

付费版

AI 互动类型自评

专业测评报告

先看你更常怎样与 AI 开始互动，再练习按任务在探索、表达与控制之间切换

根据你本次作答生成

我在AI协作中卡在角色切换上

练习在任务的不同阶段有意识地切换AI协作角色：先让AI提供创意，再主动审校并决定采纳程度。

报告类型

个性化专业报告

语言

zh-CN

版本

2026.07

01 先看结论

本次作答提示：我在AI协作中卡在角色切换上

你在需要新想法时会主动让AI参与，但在关键细节上仍倾向亲自把关，这导致角色切换不够明确。

现在最值得做的一件事

今天挑选一项正在进行的工作，先用AI生成三个创意方案，然后逐一手动检查并选出最合适的一个。

四个最值得留意的情境

1

你收到同事推荐的最新AI文档分析插件

先在工作日志中记录一次你尝试新AI工具的具体步骤和感受

2

团队正在讨论新产品概念

先记录并在团队头脑风暴时把AI生成的一个创意写在白板上，并说‘我们可以进一步讨论这个方向’

3

需要提交季度业绩报告

完成AI生成的报告后，先用5分钟手动检查关键数据，再发送给上级

4

新项目需要快速数据洞察

在项目启动会议上说明‘我们将使用AI进行数据分析，但所有结论都由人工复核后确认’

如何阅读这份报告

01

看懂结果

先抓住核心答案

02

定位场景

识别卡住的时刻

03

拿走工具

直接使用步骤与脚本

04

验证变化

用 30 天记录修正

02 关键组合

你不是一种固定的 AI 类型；更重要的是能否按任务目的与风险切换互动方式

虽然你经常在创意阶段使用AI，但在实际操作中仍会频繁手动审校，这表明你可以在同一任务中灵活切换“创意触发者”和“审校守护者”两种角色。

01

提升技术开放度的第一步

当你对新AI工具的兴趣不强时，先给自己设定一个简短的探索时间，而不是完全放弃尝试。

什么时候容易出现

你在日常工作中很少主动使用新AI软件且几乎不了解同事推荐的工具。

如果一直这样处理

如果继续回避，可能错失提升工作效率的机会。

可以改变的一步

先记录并主动搜索并尝试一款新发布的AI工具

02

让想象力表达更可见

当你偶尔使用AI灵感但不常分享时，记录下来可以帮助你在需要创意时快速回顾。

什么时候容易出现

你经常把AI生成的内容当作灵感但在团队讨论中只偶尔提及。

如果一直这样处理

不记录可能导致好点子被遗忘，影响创意输出。

可以改变的一步

在头脑风暴时把AI生成的点子写在便签上并标记‘可进一步探索’

03

平衡控制偏好与效率

当你倾向于保留手动审校环节且有时确认关键步骤时，明确这一步骤可以平衡自动化与控制感。

什么时候容易出现

你经常保留手动审校且有时亲自确认关键步骤。

如果一直这样处理

跳过校对可能导致错误进入最终成果，增加返工成本。

可以改变的一步

在AI输出后先进行一次手动校对，再决定是否采用

04

任务阶段间的角色切换指南

结合你在策划时喜欢加入AI建议、常用AI灵感以及保留审校环节，帮助你在不同任务阶段切换角色。

什么时候容易出现

你在策划阶段常加入AI建议，并在创意和审校上都有一定使用频率。

如果一直这样处理

不明确分工可能导致角色冲突或任务重复。

可以改变的一步

在项目策划阶段先列出AI可提供的三类支持（数据、创意、流程），并决定哪部分需要人工审查

03 真实情境

四种最容易让好奇、想象或控制偏好用错位置的任务

先看任务是探索、表达、整理还是真实决策，再决定要开放多少、约束多少、由谁复核。

01

你收到同事推荐的最新AI文档分析插件

先出现的信号：同事在会议上展示插件功能并邀请你试用

你可能会这样应对：你感到有些不安，想了解该工具的实际价值

容易造成的误读：同事可能会认为你对新工具不感兴趣或不愿意学习

这时先做：先在工作日志中记录一次你尝试新AI工具的具体步骤和感受

02

团队正在讨论新产品概念

先出现的信号：团队成员提出需要更多创新点子

你可能会这样应对：你有点犹豫，但想让AI的想法被考虑

容易造成的误读：同事可能认为你在依赖AI而不是自己的想法

这时先做：先记录并在团队头脑风暴时把AI生成的一个创意写在白板上，并说‘我们可以进一步讨论这个方向’

