

汽车行业深度报告

AI 系列深度 18 期：生成智能如何引入 Transformer？

增持（维持）

2026 年 08 月 07 日

证券分析师 黄细里

执业证书：S0600520010001

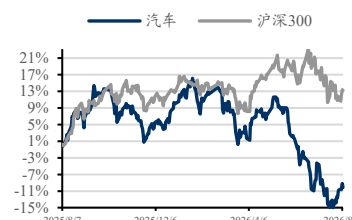
021-60199793

huangxl@dwzq.com.cn

投资要点

- **生成式视觉的本质，是学习图像数据背后的分布规律，并在文本条件下生成新内容。**本轮生成式视觉的核心范式迁移发生在 GAN 与 Diffusion 之间；扩散以更稳定的加噪—去噪学习确立生成主干，其潜空间建模与交叉注意力机制，也为 Transformer 的介入预留了接口。不同于语言领域中 Transformer 对 RNN 的快速替代，其在视觉生成任务中沿序列生成、条件编码、骨干替换与多模态扩展四个环节渐进渗透，
- **以代表论文为锚，我们将 2013 年以来生成式视觉的发展划分为五个阶段：**VAE 与 GAN 奠基（2013—2015 年）、GAN 主导高真实感生成（2016—2019 年）、扩散复兴与理论统一（2020—2021 年）、文生图爆发（2022 年）、DiT 与视频/多模态扩展（2023 年至今）。
- **Transformer 进入生成式视觉领域，经历了由局部模块引入到核心架构替换的渐进式演进。**早期阶段，Transformer 主要作为自回归生成器，将文本与图像表示统一为离散 token 序列，实现视觉内容的序列化建模；2021 年 DALL·E 将文本 token 与图像 token 联合建模，证明图像可被 token 化、并以类似语言模型的方式生成。随后，随着 CLIP 等视觉语言模型的发展，Transformer 进一步承担文本语义编码与条件信息注入功能，通过跨模态对齐提升对文本指令的理解能力；Imagen 与 Stable Diffusion 则推动大语言模型编码器与扩散模型相结合。
- **随着视觉数据 token 化与大规模训练范式的发展，Transformer 开始进入生成模型的核心模块。**2023 年 DiT 以 Transformer 替代传统扩散模型中的 U-Net 骨干网络，在潜空间图像 patch 上建模，验证了视觉生成同样具备规模化扩展潜力；不过其改变的是网络骨干，而非扩散去噪范式本身。进一步来看，Transformer 正逐步成为视频与多模态生成系统的统一建模框架，推动生成模型由单图创作向连续时空与多模态交互演进。整体而言，相比语言领域对 RNN 架构的全面替代，Transformer 在视觉生成领域更多体现为渐进式渗透：从生成器、条件编码器到核心骨干网络逐步深化。
- **我们认为，数字 AI 底座模型发展路径已逐渐清晰、价值在于确定性，但其难以实现物理世界的建模闭环；我们认为物理 AI 将成为物理世界的 AI 解决方案，从而在 AI 的当前发展阶段成为增量更大的方向，智能驾驶作为最先跑通闭环的代表场景，我们认为应重点关注。**
- **风险提示：**技术迭代与产品化落地不及预期的风险；视频与多模态生成的时间一致性与可控性不及预期的风险；大模型厂商 ARR 增速放缓、云服务商资本开支不及预期的风险；智能驾驶及机器人商业化落地不及预期的风险。

行业走势



相关研究

《AI 系列深度 17 期：视觉智能为何引入 Transformer？》

2026-08-06

《AI 系列深度 16 期：语言智能为何从 RNN 转向 Transformer？》

2026-08-06

内容目录

1. 复盘生成式视觉：从学习数据分布到自然语言生成内容	4
1.1. 生成式任务的本质与起点	4
1.2. 五阶段框架：从 GAN 主导到 Diffusion 接棒	4
1.3. 标志性成果与学术演进	4
2. GAN 路线：对抗生成的原理、巅峰与局限	6
2.1. 对抗生成的基本原理	6
2.2. 从 DCGAN 到 StyleGAN 的演进	7
2.3. 结构性局限	7
3. Diffusion 路线：去噪生成如何完成范式替代	8
3.1. 去噪生成的基本原理	8
3.2. Diffusion 对 GAN 的范式替代	9
3.3. 效率与算力局限	10
4. Transformer 如何一步步融入图像生成	10
4.1. 环节一：自回归序列生成器	11
4.2. 环节二：文本条件编码与注入	11
4.3. 环节三：扩散骨干替换（DiT）	11
4.4. 环节四：多模态统一与时空扩展	12
4.5. 核心结论：Diffusion 为主角、Transformer 为工程化骨干	13
5. 风险提示	13

