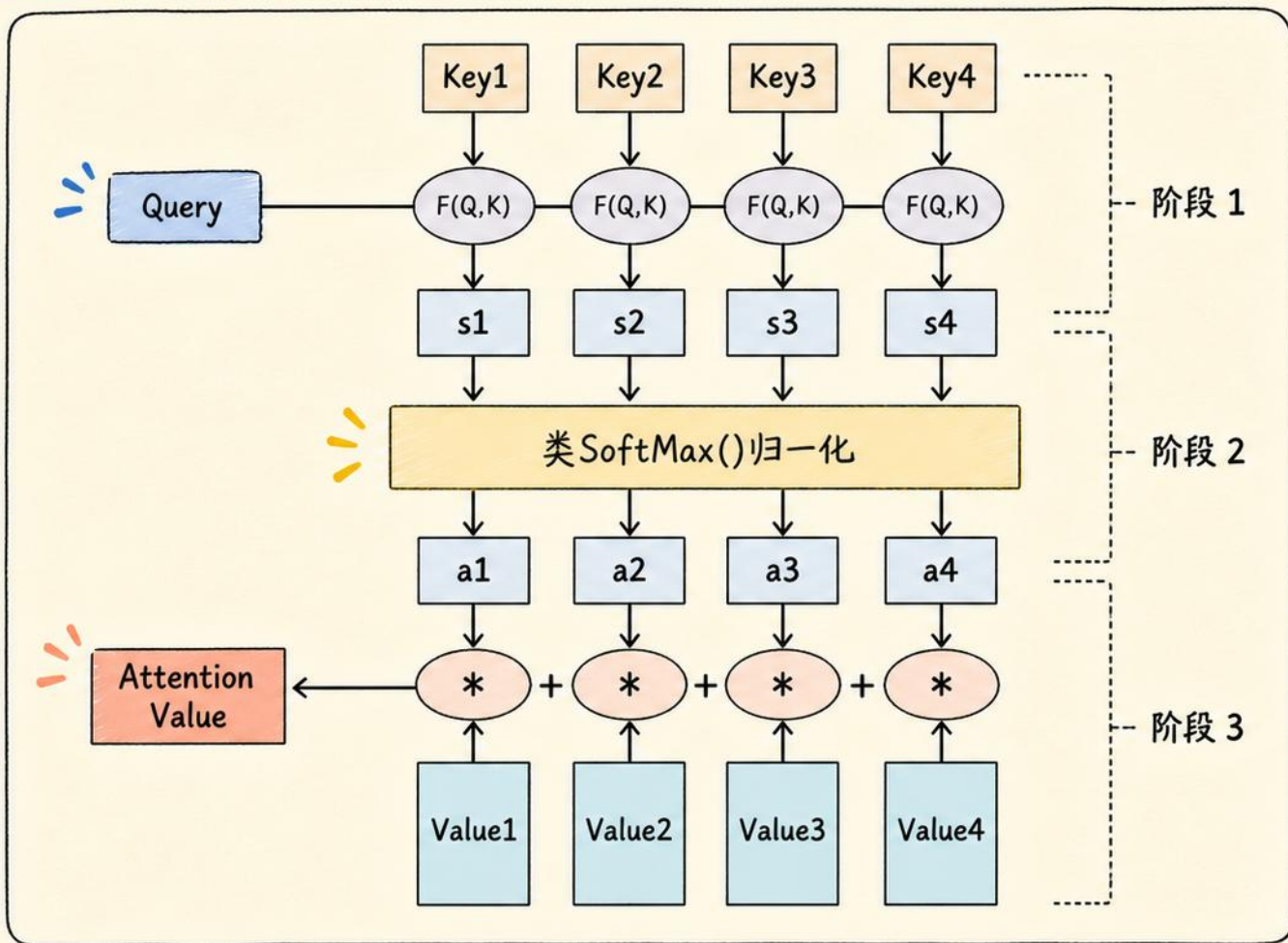
静态向量
(忽略上下文)

RNN / LSTM
(建模顺序, 感知上下文)

ELMo
(双向多层, 动态表示)

词义由上下文决定!

