

付费版

AI 协作与复核准备度检查

专业测评报告

先按任务风险分工，再把上下文、复核、数据边界和最终责任写清

根据你本次作答生成

你已经会把任务交给 AI，但证据、数据和最终责任还没有跟上交接速度

先选一个正在重复使用的任务，按风险写清允许输入的资料、必须核对的来源、停止条件、复核人和最终批准人。

报告类型

个性化专业报告

语言

zh-CN

版本

2026.07

01 先看结论

本次作答显示：你已经会把任务交给 AI，但证据、数据和最终责任还没有跟上交接速度

你能判断部分任务需要保留现场知识，也会说明使用者和场景；结果不理想时，也能找到目标或格式问题。更容易断开的地方是打开原始来源、主动寻找反例、确认资料授权、记录采纳理由和保留纠正路径。

现在最值得做的一件事

今天为一个真实任务填一张六栏交接卡：用途、允许资料、质量标准、原始来源、停止条件、最终批准人。

四个最值得留意的情境

1

用 AI 起草一份给客户看的政策说明。

先标记为草稿，打开每条关键来源，指出仍未核实的内容，并确认具备相关知识的复核人与批准人。

2

需要把一段客户资料交给工具整理。

先暂停输入，确认资料类型和授权；如果任务可以完成，改用占位符、摘要或最少字段，并记录允许范围。

3

同一种内容任务再次出现，团队准备沿用上次方法。

先记录上次有效输入、失败例子、关键差异和转人工位置，再用同一检查标准完成本轮。

4

一项 AI 参与的结果会影响员工、客户或其他个人。

先确认具体批准人，记录采用理由，并向接收者说明怎样提出异议、由谁人工复查和怎样纠正。

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

问题不只是“会不会提示”，而是交接、证据、数据和批准链有没有一起闭环

再写一轮更长的提示，不一定能补上当前缺口；当问题在资料权限、独立来源和批准责任时，继续优化措辞反而可能让输出更流畅，却更难看见没有被核对的部分。

01

任务边界已经出现，但用途不清会让草稿悄悄变成答案

你能保留需要现场知识和责任的部分，不过输出用于发散、草稿、核对还是决定参考，并不总在开始前被说清。同一份内容可能在流转中被赋予更高权重。

什么时候容易出现

一份内部草稿被转发给客户、主管或其他系统，而接收者不知道它是否经过来源检查和人工批准。

如果一直这样处理

低风险试写可能越过原本边界，成为影响他人的正式材料。

可以改变的一步

先在任务名称和产出顶部说明用途、禁止用途、复核状态和最终批准人。

02

提示可以继续优化，但缺失来源与资料权限不能靠措辞补齐

结果不理想时，你能识别部分目标和格式问题，也会减少一部分输入；但若原始来源、授权范围或领域知识缺失，再换问法仍不会让这些条件自动出现。

什么时候容易出现

输出看起来越来越完整，却没有人打开法规、数字、引用或确认客户资料能否进入当前工具。

如果一直这样处理

文字质量改善会遮住输入权限和事实核验仍为空白，增加无记录放行。

可以改变的一步

把问题改为三步：先确认允许资料，再列原始来源，最后指定具备相关知识的复核人。

03

会定位失败还不够，最终责任必须能被追到一个人和一次决定

你可以判断问题来自目标、背景或格式，但版本差异、采用理由、异议入口和纠正记录较少。出现问题后，团队可能知道哪里不理想，却不知道谁曾批准继续。

什么时候容易出现

同类任务反复重做，大家记得最终版本，却找不到舍弃过哪些建议、为什么采用或由谁放行。

如果一直这样处理

返工原因无法复用，受影响的人也缺少明确的人工复查路径。

可以改变的一步

先为每次重要放行记录版本、采用与舍弃理由、复核人、批准人和纠正入口。

03 真实情境

四种最容易让流畅输出越过控制边界的工作情境