03

需要提交季度业绩报告

先出现的信号：系统提示报告已自动生成并准备发送

你可能会这样应对：你担心AI的细节可能有误，需要确认

容易造成的误读：上级可能认为你对AI不信任，导致对你的效率产生疑虑

这时先做：完成AI生成的报告后，先用5分钟手动检查关键数据，再发送给上级

04

新项目需要快速数据洞察

先出现的信号：项目经理提出使用全自动AI流程的建议

你可能会这样应对：你想确保团队对AI使用的范围和审查流程都有共识

容易造成的误读：项目成员可能误以为你不愿意让AI发挥全部作用

这时先做：在项目启动会议上说明‘我们将使用AI进行数据分析，但所有结论都由人工复核后确认’

04 反应路径

任务目的→互动入口→生成方式→控制边界→人工检查→下一次切换

趣味原型只整理当前自报偏好，不评定 AI 能力、生产力或专业胜任力。

01	报告出现新数据时	注意力先去哪里 在项目策划阶段看到AI生成的分析报告	你会注意到 注意到报告中出现大量未核实的数据	原来的应对 继续使用报告作为主要决策依据	继续下去的成本 若不核实，可能导致后续方案基于错误信息而需返工	新的中断动作 先在两分钟内记录报告中未核实的关键点并标记为‘待核实’
02	同事推荐新工具时	注意力先去哪里 同事在会议中提出使用新AI工具	你会注意到 感到有些不安，想了解工具细节	原来的应对 保持沉默或仅表面同意	继续下去的成本 若不明确需求，后续使用可能出现功能不匹配或额外学习成本	新的中断动作 先在两分钟内写下三条该工具可能提升的具体工作环节
03	AI想法被提出时	注意力先去哪里 在创意头脑风暴时AI提供了奇思妙想	你会注意到 对AI想法产生兴趣但又担心缺乏个人色彩	原来的应对 直接引用AI想法作为团队输出	继续下去的成本 若直接使用，团队可能误以为缺乏个人创新，影响成员参与感	新的中断动作 先在两分钟内将AI想法转化为一句个人化的阐述再提交

05 使用说明

你的 AI 互动使用说明：什么适合放开，什么必须收束

我自然的入口

保持人与AI的协作透明度，确保每一步都有明确的人工确认或反馈。

目标不清时

对AI生成的创意保持开放但不盲从，必要时进行二次验证或改编。

适合探索时

在每个任务阶段设定明确的AI使用时长或步骤上限，超出后必须人工审查。

需要控制时

当你感到对AI的控制不足或创意受限时，先暂停使用，记录具体情境，再决定是否继续或切换角色。

有效的人工检查

当AI提供的建议被采纳并提升效率时，记录成功案例，强化正向体验。

容易卡住

如果AI的输出不符合预期，及时回退到手动模式并记录原因，避免重复同样的错误。

这份报告建议你优先做的决定

练习在任务的不同阶段有意识地切换AI协作角色：先让AI提供创意，再主动审校并决定采纳程度。

06 行动协议

开始一个 AI 任务前，用五步选择互动角色

写清目的、风险、允许资料、输出标准和人工接手点，再决定这次偏探索还是偏控制。

01

限定当前问题

写下当前你在工作中使用 AI 的具体任务，例如数据分析、文案生成或原型设计。

先暂停：如果在记录时发现任务本身不涉及 AI，暂停本流程并重新选择需要 AI 支持的任务。

02

分开事实与解释

把刚才记录的任务分成两类：需要你主动探索新工具的环节（如尝试新软件）和需要你保持人工核对的环节（如审校输出）。

先暂停：若两类环节都找不到，请回到步骤 1，重新确认任务范围。

03

核对未知与反例

核对以下两组证据：- 你在 相关作答（经常在策划阶段加入 AI 建议）和 相关作答（经常保留手动审校）中表现出同时愿意尝试和保持控制。- 你在 相关作答（经常把 AI 生成的内容当作灵感）和 相关作答（几乎没有主动检查细节）中表现出想象力与低检查倾向。记录哪些证据与你的实际感受不符，作为反例或边界条件。