图表目录

图 1: 生成式视觉学术派系与技术演进时间轴	6
图 2: GAN 对抗训练原理示意	7
图 3: 扩散模型前向加噪—反向去噪过程及潜空间扩散示意	9
图 4: Transformer 融入图像生成历程时间轴	11
图 5: DiT 与 U-Net 扩散骨干结构对比示意	12
表 1: 生成式视觉发展历程梳理	4
表 2: 生成式视觉标志性成果梳理	5
表 3: GAN 路线代表论文与核心贡献	7
表 4: Diffusion 路线关键论文与演进	10
表 5: Transformer 在生成式视觉任务中的四个介入环节	13

1. 复盘生成式视觉：从学习数据分布到自然语言生成内容

1.1. 生成式任务的本质与起点

与图像分类、检测、分割等判别式任务不同，生成式任务的目标不是给定图像并输出标签，而是学习图像数据背后的生成规律，并在文本等条件约束下生成此前不存在的新图像。换言之，模型不是在“识别图像”，而是在学习大量图像背后的结构、风格与语义分布，并据此生成新的视觉内容。其难点在于：自然图像处于极高维空间，真实分布无法显式写出，如何用神经网络近似这一分布并高效采样，构成十余年来生成式视觉研究的主线。

2013—2015 年是这一主线由传统概率图模型转向深度神经网络端到端训练的起点。2013 年，Kingma 与 Welling 提出变分自编码器（VAE），以变分推断与重参数化技巧训练潜变量生成模型，为生成模型提供了清晰的“编码—潜变量—解码”框架。

1.2. 五阶段框架：从 GAN 主导到 Diffusion 接棒

我们梳理了 2013 年以来生成式视觉领域代表性成果，依据主导生成范式的更替，将其发展划分为五个阶段：从奠基期多条路线并立，到 GAN 主导高真实感生成，再到扩散模型复兴并完成范式替代，继而自然语言成为生成入口，最后骨干网络 Transformer 化、生成对象向视频与多模态扩展。

表1：生成式视觉发展历程梳理

阶段	时间	阶段名称	主导范式	代表成果	阶段定位
1	2013—2015	深度生成模型奠基期	VAE/GAN/ 早期 Diffusion/自回归	VAE、GAN、早期 Diffusion、DCGAN	生成模型开始由深度神经网络端到端训练
2	2016—2019	GAN 主导的逼真图像生成期	对抗生成	pix2pix、CycleGAN、ProGAN、StyleGAN	图像生成第一次实现高真实感
3	2020—2021	Diffusion 复兴与范式替代期	去噪扩散生成	DDPM、DDIM、Score-based、Guided Diffusion	生成范式从对抗训练转向逐步去噪
4	2022	文本驱动图像生成爆发期	Text-to-Image Diffusion	DALL·E2、Imagen、Latent Diffusion、Stable Diffusion	自然语言成为图像生成入口
5	2023 至今	可控生成、Diffusion Scaling 与视频/世界模拟期	Control+DiT+Video Diffusion	ControlNet、SDXL、DiT、Sora、Lumiere、Veo	从单图创作走向可控视觉基础模型与连续时空生成

数据来源：《Generative Adversarial Networks》，《Denoising Diffusion Probabilistic Models》等其他文献，东吴证券研究所

1.3. 标志性成果与学术演进

我们进一步梳理了生成式视觉发展中的标志性成果，这些成果构成各阶段范式更替的关键节点。VAE 与 GAN 确立现代深度生成的两条基本路线；StyleGAN 代表 GAN 路线在高真实感人脸生成上的巅峰；DDPM 与 Score-based SDE 标志扩散模型复兴与理论统一；CLIP 提供语言与视觉对齐基础；Latent Diffusion 与 Stable Diffusion 把扩散过程移入潜空间并开源化，推动文生图大规模落地；DiT 则把扩散骨干迁移到 Transformer。