低风险草稿、含敏感资料的输入和会影响他人的决定，需要不同强度的交接与复核。先分级，再决定怎样使用。

01

用 AI 起草一份给客户看的政策说明。

这时先做：先标记为草稿，打开每条关键来源，指出仍未核实的内容，并确认具备相关知识的复核人与批准人。

先出现的信号：材料包含事实和规则，但任务用途、原始来源、适用日期与最终批准人尚未完整记录。

你可能会这样应对：你会提供使用场景和部分标准，但可能较少逐项打开原始来源，也不总由领域人员检查。

容易造成的误读：版式和语气完整可能让接收者以为内容已经完成专业审查。

02

需要把一段客户资料交给工具整理。

这时先做：先暂停输入，确认资料类型和授权；如果任务可以完成，改用占位符、摘要或最少字段，并记录允许范围。

先出现的信号：完整原文最方便，但资料所有者是否允许、当前工具是否获准和哪些字段可以删除仍不清楚。

你可能会这样应对：你有时会使用摘要或较少资料，但敏感识别与授权确认还不稳定。

容易造成的误读：事后删除输入可能被误认为已经解决最初的权限问题。

03

同一种内容任务再次出现，团队准备沿用上次方法。

这时先做：先记录上次有效输入、失败例子、关键差异和转人工位置，再用同一检查标准完成本轮。

先出现的信号：上次采用的输入和最终版本还在，但失败例子、版本差异、检查清单和人工接手点没有被保存。

你可能会这样应对：你能指出结果缺的是目标、背景或格式，却较少比较哪一处修改真正改善结果。

容易造成的误读：团队可能把“上次能用”理解成方法在新资料和异常情况下也稳定。

04

一项 AI 参与的结果会影响员工、客户或其他个人。

这时先做：先确认具体批准人，记录采用理由，并向接收者说明怎样提出异议、由谁人工复查和怎样纠正。

先出现的信号：任务有复核人，但采用理由、最终批准、异议入口和纠正方式还没有同时写清。

你可能会这样应对：你有时能指出复核与批准人，却较少记录采纳与舍弃理由，也没有稳定保留人工纠正路径。

容易造成的误读：“有人看过”可能被当成责任已经明确，受影响的人却不知道怎样请求复查。

04 反应路径

任务分级→上下文交接→生成→证据复核→人工批准→记录

每一步都要能回答谁做、查什么、何时停止、谁负责；自报高分不能证明模型或流程已经安全。

01	输出读起来流畅、结构也完整。	注意力先去哪里 表达质量容易被当成事实与专业判断已经可靠。	你会注意到 注意更容易停在成品感，原始来源、反例和适用范围退到后面。	原来的应对 减少进一步核对，或用另一次生成代替领域人员检查。	继续下去的成本 无法核实的事实可能继续流转，增加重要决定缺少可追溯依据的风险。
新的中断动作 先暂停放行，打开一个关键原始来源，写下一条反例和一项仍未知条件。					
02	完整资料能让整理任务更省事，但权限尚未确认。	注意力先去哪里 先完成再删除看起来像是最省时间的路径。	你会注意到 注意落在任务完成速度，敏感字段、所有者授权和允许工具还没有逐项确认。	原来的应对 可能先输入完整内容，之后再决定是否删除或改用摘要。	继续下去的成本 资料在最初进入工具时已经超过未知边界，导致后续删除也不能改变最初的输入。
新的中断动作 先停止输入，把资料替换成占位符，并确认授权范围和允许使用的工具。					
03	结果出错或受到质疑，需要说明为什么采用。	注意力先去哪里 “这是 AI 生成的”可能被用来解释来源，却没有说明谁选择和批准。	你会注意到 团队能看到最终文本，但版本、舍弃理由、批准人与异议路径不完整。	原来的应对 重新生成一个版本，原来的决定过程和和责任位置继续没有记录。	继续下去的成本 同类错误可能重复，受影响的人也难以找到可以纠正决定的人。
新的中断动作 先记录谁采用、谁复核、谁批准和怎样纠正，再决定是否重新生成。					

