

Case 1

Boundary-First Preference

Local concept-clarification cues are generalized into a reusable evidence-boundary guardrail.

Job 2

容易误读的点：概念澄清

- “标志”是组织原则而非单一靶点列表；同一标志可由多条路径实现。
- “使能特征”非“标志能力”，其作用是加速/促成多标志的获得。
- “新兴标志”意指潜在普遍性与概念扩展，非既定普适结论。
- 肿瘤微环境并非被动背景，而是积极参与者；忽视其作用可能导致对干预效果的误判。

Concept-clarification cue

Job 4

证据局限与易误读澄清（基于原文）

证据局限（要点）

- 框架属性：为机制与结构化理解提供“地图”，非定量风险预测工具。
- 情境依赖：微环境互作与子回路串扰高度依赖组织情境与进展阶段。
- EMT 相关：EMT 可部分发生且可逆；不同肿瘤类型/阶段表现不一。
- 侵袭模式：集体/变形虫式侵袭的普遍性与分子机理仍未解决。
- 靶向示例：治疗图谱中的药物仅为示例；存在差异、演化与耐受可能。

易误读点/边界提示（对照）

易误读点	边界提示
将“六大标志”图当作定量风险评分或诊疗指南	概念性框架，用于组织理解
默认 EMT 在所有病例中普遍且完全发生	原文指出部分 EMT 与可逆性（MET）
把集体/变形虫式侵袭视为已明确确定的普适机制	作者明确其普遍性与机制仍未解决
把微环境与信号网络视为静态、孤立的通路	原文强调动态重构与广泛串扰

Evidence-boundary table

Job 6

证据局限与边界：先澄清，再传播

- 文献属性：观点性综述/框架性总结，非流行病学发生率/疗效对比
- 图示属性：概念化表达，作者指出子电路之间存在大量交叉，非精确路径图
- 角色不确定性：肿瘤相关淋巴作用仍不清；周细胞覆盖差→血行播散属假说，尚待证实
- 推广边界：图6药物为“示例”，并非完整或通用目录

易误读点	边界提示
将六大特征当作每一肿瘤的必备清单	框架用于“组织化理解”，实际表现受细胞来源、微环境与进展阶段影响（文中描述）
把图2视为定量通路图	作者强调为概念化“集成电路”，存在大量交叉与外界输入
认为淋巴管一定促进肿瘤生长	文中称其作用“仍不清”；边缘淋巴生长常见，但功能界定有限
认为周细胞覆盖差必然导致转移	标注为“假说，尚待证实”，不可视作确证结论
将图6药物列表视为标准治疗目录	图标注“仅示例+深管线”，不可等同于通用或本院目录

来源：Hanahan & Weinberg 2011（正文与图注表述）。本页仅澄清证据边界，避免误读。

Inherited guardrail

Case 2

Action-Closure Preference

Next-step questions are generalized into a Problem-Owner-Timeline table.

Job 1

Closing | Concise Conclusions and Next Questions

Framework

- Core takeaway: interactions among hallmarks and enabling traits organize the review.
- Evidence level: mainly mechanistic synthesis; cross-tumor generalization should remain cautious.

Therapy Map

- Do not introduce extra-mural metrics or thresholds beyond the paper.
- Retain source traceability to the original figures and wording only.

Microenvironment

- Tumor microenvironment and multicellular interactions are essential contextual elements.
- Boundary: the paper does not provide quantitative effect sizes or a unified classification standard.

Data Needs

- Quantitative validation is needed for interaction hypotheses and comparative claims.
- Heterogeneity and reproducible pathways across tumor and organ contexts remain open questions.

Question cards

Job 3

Action Table (Source-Bound)

Evidence boundary: no external facts; do not infer risk or efficacy data.