表2：生成式视觉标志性成果梳理

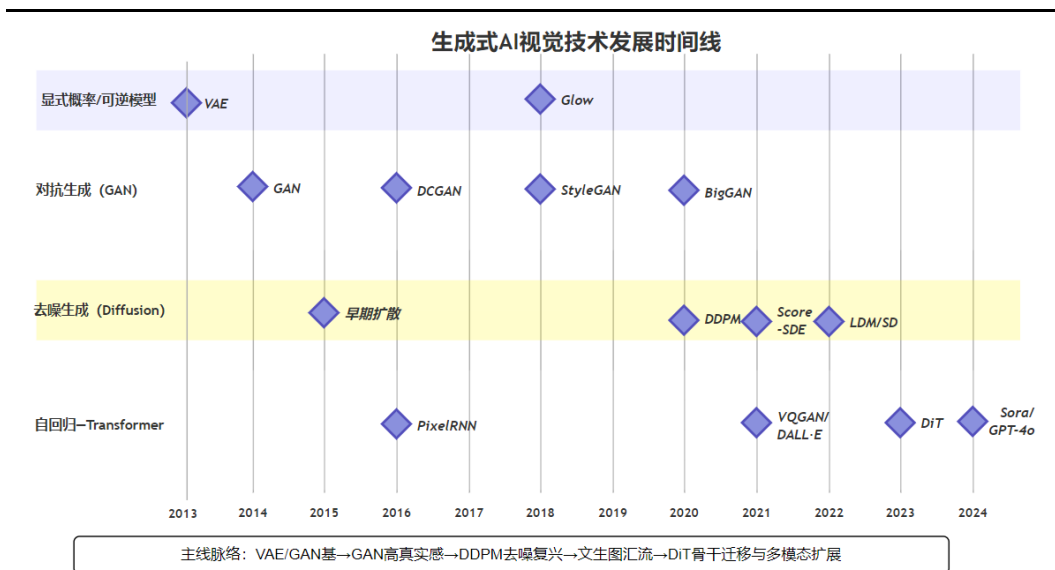
时间	论文/成果	代表人物/机构	核心贡献	阶段定位
2013/14 年	VAE (Auto-Encoding Variational Bayes)	Kingma,Welling	重参数化技巧，连续潜空间生成	现代深度生成起点
2014 年	GAN (Generative Adversarial Nets)	Goodfellow et al.	对抗生成框架，隐式建模分布	视觉真实感路线起点
2018/19 年	StyleGAN	NVIDIA	基于风格的生成器，可控潜空间	GAN 时代巅峰
2020 年	DDPM	Ho,Jain,Abbeel	去噪扩散复兴，高质量图像合成	Diffusion 成为主线
2021 年	Score-Based Modeling through SDEs	Song et al.	统一 score 与 diffusion 框架	Diffusion 理论统一
2021 年	CLIP	OpenAI	图文对比学习，语言控制视觉概念	文本—视觉对齐
2022 年	Latent Diffusion Models	Rombach et al.	潜空间 diffusion+交叉注意力	Stable Diffusion 基础
2022 年	Stable Diffusion	Stability AI/CompVis	开源化、消费级 GPU 可运行	生成式视觉应用爆发
2023 年	DiT	Peebles,Xie	Transformer 替代 U-Net，扩散模型 scaling	生成 backbone 迁移

数据来源：《Auto-Encoding Variational Bayes》，《Scalable Diffusion Models with Transformers》等其他文献，东吴证券研究所

从技术路线演进看，生成式视觉大致沿四条线索展开：一是以 VAE、Glow 为代表的显式概率/可逆模型线；二是以 GAN 及其变体为代表的对抗生成线；三是以早期扩散、DDPM、Score-based 模型为代表的去噪生成线；四是以 PixelRNN/CNN、VQGAN、DALL·E1 为代表的自回归序列建模线。

2021—2022 年前后，技术路线出现汇合：CLIP 提供文图对齐，自回归线贡献图像 token 化思想，二者与去噪生成线结合，催生以文本为入口、以扩散为主干的文生图体系；2023 年后，自回归线所依托的 Transformer 进一步进入扩散主干（DiT），并向视频与多模态扩展。

图1：生成式视觉学术派系与技术演进时间轴



数据来源：《Generative Adversarial Networks》，《Denoising Diffusion Probabilistic Models》等其他文献，东吴证券研究所绘制

从标志性成果的时间分布看，GAN路线代表性成果集中于2014—2019年，并在2018年前后的 StyleGAN、BigGAN 达到高真实感高点；2020年后，新增标志性成果几乎全部出现在扩散与文生图路线。我们认为，这种“一条路线成果增量放缓、另一条路线成果密集涌现”的分布，是判断范式更迭的客观信号。