05 使用说明

你的 AI 协作说明：可以委派什么，必须由谁核对与批准

可委派的任务

把任务交给我或工具前，请先说明真正使用者、用途和出错影响。

必须补足的上下文

背景或来源不全时，我需要把未知项列出来，不用更流畅的文字填补空白。

证据怎么核对

未经确认的个人、客户、员工和公司敏感资料不进入当前工具；能用摘要或占位符时只给最少信息。

哪些数据不能输入

速度与核验冲突时，高影响结果先停止放行，由具备知识和权限的人决定。

失败怎样留下记录

允许资料、质量标准、原始来源、反例、停止条件、复核人与批准人写清后，协作才可检查。

最终由谁批准

只说“有人看过”、让另一个 AI 再看一遍，或把流畅表达当成可靠，会遮住责任空白。

这份报告建议你优先做的决定

先选一个正在重复使用的任务，按风险写清允许输入的资料、必须核对的来源、停止条件、复核人和最终批准人。

06 行动协议

重要任务开始前，用六步建立可复核的人机 workflow

先分任务风险，再明确允许的资料、输出标准、证据来源、停止条件和人工负责人。

01

标记任务风险

记录结果会影响谁、出错后能否发现和撤回，再区分低、中、高影响任务。

先暂停：若影响范围或出错后果无法说明，停止扩大委派并由实际负责人确认。

02

限定用途与允许资料

写下产出用途、禁止用途、允许输入的资料和可以替换成占位符的字段。

先暂停：如果资料授权、敏感程度或允许工具不清楚，停止输入并联系资料所有者或合规负责人。

03

补足上下文和质量标准

说明使用者、场景、事实限制、时限，并提供一个可接受和一个不可接受例子。

先暂停：若不可接受错误仍无法定义，中断生成并先与使用者确认。

04

核对来源与反例

打开关键原始来源，记录适用日期、范围和一个可能失败的条件。

先暂停：如果关键事实无法独立检验或需要专业判断，停止自动放行并转交具备相应知识的人。

05

指定停止、复核和批准

写下必须停止的异常、负责复核的人、最终批准人和可以否决的条件。

先暂停：若复核人没有足够资料、时间、能力或否决权，停止进入正式使用。

06

记录采用、舍弃与纠正

记录采用和舍弃了哪些建议、理由、版本，以及受影响的人怎样请求人工复查。

先暂停：出现法律、医疗、财务、招聘、安全或其他高影响决定时，停止仅凭自评与生成内容，使用当地规则要求的专业程序。

07 可直接使用

不必只说“用 AI 做一下”，可以直接这样交接、复核和升级

收到一句“用 AI 做一下”的任务。

容易让事情更难 可以，我先生成一版再说。

可以说 “开始前请确认：这份内容用于发散、草稿、核对还是决定参考？哪些资料可以输入，哪些错误不能接受，最后由谁批准？”

这句话的作用 用途、资料边界、质量底线和最终责任在开始前被明确。

输出包含数字、法规或引用，需要进入正式材料。

容易让事情更难 它写得很完整，我再让另一个 AI 检查一下。

可以说 “我先把关键数字和规则逐项连到原始来源，再写一个可能不适用的条件；完成后请具备相关知识的人决定是否采用。”

这句话的作用 原始来源、反例、适用范围和专业复核成为放行条件。

任务需要使用客户或公司资料，但权限不清。

容易让事情更难 我先放进去试，之后删掉就行。

可以说 “这份资料的授权范围还没确认，我不先输入。请确认当前工具是否获准；如果可以用摘要或占位符完成，我只提供完成任务所需的最少字段。”

这句话的作用 停止动作、授权问题、替代方式和最少信息范围都可以被核对。

重要结果准备发布，需要指定责任和纠正路径。

容易让事情更难 大家都看过了，应该可以发。

可以说 “请在发布前写明复核人和最终批准人，并记录采用与舍弃的建议。如果有人受到影响，也请提供人工复查和纠正入口。”