Attention = 相关性加权



1 Q-K 打分

Query 与每个 Key 计算相关性分数
得到 $s_1, s_2, s_3, s_4 \dots$

2 Softmax

对分数做 Softmax 归一化,
得到注意力权重 $a_1, a_2, a_3, a_4 \dots$

3 加权求和

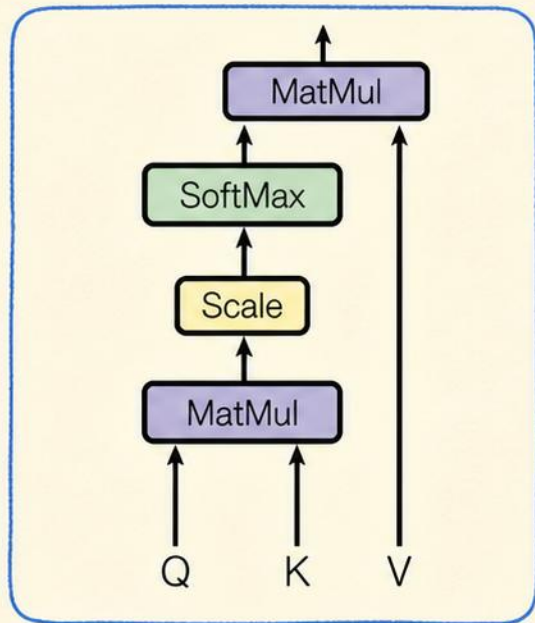
用权重乘以对应的 Value,
再把结果加和, 得到输出



$$\text{softmax}(QK^T)V$$

Attention 家族

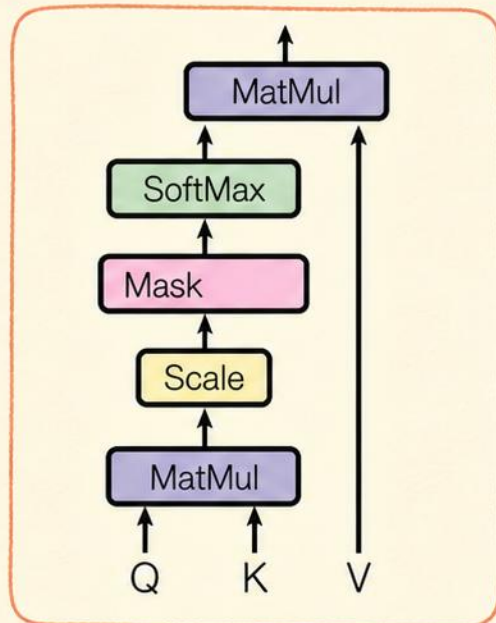
Self-Attention



全局互看