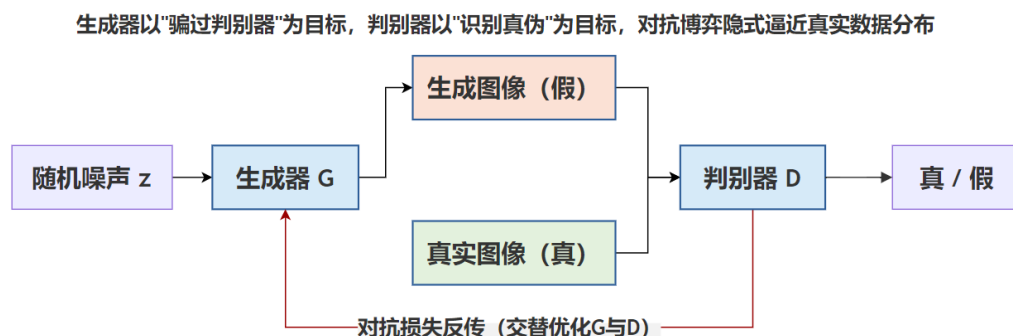
2. GAN 路线：对抗生成的原理、巅峰与局限

2.1. 对抗生成的基本原理

GAN 是 2016—2019 年生成式视觉的主线。GAN 的路线演进可概括为一条由基础框架走向高真实感、高可控的路径：GAN（对抗框架）→DCGAN（卷积化稳定）→pix2pix/CycleGAN（条件与非配对图像翻译）→ProGAN/BigGAN/StyleGAN（渐进式高分辨率、大规模类别条件、可控潜空间）。

从原理看，GAN 的核心是生成器与判别器的对抗博弈。生成器从随机噪声出发生成图像，判别器判断输入是真实样本还是生成样本；两者交替优化，生成器以“骗过判别器”为目标，判别器以“识别真伪”为目标，直到生成样本难以与真实样本区分。可以把它理解为“造假者”和“鉴定师”相互较量、共同提高。这一框架的重要意义在于，它不需要显式写出复杂的数据概率分布，而是通过“造假—鉴假”的博弈隐式逼近真实数据分布，从而绕开传统概率生成模型中难以处理的似然计算。

图2：GAN 对抗训练原理示意



数据来源：《Generative Adversarial Networks》，东吴证券研究所绘制

2.2. 从 DCGAN 到 StyleGAN 的演进

在基础框架之上，GAN 通过一系列结构与训练改进走向实用。2015 年，DCGAN 以全卷积结构稳定 GAN 训练，使其成为图像生成的标准架构；2017 年，pix2pix、CycleGAN 把 GAN 推广到图像到图像翻译，分别解决配对与非配对条件下的域转换；2018—2019 年，BigGAN 以大规模类别条件实现高保真合成，StyleGAN 则以基于风格的生成器引入可控潜空间，大幅提升人脸生成质量，标志 GAN 路线达到巅峰。

表3：GAN 路线代表论文与核心贡献

时间	代表论文	代表人物/机构	核心贡献	在 GAN 发展中的位置
2014 年	GAN	Goodfellow et al.	对抗生成框架，隐式建模数据分布	对抗生成路线起点
2015/16 年	DCGAN	Radford et al.	全卷积结构稳定 GAN 训练	GAN 图像生成标准化
2017 年	pix2pix	Isola et al.	条件 GAN 做配对图像到图像转换	可控视觉生成
2017 年	CycleGAN	Zhu et al.	循环一致性实现非配对域转换	图像风格/域迁移
2018/19 年	BigGAN	Brock et al.	大规模类别条件高保真合成	GAN 高保真高峰
2018/19 年	StyleGAN	NVIDIA	基于风格的生成器，可控潜空间，高质量人脸	GAN 时代巅峰

数据来源：《Generative Adversarial Networks》，《A Style-Based Generator Architecture for Generative Adversarial Networks》等其他文献，东吴证券研究所