先暂停：如果发现所有证据都与实际感受一致，直接进入下一步；若出现明显不符，暂停并在团队或导师处讨论该任务的 AI 角色定位。

04

选择低风险动作

选择一个低风险的角色切换动作：在下一个相似任务的策划阶段，先使用 AI 生成创意（参考 相关作答），随后在输出完成后立即进行手动审校（参考 相关作答）。

先暂停：如果在切换过程中出现关键错误或时间超出预期，立即停止并记录问题，以便后续调整。

05

记录结果

记录结果：- AI 生成的创意数量或质量（可用 1-5 评分）- 手动审校所花时间（分钟）- 任务完成的整体满意度（1-5）

先暂停：若记录显示满意度低于 3 分且审校时间超过预期 30%，则进入调整或升级阶段。

06

决定继续、调整、停止或升级

根据记录决定下一步：- 若满意度高且审校时间合理，继续在其他任务中重复相同切换模式。- 若满意度一般或审校时间偏长，尝试在下次任务中先加强人工核对（参考 相关作答）再使用 AI 创意。- 若出现严重错误或团队反馈负面，暂停自行实验，寻求同事或专业顾问的帮助。

先暂停：如果在三次循环后仍未达到预期效果，升级至专业 AI 协作培训或外部顾问支持。

07 可直接使用

不必只说“帮我做”，可以这样交接探索、表达与收束任务

在团队会议核对任务分配时。

容易让事情更难 直接说“我不想用AI”。

可以说 “我注意到我们可以尝试用AI来快速生成数据分析，这样可以节省时间，你们觉得怎么样？”

这句话的作用 大家对AI的使用范围和目标更明确了。

在头脑风暴环节提出创意。

容易让事情更难 只说“我不擅长AI”。

可以说 “我刚用AI生成了几个想法，想和大家分享并一起改进。”

这句话的作用 团队对AI提供的创意有了共识并决定进一步探索。

在项目进度检查时。

容易让事情更难 直接要求AI全权完成任务。

可以说 “我会先让AI生成草稿，然后我们一起核对关键细节，确保质量。”

这句话的作用 人工审校与AI输出的配合让结果更可靠。

任务结束后进行复盘。

容易让事情更难 忽视对AI使用的回顾。

可以说 “我们回顾一下这次AI的使用情况，哪些环节顺畅，哪些需要手动干预？”