Problem / Question	Owner	Timeline
Add evidence boundary / emerging labels under relevant page titles and retrofit existing pages	Content team; medical review	Before this release
Turn microenvironment and intracellular-circuit content into paired pages with cross-references	Layout / content team	Next release cycle
Standardize terminology for emerging vs enabling characteristics and replace inconsistent labels	Content team	Before this release
Add a training comparison page distinguishing hallmark, emerging hallmark, and enabling characteristic	Training owner; medical review	Before next training
Add a footnote that the integrated-circuit figure is a simplified schematic with crosslink	Content team	Before this release
Add a note on pages using drug-category figures: conceptual example, not efficacy or risk data	Content team	Before this release

Owner/timeline table

Job 6

Next Questions and Action Checklist | Quality and Safety Closure

Question to Align	Owner	Timeline
Standardize terminology: six hallmarks, enabling traits, emerging hallmarks, and microenvironment	Education and QI group	Publish a terminology sheet within 2 weeks
Add a dedicated evidence-boundary page to internal training materials	Education and QI group	Release before the next training session
Build a local high-resolution asset library for Figures 1, 3, 4, 5, and 6 with sources and legend highlights	Information team and education team	Archive within 2 weeks
Add unresolved or hypothesis labels to pages involving lymphatic or pericyte content	Responsible deck owners	Revise current decks within this month
Standardize a fixed microenvironment-role page citing Figures 4 and 5 in all related decks	Education team and relevant disciplines	Complete coverage this quarter

Note: these are internal quality and communication workflow actions grounded in the attached framework; they are not clinical treatment recommendations.

Execution closure: every action should have a named owner and a due date, be tracked weekly, and be checked in the next review cycle. Source: Hanahan and Weinberg 2011 (Cell).

Case 3

Module-Interface Preference

Architecture content is generalized into a Module-Input-Output-Responsibility view with explicit component boundaries

Job 1

整体架构：模块与数据流

模块	输入/输出	职责	
编码器 (Encoder, >N)	输入：词嵌入与位置编码叠加；输出：上下文编码 Z	由多头自注意力 + 逐位置前馈网络构成，配合残差与层归一化，抽取源序列全局依赖	→
解码器 (Decoder, >N)	输入：移位后的目标嵌入与编码器输出 Z；输出：解码表示	包含掩蔽自注意力与编码器-解码器注意力，以及逐位置前馈网络；通过掩蔽保证自回归	→
多头注意力 (Attention)	输入：Q、K、V；自注意力时 Q=K=V=当前层表示；跨注意力时 Q=解码器表示，K=V=Z；输出：聚合后的上下文	并行多头在无循环/卷积下建模长程依赖，聚合全局信息	→
前馈网络 (Feed-Forward, FFN)	输入：每个位置的隐藏表示；输出：变换后的隐藏表示	两层逐位置前馈（含 ReLU），各位置相同且独立应用，增强非线性表达	→

2 / 10

Module/IO Organization

Job 2

模块清单（IO 与职责）

模块	输入	输出	责任
词嵌入 + 位置编码	token 序列	d_model 维表示	将离散 token 映射到向量并注入位置信息
多头自注意力（Encoder）	序列表示	上下文混合后的表示	建模同层全局依赖
逐位置前馈网络	序列表示	非线性变换后的表示	提升特征表达、逐位置独立
编码器层（堆叠）	序列表示	编码后的序列表示	叠加注意力与前馈
掩蔽自注意力（Decoder）	部分目标前缀	因果约束的表示	防止偷看未来
编码器-解码器注意力	解码表示 + 编码表示	融合源端信息的表示	跨序列对齐与条件化
解码器层（堆叠）	解码表示	目标侧上下文表示	整合自注意与跨注意
输出投影 + softmax	解码表示	词分布	生成下一词概率

IO-Responsibility Schema

Job 6

架构总览（模块边界与职责）要点

- Encoder-Decoder 双栈 + 多头自注意力 + 前馈网络，通过残差与层归一化稳定训练（§3.1）。
- 关键超参（Base）：N=6；d_model=512；d_ff=2048；h=8；d_k=d_v=64；位置编码：正弦。

模块 IO 与职责

- Encoder Layer：输入=上一层隐表示；输出=同维隐表示；职责=跨位置建模（自注意力）+ 按位非线性（前馈）。
- Decoder Layer：输入=先前目标前缀、编码器输出；输出=下一步预测上下文表示；职责=在因果性下融合源与目标前缀。