2.3. 结构性局限

尽管 GAN 曾在高真实感图像生成上领先，但其固有局限同样明显，也正是后续扩散模型取而代之的原因。

其一，训练不稳定：生成器与判别器的对抗优化难以平衡，对超参数与网络结构敏感，容易不收敛或震荡。

其二，模式崩溃：生成器可能只覆盖真实分布的一部分模式，导致样本多样性不足。

其三，缺乏显式似然：GAN 隐式建模分布，难以直接评估样本概率，模型比较与训练监控较困难。

其四，语义与文本条件控制较弱：早期 GAN 以类别标签为主要条件，自然语言驱动生成能力不够自然。

GAN 的局限，尤其是训练稳定性与样本多样性，构成 GAN 路线天花板，并为 2020 年后扩散模型复兴留出空间。

3. Diffusion 路线：去噪生成如何完成范式替代

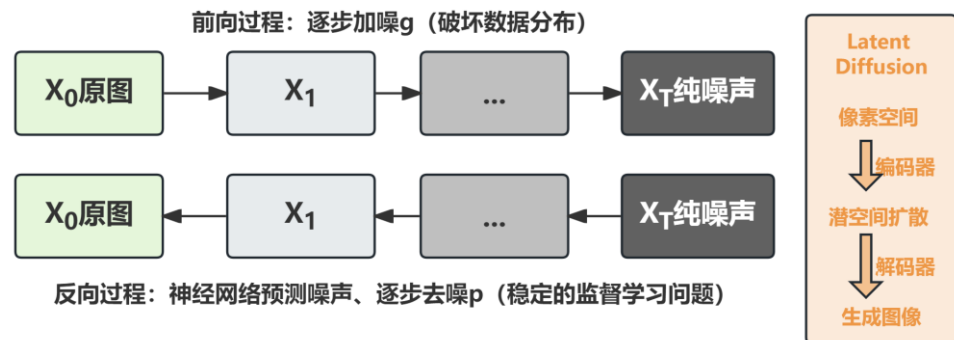
3.1. 去噪生成的基本原理

2020—2021 年是生成式视觉的范式替代期。扩散模型兴起，代表生成范式从对抗生成转向去噪生成。

从原理看，扩散模型的思想可追溯至 2015 年 Sohl-Dickstein 等人基于非平衡热力学提出的早期工作：先通过多步加噪逐步破坏数据分布，再学习反向的去噪恢复过程。这一思想在 2020 年由 Ho、Jain、Abbeel 的去噪扩散概率模型（DDPM）真正激活：模型在前向过程中对图像逐步加入高斯噪声，直至接近纯噪声；在反向过程中训练神经网络预测每一步所加噪声或去噪方向，从而从纯噪声出发逐步还原出图像。

换言之，扩散模型像是先把一张清晰图片一步步“弄花”，再训练模型学会一步步“擦干净”。其关键意义在于，把生成过程转化为稳定的监督学习问题，即给定加噪图像，学习如何去噪，这与 GAN 的对抗博弈形成鲜明对比。2021 年，Song 等人把分数匹配、扩散概率模型与连续时间随机微分方程（SDE）统一于同一框架，从理论上解释如何从噪声分布反向生成数据。

图3：扩散模型前向加噪—反向去噪过程及潜空间扩散示意



数据来源：《Denoising Diffusion Probabilistic Models》，《High-Resolution Image Synthesis with Latent Diffusion Models》等其他文献，东吴证券研究所绘制

3.2. Diffusion 对 GAN 的范式替代

我们认为，扩散模型在数年内取代 GAN 成为主线的原因同样是在于很好的解决了 GAN 的几项核心局限。

其一，训练稳定：去噪是标准监督回归问题，不存在对抗优化的平衡难题，训练更稳定、更易复现。

其二，覆盖度好：扩散过程对整个数据分布建模，显著缓解模式崩溃，样本多样性更佳。

其三，可控性强：去噪的迭代结构便于在每一步注入条件信息，为后续文本条件生成提供天然接口。

学术与产业信号相互印证：2021 年，Dhariwal 与 Nichol 以“扩散模型在图像合成上超越 GAN”为题给出基准证据；2022 年起，主流文生图产品（DALL·E2、Imagen、Stable Diffusion）几乎全部转向扩散路线，标志这一更迭在学术与应用两端同时验证。

表4: Diffusion 路线关键论文与演进

时间	代表论文	代表人物/机构	核心贡献	在 Diffusion 发展中的位置
2015 年	早期扩散模型	Sohl-Dickstein et al.	非平衡热力学，加噪—反向恢复思想	Diffusion 前史
2020 年	DDPM	Ho,Jain,Abbeel	去噪扩散复兴，预测噪声、稳定监督式生成	Diffusion 成为主线
2021 年	Score-Based SDE	Song et al.	统一 score matching、扩散与连续 SDE	Diffusion 理论统一
2021 年	DDIM	Song,Meng,Ermon	确定性采样、减少步数加速推理	采样效率改进
2021/22 年	Classifier-Free Guidance	Ho&Salimans	无需分类器的条件引导	条件生成可控性
2022 年	Latent Diffusion	Rombach et al.	潜空间扩散+交叉注意力接条件	Stable Diffusion 基础