这句话的作用 对AI的使用边界和改进点更清晰，下一次可以更自然切换角色。

08 安全边界

这些结果不能证明能力或安全

互动偏好不是 AI 素养、绩效、模型安全、合规或专业能力认证。

01 限定当前问题

如果在记录时发现任务本身不涉及 AI，暂停本流程并重新选择需要 AI 支持的任务。

02 分开事实与解释

若两类环节都找不到，请回到步骤 1，重新确认任务范围。

03 核对未知与反例

如果发现所有证据都与实际感受一致，直接进入下一步；若出现明显不符，暂停并在团队或导师处讨论该任务的 AI 角色定位。

04 选择低风险动作

如果在切换过程中出现关键错误或时间超出预期，立即停止并记录问题，以便后续调整。

高影响任务仍需现实权限与专业复核

涉及医疗、法律、财务、招聘、安全或敏感资料时，请遵守当地规则、组织政策和资料授权，并由具备能力的人最终批准。

09 已有资源

你已有的互动资源，以及角色不切换时的成本

01

规划与沟通能力

它有用的时候：任务要求高准确性或合规审查时

容易用过头：过度审校AI输出导致任务进度放慢

别人可能怎么理解：对方可能会觉得你不信任AI的能力，认为你在拖延

这样校正：在需要快速决策时，先列出两三个关键点，而不是先完整审阅AI生成的全部内容

02

自主取舍与边界能力

它有用的时候：项目已有明确技术路线且变更成本高时

容易用过头：过度回避新工具使创新机会受限

别人可能怎么理解：同事可能会误以为你对新技术缺乏兴趣，导致合作机会减少

这样校正：先当出现新AI工具时，先设定一分钟的探索时段，而不是完全放弃尝试

03

复盘与调节能力

它有用的时候：需要在短时间内产生多方案时

容易用过头：较多依靠AI创意导致个人表达被稀释

别人可能怎么理解：团队成员可能会认为你对AI创意不够重视，觉得你缺乏想象力

这样校正：在创意讨论时，将AI提供的想法标记为‘可供参考’，随后自行补充个人观点，而不是直接采用

10 行动计划

用 30 天验证：哪一次角色切换真正改善了任务

这 30 天只验证一件事

结合本次作答显示的技术开放度、想象力表达和控制偏好，你可以在不同任务阶段（探索-生成-审校）尝试更自然的AI协作角色，并通过有意识的切换提升协作流畅度。

开始前：第1周记录你在日常工作中主动尝试新AI工具的次数、在头脑风暴时引用AI创意的次数、以及手动审校AI输出的次数，作为本次作答的初始行为基线。

第 2 周

第2周加入低风险动作：在一次例行任务中，先使用AI生成初稿，再自行添加两处手动修改后提交。

要记下来：记录是否完成AI生成、手动修改次数以及完成任务所用总时长。

第 3 周

第3周在真实项目的策划阶段重复第2周的流程，并在团队头脑风暴时主动提出AI提供的一个创意点。

要记下来：记录AI创意被采纳次数、团队反馈以及个人对创意价值的主观评分。

第 4 周

第4周比较第1-3周记录的三项指标，决定是否保留、调整或停止该协作模式。

要记下来：对比技术探索频率、想象力表达使用度、控制介入次数的变化幅度，并记录主观满意度（1-5分）。

第 4 周

如果第4周指标未出现预期提升，暂停该实验并在下次任务中重新评估AI角色选择。

要记下来：记录暂停决定的原因和下一步计划。

技术探索频率

记录每日主动尝试新AI软件的次数（包括阅读功能介绍或实际操作）

希望看到：相较第1周基线增加

想象力表达使用度

记录在团队讨论或个人创作时引用AI生成的创意或素材的次数

希望看到：相较第1周基线增加

控制介入次数

记录手动审校或调整AI输出的次数（包括检查细节、确认关键步骤）

希望看到：相较第1周基线保持或略增

11 保存与分享

把这一页存下来，作为 AI 任务起步卡

这次目的

如果您在任务开始时先确认 AI 能提供的价值（如数据分析或创意提示），随后在关键审校环节介入检查，就能兼顾开放探索与控制需求。

我先探索

面对不确定的 AI 工具时，先进行小范围试用（如一次性生成报告），再评估是否值得在更大范围推广，可降低尝试成本并保持对结果的掌控感。

我会约束

您在技术开放度上表现出“有时”尝试新 AI 和“经常”在策划阶段加入 AI 建议，但在日常工作中主动使用新工具的频率较低。在想象力表达上，您“经常”把 AI 生成的内容当作二次创作的灵感，但对在社交媒体分享或将 AI 想法转化为原型的兴趣较弱。控制偏好方面，您“经常”保留手动审校 AI 输出，而对自行设定规则或检查细节的意愿相对较低。

不能输入

当你感到对 AI 的控制不足或创意受限时，先暂停使用，记录具体情境，再决定是否继续或切换角色。

谁来检查

这说明您在项目关键阶段倾向于让 AI 提供创意和数据支持，但在执行细节上仍希望保持人工把关。

何时接手

当任务需要持续、低门槛的 AI 操作时，您可能会因缺乏主动尝试而错失效率提升的机会；而在需要快速迭代创意时，较少在社交平台或原型阶段使用 AI 可能限制了灵感的扩散。

我的一句话行动原则

“遇到需要决定是否使用 AI 辅助时，先评估任务是否属于“创意提供”或“细节审校”。若是创意提供，先尝试一次 AI 生成，若是细节审校，则保留手动检查。”

我的补充

12 互动工具

AI 互动角色混合图：探索、表达与控制不是固定类型

把技术开放、想象表达和控制偏好放进具体任务，看到你当前更常从哪个入口与 AI 协作。

01 提升技术开放度的第一步

当你对新AI工具的兴趣不强时，先给自己设定一个简短的探索时间，而不是完全放弃尝试。

可观察时刻

你经常把AI生成的内容当作灵感但在团队讨论中只偶尔提及。

需要保留的边界

如果继续回避，可能错失提升工作效率的机会。

下一步验证

先记录并主动搜索并尝试一款新发布的AI工具

02 让想象力表达更可见

当你偶尔使用AI灵感但不常分享时，记录下来可以帮助你快速回顾。

可观察时刻

你经常把AI生成的内容当作灵感但在团队讨论中只偶尔提及。

需要保留的边界

不记录可能导致好点子被遗忘，影响创意输出。

下一步验证

在头脑风暴时把AI生成的点子写在便签上并标记‘可进一步探索’