任意位置可关注序列中所有位置，捕捉全局依赖关系。

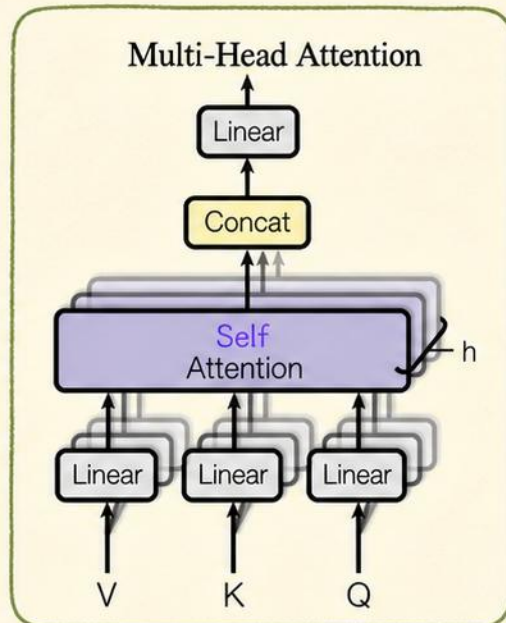
Masked



遮住未来

通过因果 Mask 防止看到未来信息，保障自回归生成的正确性。

Multi-head



多头并行

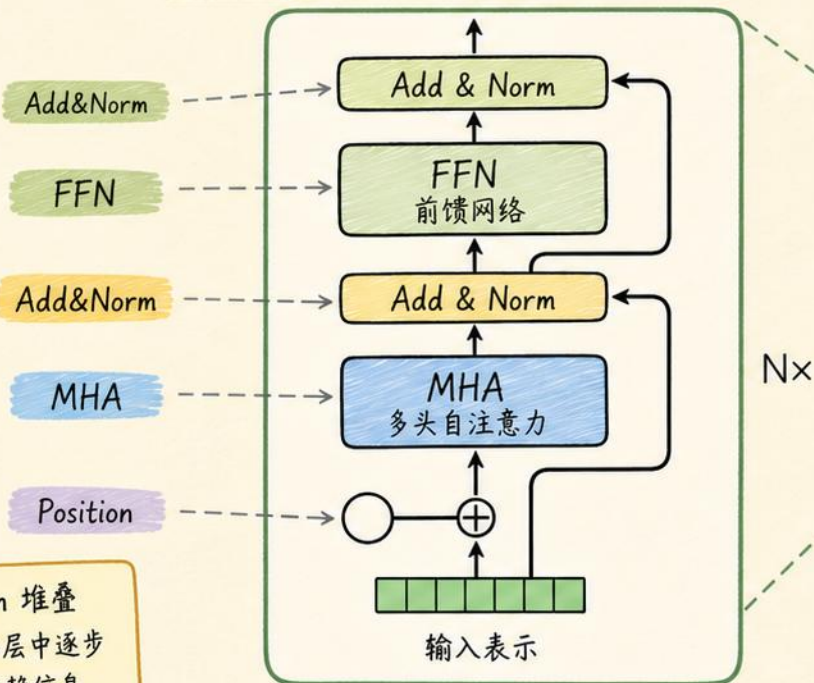
多个注意力头并行学习不同子空间，增强表示能力与建模灵活性。

Transformer 总装

核心理念：

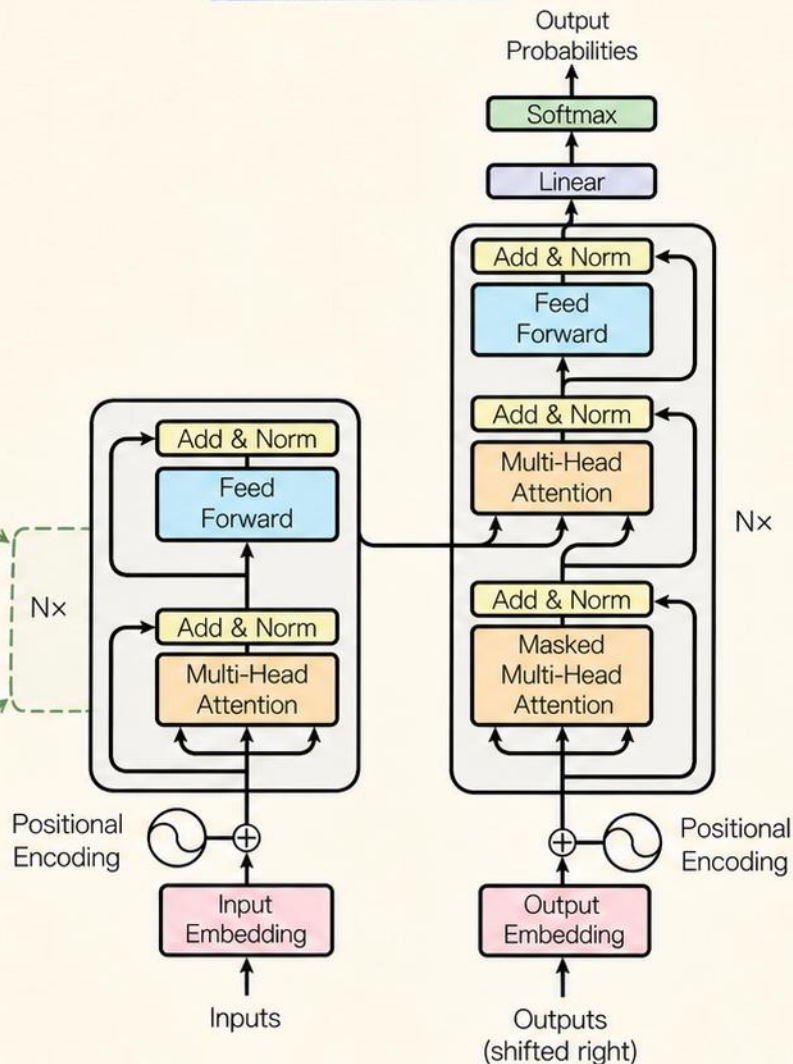
用注意力(信息交互) + 前馈网络(特征变换)
在残差连接与归一化下反复堆叠，构成强大的序列建模器。

Encoder 单层结构 (放大看)



