

能生成，不等于能成为 AI 产品

—— 一次从“可运行 Demo”到“可验证产品”的项目复盘

文章类型	记录时间	案例来源
AI 产品观察	2026 年 7 月 23 日	AI Game Content Copilot

核心判断：模型证明的是“能力存在”，产品需要证明的是“用户能够在可接受的成本与风险下，稳定完成任务”。

第一次做大模型应用时，我很容易把“成功生成结果”当作项目已经完成：用户输入一句话，模型返回一段看起来合理的内容，页面没有报错，整个流程也能顺利运行。但随着测试次数增加，我逐渐发现，能够生成只是 AI Demo 的起点，而不是 AI 产品的终点。

输出是否稳定、用户是否知道应该怎样输入、错误结果能否被识别、失败后是否存在恢复路径、响应时间和调用成本能否接受，这些问题都不会因为模型成功返回了一次答案而自动消失。模型提供的是一种可能性，而产品真正需要解决的，是如何把这种带有不确定性的能力，转化为用户可以理解、控制并重复获得的价值。

从“模型能生成”到“用户可重复获得价值”

AI 产品化不是增加一个按钮，而是补齐任务闭环。



产品经理的工作：把模型的不确定性，转化为可理解、可控制、可评测、可持续的用户价值。

图 1 从模型能力到重复用户价值的产品化路径（作者整理）

01 从“跑起来”到“任务闭环”

在着手开发 AI Game Content Copilot 时，我最初的目标很简单：先做出一个能够运行的原型。我完成产品定位后，借助 Cursor 搭建了第一版 Demo，并验证了“输入—生成—输出”的基本链路。看到模型成功输出一次内容时，我一度认为核心任务已经完成。

后来重新以用户任务为中心审视这个项目，我才意识到：这只是完成了技术链路验证，还没有建立产品闭环。原型证明了模型可以生成内容，却没有回答三个更关键的问题——输出是否持续可用、用户如何判断结果质量，以及完成一次有效任务究竟需要付出多少时间与成本。

我在初版中发现的四类问题

问题类型	初版表现	产品经理需要继续追问
结果稳定性	同一档位能够生成 A/B/C 三个版本，但部分字段偶尔缺失，输出结构和内容完整性不稳定。	记录字段缺失率、有效输出率和重复生成波动。
结果可判断性	虽然提供了多个版本，但缺少清晰的版本定位与评分标准，用户难以判断哪个结果更适合继续使用。	为版本定义目标，并建立内容质量与场景适配评分。
延迟与成本	高能力模型在长文本和剧情表达上表现突出，但普通任务的等待时间和调用成本偏高。	比较首 Token 延迟、完整任务时间和单位成功任务成本。
失败恢复与可观测性	发生超时、字段缺失或低质量输出时，系统缺少明确提示，也没有形成 Bad Case 记录。	增加错误分类、重试/降级策略与生成日志。

复盘原则：不要用“模型曾经生成过一个好答案”证明产品可用，而要用一组可重复的任务、指标和失败样本说明产品是否可靠。

02 什么才算“成为 AI 产品”

基于这次复盘，我把 AI 产品化拆成六个相互约束的维度：可用性、稳定性、可控性、可评测性、成本效率与风险边界。它们不是孤立的检查清单，而是围绕用户任务形成的一组平衡关系。某个维度表现突出，并不能自动弥补其他维度的缺失。



图 2 AI 产品六维产品化框架（作者整理）

维度	产品问题	可观察指标示例
可用性	用户能否理解输入方式，并在较低学习成本下完成核心任务。	任务完成率、输入修改次数、操作路径长度
稳定性	产品是否会超时、崩溃、缺字段，输出质量是否出现大幅波动。	有效输出率、字段完整率、错误率、P95 延迟
可控性	产品是否能够约束格式、长度、风格和风险边界，失败后是否可以恢复。	格式遵循率、重试成功率、人工接管率
可评测性	团队能否定义好结果，并通过测试集与 Bad Case 迭代。	评分一致性、场景通过率、回归测试通过率
成本效率	完成一次有效任务的时间、Token 和基施成本是否合理。	单位成功任务成本、平均生成时长、资源利用率
风险边界	隐私、安全、内容合规与模型升级风险是否被识别和控制。	违规率、敏感信息暴露率、版本回归异常数

03 我的项目如何发生变化

重新定义完成标准后，我不再把“模型返回内容”视为终点，而是围绕输入、生成、评价、失败恢复与数据记录重构流程。此次改造的重点不是堆叠功能，而是让每个环节都能够被观察、被判断和被迭代。