03 平衡控制偏好与效率

当你倾向于保留手动审校环节且有时确认关键步骤时，明确这一步骤可以平衡自动化与控制感。

可观察时刻

你经常保留手动审校且有时亲自确认关键步骤。

需要保留的边界

跳过校对可能导致错误进入最终成果，增加返工成本。

下一步验证

在AI输出后先进行一次手动校对，再决定是否采用

04 任务阶段间的角色切换指南

结合你在策划时喜欢加入AI建议、常用AI灵感以及保留审校环节，帮助你在不同任务阶段切换角色。

可观察时刻

你在策划阶段常加入AI建议，并在创意和审校上都有有一定使用频率。

需要保留的边界

不明确分工可能导致角色冲突或任务重复。

下一步验证

在项目策划阶段先列出AI可提供的三类支持（数据、创意、流程），并决定哪部分需要人工审查

13 互动工具

任务切换卡：什么时候尝试，什么时候收束与人工接手

低风险探索可以更开放；涉及事实、资料、影响他人的输出时，要提高边界、复核和责任要求。

01

你收到同事推荐的最新AI文档分析插件

已观察信号：同事在会议上展示插件功能并邀请你试用

需要核对：同事可能会认为你对新工具不感兴趣或不愿意学习

低风险下一步：先在工作日志中记录一次你尝试新AI工具的具体步骤和感受

02

团队正在讨论新产品概念

已观察信号：团队成员提出需要更多创新点子

需要核对：同事可能认为你在依赖AI而不是自己的想法

低风险下一步：先记录并在团队头脑风暴时把AI生成的一个创意写在白板上，并说‘我们可以进一步讨论这个方向’

03

需要提交季度业绩报告

已观察信号：系统提示报告已自动生成并准备发送

需要核对：上级可能认为你对AI不信任，导致对你的效率产生疑虑

低风险下一步：完成AI生成的报告后，先用5分钟手动检查关键数据，再发送给上级

04

新项目需要快速数据洞察

已观察信号：项目经理提出使用全自动AI流程的建议

需要核对：项目成员可能误以为你不愿意让AI发挥全部作用

低风险下一步：在项目启动会议上说明‘我们将使用AI进行数据分析，但所有结论都由人工复核后确认’

14 结论依据

这些结论来自哪些作答，以及它不能认证什么

先核对测量边界和本次原始分，再查看每条结论引用了哪些具体作答。

本次维度

技术开放度

9

面对新 AI 工具与任务时愿意尝试、探索和学习程度。

原始分 · 仅在当前量表内查看

想象力表达

7

使用 AI 扩展想法、表达方式与创意选项的倾向。

原始分 · 仅在当前量表内查看

控制偏好

9

在任务中设置约束、核对过程并保留人工控制的倾向。

原始分 · 仅在当前量表内查看

怎样使用这份结果

原始分由 PDF 测量页直接呈现，仅作本次作答的记录。

本报告仅依据您在本次测评中选择的方案进行解释，未使用任何人口常模或外部比较。

- 只能反映您当前对 AI 工具的尝试、创意使用和控制倾向。
- 不提供职业匹配、人格标签或未来表现预测。
- 结论基于已选答案的直接证据，未涉及未选择的选项。

未对您进行临床诊断、疾病标记或跨群体排名。

报告编号：AI-INTERACTION / SAMPLE-V1 · 完成日期：2026-07-23

15 结论依据

逐条核对：结论与作答依据

每条结论同时保留支持线索和反例。它们来自本次作答，不是人口常模或诊断证据。

	本次作答显示你在技术开放度上既有尝试新AI工具的意愿（如相关作答经常使用AI生成数据分析），也会在日常工作中保持保守（如相关作答很少主动尝试）	第 1 题 第 2 题
也要看到：	第 3 题	
	本次作答说明你在想象力表达上会利用AI获取灵感（相关作答经常二次创作），但对将AI创意转化为产品或公开分享持保留态度（相关作答很少）	第 6 题 第 7 题
也要看到：	第 8 题	
	本次作答表明你倾向于在关键环节保留人工审校（相关作答经常保留手动审校），即使在规则设定上不常主动干预（相关作答几乎没有）	第 11 题 第 12 题
也要看到：	第 13 题	
	本次作答提示你在控制偏好上会在关键步骤保持人工确认（相关作答有时确认），而不倾向于让AI全权决定（相关作答很少）	第 11 题 第 12 题
也要看到：	第 13 题	