Attention 堆叠
让模型在多层中逐步
聚合长程依赖信息。

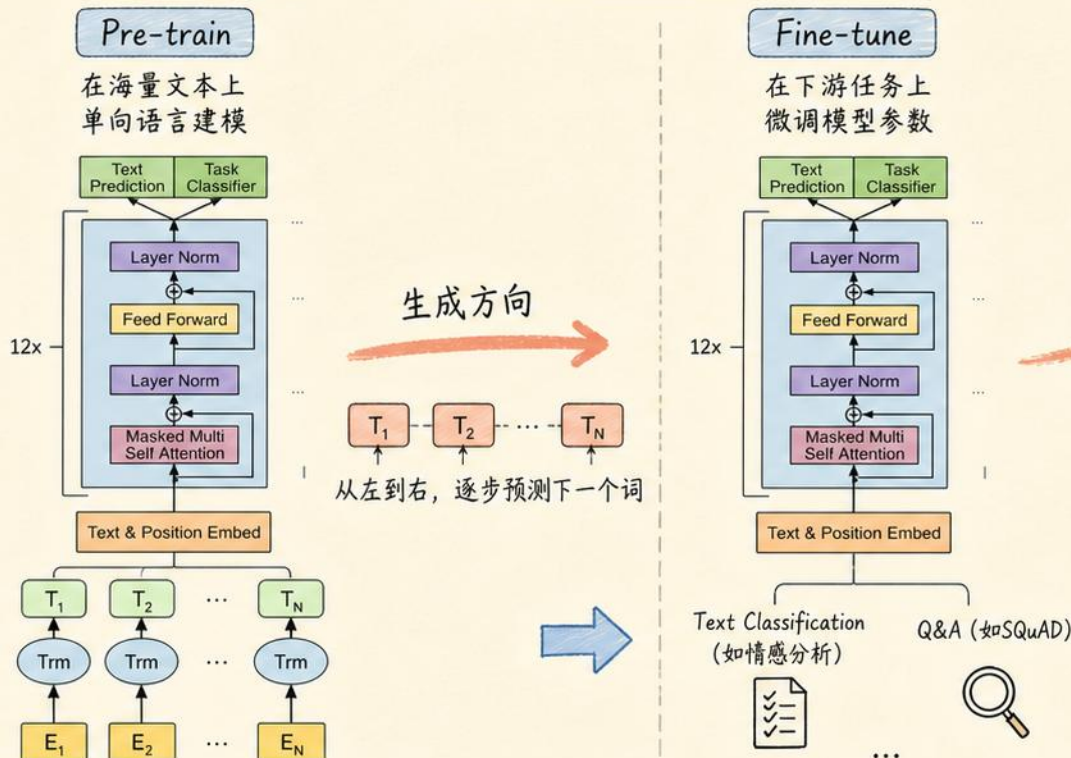
Transformer 总体结构



GPT: 单向 LM 路线

单向 LM

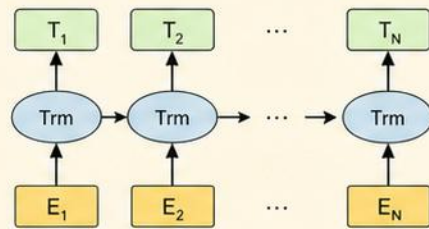
Transformer Decoder



GPT vs BERT: 结构对比

GPT (Decoder-like)

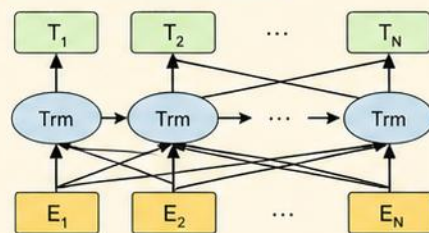
单向生成



- ✓ 单向 LM
- ✓ 从左到右
预测下一个词
- ✓ 生成式

BERT (Encoder-like)

双向理解



- ✓ 双向理解
- ✓ 利用上下文
- ✓ 更擅长理解
(非生成)



GPT 路线: 先在大规模文本上做单向语言建模 (Pre-train), 再在具体任务上微调 (Fine-tune), 以生成式方式预测下一个词。

BERT: 双向语义理解

MLM

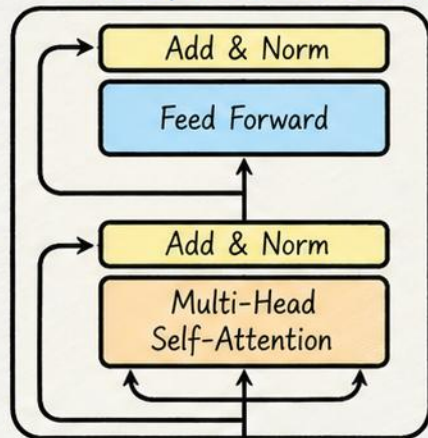
掩码语言模型

预测被遮盖的词

my dog is cute
↓
my [MASK] is cute
↓
预测: dog

$L \times$

Transformer Encoder



深层双向特征

NSP

下一句预测

判断句子B是否为A的下一句

句子 A my dog is cute
句子 B he likes dogs
↓
IsNext ?
(Yes / No)