依据：§3.1–3.3

2 / 10

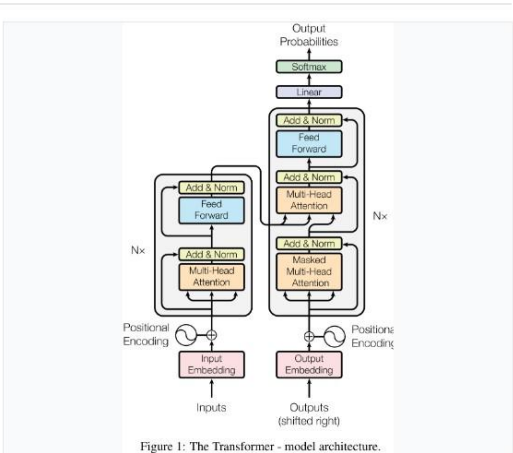


Figure 1: The Transformer - model architecture.

Case 4

Implementation-Checklist

Mechanism explanations become verifiable implementation checks.

Job 3

实现风险与检查清单（仅据论文）

- 自回归屏蔽：确保解码端 masked 自注意力严格屏蔽未来位置信息
- 复杂度与内存：自注意力对长序列呈 $O(n^2)$ 资源消耗（论文 Table 1）
- 注意力缩放与数值：按维度进行缩放，保持数值稳定
- 位置编码注入：按论文加入（正弦）位置编码，并与嵌入同维度相加
- 残差与归一化：各子层后应用残差连接与 LayerNorm，维度保持一致

边界声明：以上仅复述论文内机制与证据，不引入论文外假设或生产级推断。

Risk Checks

Job 5

实现检查清单（实现/复现必查）

仅基于原文机制与实验设置，面向工程落地的可验证要点

Tensor 形状

- 嵌入与位置编码输出形状一致：(n, d_model)
- Q/K/V 线性映射每头维度 d_k, d_v；Self-Attn 输入/输出 (n, d_model)
- 多头拼接后经 Wo 投影回 (n, d_model)；FFN 按位置独立

并行与掩码

- 层内注意力在序列维度并行；Encoder/Decoder 层可并行堆叠
- 训练时 Decoder 借助自回归 Mask 并行计算所有位置
- Padding Mask 需在 Self/Cross-Attn 全链路保持一致

依赖与顺序

- 位置编码（正弦/学习式）必备，遵循论文设置
- 各子层采用 Add & Norm（残差 + 层归一化）
- Encoder 堆叠 → Decoder 堆叠（含自注意力与交叉注意力）

参与与训练

- h、d_model、d_ff、d_k/d_v 等超参一致性
- Dropout、Label Smoothing、学习率 Warmup 按论文配置
- 源/目标词嵌入与输出层对齐（权重共享时保持一致）

实现风险

- 自注意力开销 $\sim O(n^2 \cdot d)$ ；长序列显存/时间需评估
- 掩码实现错误可能泄露未来信息（解码侧）
- 头维/拼接/投影维度不一致导致运行错误（多头 → 合并 → Wo）
- LayerNorm/残差位置不当影响训练稳定性

Checklist Schema

Job 6

实施检查清单（按论文可核验项）

以可验证项对齐实现（形状/并行/复杂度/依赖假设）

检查项	期望设置	证据出处
层数与维度	N=6；d_model=512；d_ff=2048	§3.1, §3.3
多头参数	h=8；d_k=d_v=64（=512/8）	§3.2.2
张量形状	$Q, K, V \in \mathbb{R}^{n \times d_k}$ （按批量矩阵计算）	§3.2.1
并行策略	样本内按位置并行（自注意力）；按头并行（多头）	§1, §3.2.2
因果依赖	解码器自注意力使用遮罩，禁止看未来	§3.2.3
位置编码	正弦位置编码（Base）；可替换为学习式（效果近似）	§3.5, 表 3（E）
复杂度量级	自注意力 $O(n^2 \cdot d)$ ；卷积 $O(k \cdot n \cdot d)$ ；RNN $O(n \cdot d^2)$	表 1
优化与正则	Adam($\beta_1=0.9, \beta_2=0.98, \epsilon=1e-9$)、warmup=4000、dropout=0.1、label smoothing=0.1	§5.3–5.4

5 / 10

Stable execution closure

Boundary-Centered Overview

Reusable Checklist