数据来源：《Deep Unsupervised Learning using Nonequilibrium Thermodynamics》，《Denoising Diffusion Probabilistic Models》等其他文献，东吴证券研究所

3.3. 效率与算力局限

扩散模型的局限主要集中在效率与算力两方面。

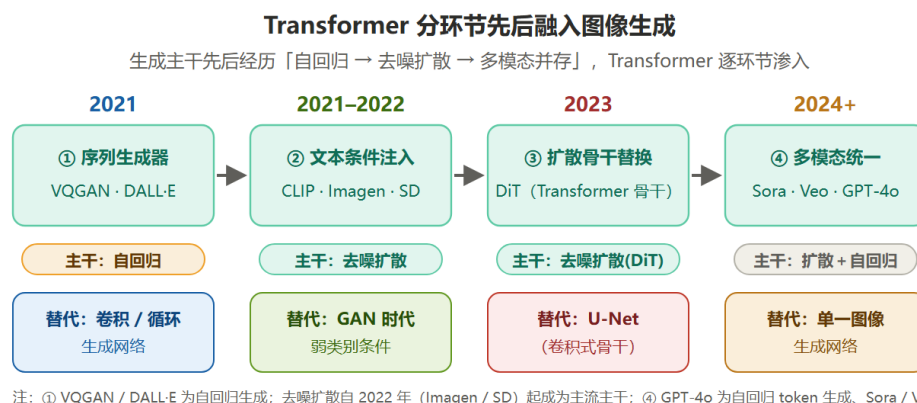
其一，采样速度慢：标准 DDPM 需要数十至上千步迭代去噪才能生成一张图像，推理成本远高于 GAN 的一次前向；为此，DDIM 等确定性采样方法被提出以减少步数。

其二，像素空间算力开销大：直接在高维像素空间做扩散，训练与推理都较昂贵。2022 年，Rombach 等人的潜在扩散模型（Latent Diffusion Models）针对该局限给出关键解法：不在像素空间、而在预训练自编码器的低维潜空间中做扩散，并通过交叉注意力接入文本等条件，从而大幅降低成本。这一改进同时打开“潜空间+文本条件”的通道，同时铺垫了 Transformer 在文本条件与骨干环节的介入。

4. Transformer 如何一步步融入图像生成

与语言、视觉理解领域中 Transformer 作为单一架构整体替换前代不同，在生成式视觉中，Transformer 并非一次性替换某个组件，而是分多个环节、在不同时间点先后介入。我们认为，Transformer 融入图像生成大致经历四个环节：自回归序列生成器、文本条件编码与注入、扩散骨干替换、多模态统一与时空扩展。

图4：Transformer 融入图像生成历程时间轴



数据来源：《Zero-Shot Text-to-Image Generation》，《Scalable Diffusion Models with Transformers》等其他文献，东吴证券研究所绘制

4.1. 环节一：自回归序列生成器

Transformer 进入图像生成的第一个环节，是作为自回归序列生成器。2021 年，OpenAI 的 DALL · E1 采用自回归 Transformer，把文本 token 与图像 token 拼接为一个序列统一建模，实现零样本文本到图像生成；同年，Esser、Rombach、Ommer 的 VQGAN 以离散化方式把图像编码为 token，再以 Transformer 建模其序列分布。这一环节中，Transformer 应用于生成骨干与序列建模，替代此前以卷积或循环结构为主的生成网络；其意义在于证明“图像可以 token 化，视觉生成也可以像语言一样被序列建模”。但自回归生成存在顺序依赖、推理效率和全局一致性方面限制，因此未成为文生图最终主线。

4.2. 环节二：文本条件编码与注入

Transformer 进入图像生成的第二个环节，是作为文本条件的编码器与注入机制。CLIP 通过大规模图文对比学习把图像与文本置于同一语义空间，使自然语言得以引用视觉概念。在此基础上，2022 年 Imagen 指出，采用大语言模型式的文本编码器对图文对齐与图像质量至关重要；GLIDE、DALL · E2 把文本条件引入扩散过程；Latent Diffusion/Stable Diffusion 则以交叉注意力把 Transformer 编码的文本表示注入扩散去噪网络。这一环节中，Transformer 应用于文本条件的编码与注入，替代 GAN 时代较薄弱的类别标签式条件机制；其意义在于使自然语言成为图像生成的主要控制接口，直接推动 2022 年文生图爆发。需强调的是，此环节中扩散仍是生成主干，Transformer 承担的是“语言理解与条件接入”的配套角色。