Token

$E_{[CLS]}$ E_{my} E_{dog} E_{is} E_{cute} $E_{[SEP]}$ E_{he} E_{likes} E_{play} $E_{\#ing}$ $E_{[SEP]}$

Segment

E_A E_A E_A E_A E_A E_A E_B E_B E_B E_B E_B

Position

E_0 E_1 E_2 E_3 E_4 E_5 E_6 E_7 E_8 E_9 E_{10}

$[CLS]$ my dog is cute $[SEP]$ he likes play $\#ing$ $[SEP]$

句子 A

句子 B

BERT 作为通用底座

集大成者

统一的预训练模型，
适配多种下游任务

通用预训练底座

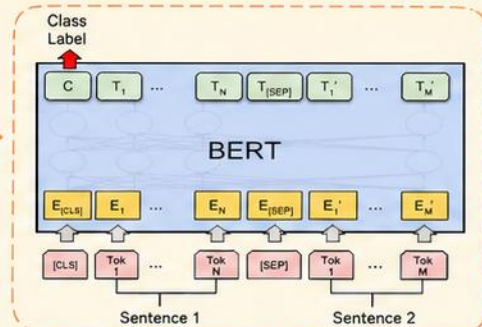
双向理解 + 深度上下文
大规模无监督预训练

BERT

Bidirectional Encoder
Representations from
Transformers

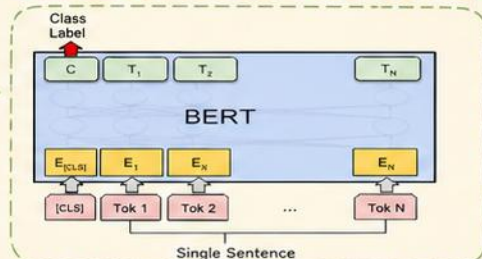
句对分类

判断句对关系



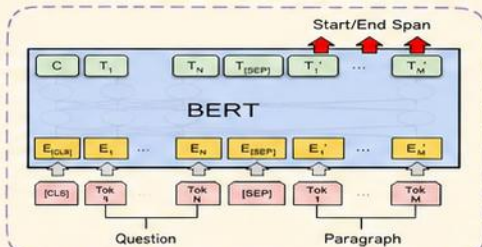
单句分类

理解整句语义



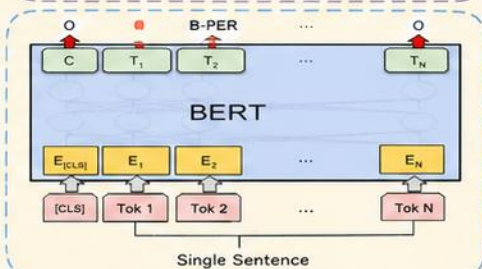
问答抽取

定位答案片段

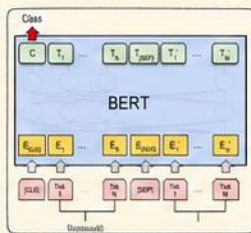


序列标注

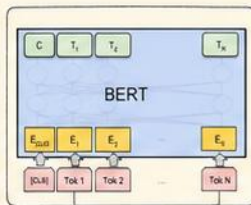
逐词标注标签



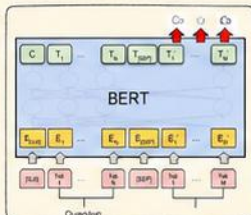
来源图示 (参考)



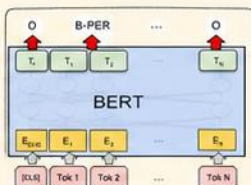
(a) 句对分类



(b) 单句分类



(c) 问答抽取



(d) 序列标注



1. 统一输入

相同的词/位置表示与编码器
适用于多种任务格式



2. 简单任务头

在预训练表示上添加轻量结构
快速适配不同下游任务



3. 两阶段训练

先大规模无监督预训练
再小规模有监督微调