环节	初版	迭代后
输入方式	自由输入，需求容易遗漏	结构化字段与示例输入，降低歧义
生成结果	单次输出或版本差异不清晰	A/B/C 多版本生成，并明确各版本定位
质量判断	依赖主观阅读	加入多维评分与人工复核入口
异常处理	超时或缺字段后只能重新生成	增加失败提示、重试与降级路径
数据记录	没有生成历史和问题沉淀	记录生成结果、Bad Case、Token 与耗时
模型策略	所有任务使用同一高能力模型	普通任务优先轻量模型，复杂任务按需升级



图 3 以 Bad Case 和指标驱动持续迭代闭环（作者整理）

关于结果：当前改造主要完成了产品机制与观测能力的补齐。由于尚未完成足够规模的对照实验，本文不虚构“提升百分比”，后续需要通过固定测试集和真实使用数据验证效果。

04 下一步：用指标验证，而不是凭感觉宣布完成

AI 产品的迭代不能只依赖“这次看起来更好”。为了判断改造是否真正有效，我会把验证拆成离线评测和在线产品指标两部分：离线评测用于控制变量、快速回归；在线指标用于观察真实用户是否更顺利地完成任务。

验证层级	指标	如何理解
离线评测	固定测试集通过率	对一组角色、剧情、广告文案与复杂约束任务进行重复测试
离线评测	结构化输出完整率	检查必填字段、格式、字数与禁止项是否满足要求
离线评测	人工评分一致性	用统一 Rubric 对相关性、创意度、可用性等维度评分
在线指标	核心任务完成率	用户是否在无需反复修改输入的情况下获得可用结果
在线指标	重试率与人工接管率	输出失败或不满意后，用户是否需要频繁重新生成
在线指标	P95 完整任务时间	从提交需求到获得完整可用结果所需时间
在线指标	单位成功任务成本	总调用成本 ÷ 成功完成的有效任务数
在线指标	用户满意度	用户对结果质量、操作成本与等待时间的综合评价

其中，我更关注“单位成功任务成本”，而不是孤立地看每百万 Token 单价。一个模型单价更低，但如果输出更长、失败重试更多，完成一次有效任务的真实成本仍可能更高。对产品而言，成本必须和任务成功一起计算。

05 AI 产品经理在其中承担什么角色

这次项目复盘让我重新理解了 AI 产品经理的工作边界。产品经理既不能把模型能力当作产品价值，也不能只通过增加流程来掩盖模型问题，而需要在用户目标、模型边界、工程实现与商业约束之间做取舍。

- 先定义用户要完成的任务，再决定是否需要 AI，以及模型承担哪一部分。
- 把“好结果”转化为可执行的评测标准，而不是停留在个人感觉。
- 为失败设计恢复路径，并把 Bad Case 变成下一轮迭代输入。
- 同时关注效果、延迟、成本与风险，避免只追求单一 Benchmark。
- 让模型、规则、工作流和人工确认各自承担最合适的任务。

06 写在最后

模型负责提供可能性，产品负责把不确定的可能性转化为用户能够重复获得的价值。一个 AI 产品的好坏不仅取决于模型能力，更取决于团队能否明确用户痛点和使用路径，建立可验证的质量标准，处理失败与异常，并把响应时间、调用成本和风险控制在接受范围内。

因此，“能生成”只能证明技术链路成立；只有当用户能够稳定完成任务，团队能够持续发现并修复问题，产品才真正开始成立。对我而言，这也是从“做出一个 Demo”走向“用品思建 AI 应用”的第一步。

结论：AI Demo 的完成标准是模型成功输出；AI 产品的完成标准，是用户能够稳定、可控、可评测地完成任务。

附：上线一个 AI 功能前，我会再检查什么

检查项	上线前问题
用户任务	目标用户、触发场景和成功标准是否清晰？不用 AI 是否也能更低成本地解决？
输入设计	用户是否知道如何输入？系统是否能补全必要字段并减少歧义？
输出质量	是否有固定测试集、评分 Rubric 和最低通过标准？
失败恢复	超时、缺字段、低质量或违规输出发生后，是否有提示、重试、降级与人工接管？
成本与性能	是否持续记录有效输出率、P95 延迟、重试率和单位成功任务成本？
风险与迭代	隐私、内容安全、权限边界是否明确？模型或 Prompt 更新后是否会做回归测试？

这份清单并不能保证一个 AI 功能一定成功，但它能够帮助我避免只用“模型能不能生成”来判断产品是否完成。