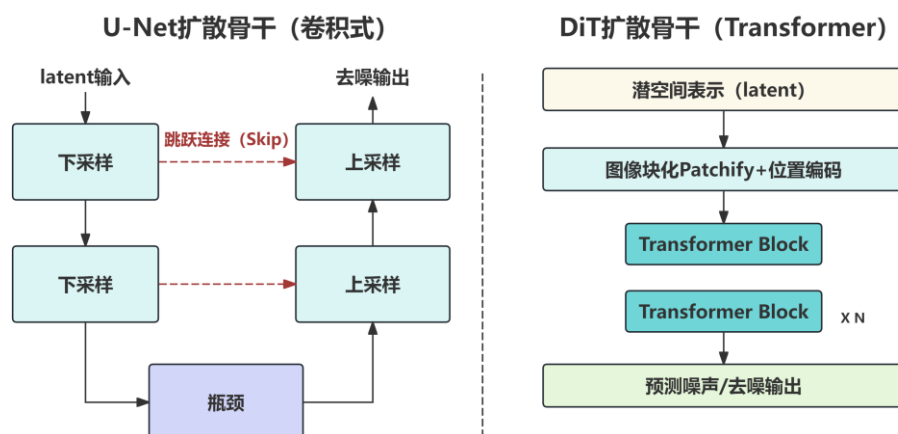
4.3. 环节三：扩散骨干替换 (DiT)

Transformer 进入图像生成的第三个环节，是替换扩散模型的骨干网络。2023 年，Peebles 与 Xie 的 DiT 用 Transformer 替代扩散模型常用的 U-Net，在潜空间的图像块

（latent patch）上建模，并展现良好的规模化特征。这一环节中，Transformer 替代的具体组件是 U-Net 这一卷积式去噪骨干；其意义在于把视觉生成接入大语言模型时代的规模化路线。

换言之，只要把视觉表示 token 化、patch 化，Transformer 同样可以作为生成模型骨干。此时，Transformer 升级的仍是扩散模型的骨干实现，而非改变去噪生成这一范式本身。

图5: DiT 与 U-Net 扩散骨干结构对比示意



数据来源:《Scalable Diffusion Models with Transformers》,《High-Resolution Image Synthesis with Latent Diffusion Models》, 东吴证券研究所绘制

4.4. 环节四：多模态统一与时空扩展

Transformer 进入图像生成的第四个环节,是作为统一多模态表示与时空建模骨干,推动生成对象从静态图像扩展到视频与多模态系统。2022 年, Ho 等人的 Video Diffusion Models 把扩散扩展到视频生成,重点解决高保真与时间一致性; Make-A-Video 把文生图进展迁移到文生视频; 2024 年, Lumiere 提出时空 U-Net 以改善全局时间一致性。2024 年以来, OpenAI 的 Sora 强调把不同类型视觉数据转化为统一表示,并将视频生成视为“世界模拟器”方向; Google DeepMind 的 Veo3 突出创意控制、原生音频与更长视频; GPT-4o 的原生图像生成则代表另一方向,即图像生成被整合进原生多模态模型,强调对话上下文、图像输入、文字渲染与多轮编辑。这一环节中, Transformer 应用于统一多模态序列表示与时空建模,使生成式视觉从单图创作走向视频与多模态创作系统。

4.5. 核心结论：Diffusion 为主角、Transformer 为工程化骨干

表5：Transformer 在生成式视觉任务中的四个介入环节

环节	代表模型（时间）	用于哪个环节	替代了什么组件	解决了什么问题
序列生成器	VQGAN 、 DALL·E1（2021）	生成骨干/序列建模	卷积或循环式生成网络	证明图像可 token 化、可序列建模
文本条件编码与注入	CLIP 、 Imagen 、 GLIDE 、 Stable Diffusion（2021—2022 年）	文本编码+交叉注意力注入	GAN 时代薄弱的类别标签条件	自然语言成为生成控制接口
扩散骨干替换	DiT（2023 年）	扩散去噪骨干	U-Net（卷积式骨干）	视觉生成接入规模化路线
多模态统一与时空扩展	Sora、Veo、GPT-4o（2024 年至今）	统一多模态表示+时空建模	单一图像生成网络	从单图走向视频与多模态创作系统

数据来源：《Taming Transformers for High-Resolution Image Synthesis》，《Scalable Diffusion Models with Transformers》等其他文献，东吴证券研究所

综合四个环节，Transformer 在生成式视觉中的融入路径较为清晰：先作为自回归序列生成器证明图像可被序列建模，再作为文本编码器与注入机制使语言成为生成入口，继而以 DiT 替换扩散骨干、接入规模化，最终作为统一多模态骨干支撑视频与多模态生成。

因此整体来看，本轮生成式视觉真正的范式更替是 GAN 到 Diffusion，即从对抗生成转向去噪生成；Transformer 的作用是工程化骨干升级与语言可控接口，而非生成范式本身的更替者。这一判断可由学术与产业两端相互印证：学术上，标志性成果的范式归属从 GAN 转向扩散，Transformer 的实质贡献集中在 DiT 这一骨干替换；产业上，主流文生图产品的生成主干仍是扩散，Transformer 主要承担文本编码、条件注入以及 DiT 路线中的骨干角色。

5. 风险提示

技术迭代不及预期的风险：数字 AI、物理 AI 技术进展低于预期，任务或应用场景拓宽速度低于预期。

大模型厂商 ARR 增速放缓的风险：若客户增长、续费率或客单价不及预期，模型厂商 ARR 增速可能放缓。

云服务商资本开支不及预期的风险：若自由现金流承压，或 AI 算力需求及投资回

报低于预期，云服务商可能下调资本开支。

智能驾驶及机器人商业化落地不及预期的风险：若产品可靠性、成本下降或用户需求低于预期，相关产品商业化进程可能放缓。



免责声明

东吴证券股份有限公司经中国证券监督管理委员会批准，已具备证券投资咨询业务资格。

本研究报告仅供东吴证券股份有限公司（以下简称“本公司”）的客户使用。本公司不会因接收人收到本报告而视其为客户。在任何情况下，本报告中的信息或所表述的意见并不构成对任何人的投资建议，本公司及作者不对任何人因使用本报告中的内容所导致的任何后果负任何责任。任何形式的分享证券投资收益或者分担证券投资损失的书面或口头承诺均为无效。

在法律许可的情况下，东吴证券及其所属关联机构可能会持有报告中提到的公司所发行的证券并进行交易，还可能为这些公司提供投资银行服务或其他服务。

市场有风险，投资需谨慎。本报告是基于本公司分析师认为可靠且已公开的信息，本公司力求但不保证这些信息的准确性和完整性，也不保证文中观点或陈述不会发生任何变更，在不同时期，本公司可发出与本报告所载资料、意见及推测不一致的报告。

本报告的版权归本公司所有，未经书面许可，任何机构和个人不得以任何形式翻版、复制和发布。经授权刊载、转发本报告或者摘要的，应当注明出处为东吴证券研究所，并注明本报告发布人和发布日期，提示使用本报告的风险，且不得对本报告进行有悖原意的引用、删节和修改。未经授权或未按要求刊载、转发本报告的，应当承担相应的法律责任。本公司将保留向其追究法律责任的权利。

东吴证券投资评级标准

投资评级基于分析师对报告发布日后 6 至 12 个月内行业或公司回报潜力相对基准表现的预期（A 股市场基准为沪深 300 指数，香港市场基准为恒生指数，美国市场基准为标普 500 指数，新三板基准指数为三板成指（针对协议转让标的）或三板做市指数（针对做市转让标的），北交所基准指数为北证 50 指数），具体如下：

公司投资评级：

买入：预期未来 6 个月个股涨跌幅相对基准在 15%以上；

增持：预期未来 6 个月个股涨跌幅相对基准介于 5%与 15%之间；

中性：预期未来 6 个月个股涨跌幅相对基准介于-5%与 5%之间；

减持：预期未来 6 个月个股涨跌幅相对基准介于-15%与-5%之间；

卖出：预期未来 6 个月个股涨跌幅相对基准在-15%以下。

行业投资评级：

增持：预期未来 6 个月内，行业指数相对强于基准 5%以上；

中性：预期未来 6 个月内，行业指数相对基准-5%与 5%；

减持：预期未来 6 个月内，行业指数相对弱于基准 5%以上。

我们在此提醒您，不同证券研究机构采用不同的评级术语及评级标准。我们采用的是相对评级体系，表示投资的相对比重建议。投资者买入或者卖出证券的决定应当充分考虑自身特定状况，如具体投资目的、财务状况以及特定需求等，并完整理解和使用本报告内容，不应视本报告为做出投资决策的唯一因素。

东吴证券研究所
苏州工业园区星阳街 5 号

邮政编码：215021

传真：（0512）62938527

公司网址：<http://www.dwzq.com.cn>