这句话的作用 复核、批准、决定理由和纠正路径不再停留在“大家都看过”。

08 安全边界

这些工作不能只凭自评或模型输出放行

法律、医疗、财务、招聘、安全与其他高影响决定，需要符合当地规则的专业判断、资料权限和独立复核。

01 标记任务风险

若影响范围或出错后果无法说明，停止扩大委派并由实际负责人确认。

02 限定用途与允许资料

如果资料授权、敏感程度或允许工具不清楚，停止输入并联系资料所有者或合规负责人。

03 补足上下文和质量标准

若不可接受错误仍无法定义，中断生成并先与使用者确认。

04 核对来源与反例

如果关键事实无法独立核验或需要专业判断，停止自动放行并转交具备相应知识的人。

流畅不等于正确，人工在环也不自动等于有效控制

如果复核人没有足够能力、证据、时间或否决权，应停止放行，并转交具备权限与专业能力的人处理。

09 已有资源

你已经具备的协作控制，以及用过头时可能增加的成本

01

按出错影响保留人工边界的取舍意识

它有用的时候：任务里既有固定步骤，也有依赖现场、关系或责任的部分。

容易用过头：如果只在脑中知道边界，没有把用途和禁止用途写给接收者，内容仍可能在流转中越界。

别人可能怎么理解：别人可能把草稿或发散材料直接当成已经核准的答案。

这样校正：在产出顶部说明用途、禁止用途、复核状态和最终批准人。

02

补充使用者和场景的表达能力

它有用的时候：需要让产出适合具体对象、渠道和工作目标时。

容易用过头：场景写得很丰富，仍不能替代事实来源、资料权限和不可接受错误。

别人可能怎么理解：一份详细交接可能被误认为所有边界都已覆盖。

这样校正：在场景之后补上原始来源、资料权限、不可接受例子和停止条件。

03

结果不理想时的复盘与定位能力

它有用的时候：需要判断问题来自目标、背景、证据、限制还是格式。

容易用过头：如果每次都临场判断而不保留版本和失败例子，同类任务仍会从头开始。

别人可能怎么理解：频繁重写可能看起来很勤奋，却没有形成可复用控制。

这样校正：把版本差异、失败例子、检查清单和人工接手点同时留下。

10 行动计划

用 30 天验证：哪一个复核或停止条件真正减少返工和未发现错误

这 30 天只验证一件事

在重复任务中加入风险分级、资料最小化、原始来源、停止条件和明确批准人，比仅优化提示更容易减少未记录的放行。

开始前：第1周记录当前任务的交接方式、输入内容、上下文完整性、来源是否可查、复核与批准流程以及出现问题的纠正方式，不假设现有流程已有效。

第 1 周

只记录当前流程，不增加新控制。

要记下来：记录任务次数、输入资料、上下文是否齐全、来源是否打开、复核与批准方式、返工结果。

第 2 周

加入任务风险标记和资料最小化清单。

要记下来：记录风险是否标记、删除或替换了哪些字段、是否因权限不清停止输入及返工变化。

第 3 周

加入原始来源、反例和停止条件。

要记下来：记录来源是否逐项打开、反例是否写下、停止条件是否触发、由谁接手及结果。

第 4 周

加入批准和纠正记录，比较前三周并保留最有效的一项控制。

要记下来：记录复核人和批准人是否明确、采用理由是否保存、纠正入口是否存在、返工结果变化。

可检查交接

第1周有多少任务写清了用途、资料边界、质量标准和禁止用途？

希望看到：相较第1周基线，增加在开始前完成可检查交接的任务数量。

独立核验动作

第1周有哪些关键事实真正打开了原始来源，并记录了反例或适用边界？

希望看到：相较第1周基线，增加打开原始来源、记录反例并转交领域专家审查的次数。

资料最小化

第1周哪些任务使用了完整资料，哪些字段其实可以删除、摘要或替换？

希望看到：相较第1周基线，增加先确认权限并仅使用最少资料的任务比例。

责任可追溯

第1周哪些重要结果能找到复核人、批准人、采用理由和纠正入口？

希望看到：相较第1周基线，增加能追溯到批准人与纠正路径的正式结果数量。

11 保存与分享

任务风险分级卡：同一个 AI 动作，风险不同就不能用同一套流程

把真实任务按影响对象、可逆性、资料敏感度和错误可见性分开，再决定 AI 能做多少、谁必须复核。

01	用 AI 起草一份给客户看的政策说明。	已观察信号：材料包含事实和规则，但任务用途、原始来源、适用日期与最终批准人尚未完整记录。	需要核对：版式和语气完整可能让接收者以为内容已经完成专业审查。	低风险下一步：先标记为草稿，打开每条关键来源，指出仍未核实的内容，并确认具备相关知识的复核人与批准人。
02	需要把一段客户资料交给工具整理。	已观察信号：完整原文最方便，但资料所有者是否允许、当前工具是否获准和哪些字段可以删除仍不清楚。	需要核对：事后删除输入可能被误认为已经解决最初的权限问题。	低风险下一步：先暂停输入，确认资料类型和授权；如果任务可以完成，改用占位符、摘要或最少字段，并记录允许范围。
03	同一种内容任务再次出现，团队准备沿用上次方法。	已观察信号：上次采用的输入和最终版本还在，但失败例子、版本差异、检查清单和人工接手点没有被保存。	需要核对：团队可能把“上次能用”理解成方法在新资料和异常情况下也稳定。	低风险下一步：先记录上次有效输入、失败例子、关键差异和转人工位置，再用同一检查标准完成本轮。
04	一项 AI 参与的结果会影响员工、客户或其他个人。	已观察信号：任务有复核人，但采用理由、最终批准、异议入口和纠正方式还没有同时写清。	需要核对：“有人看过”可能当成责任已经明确，受影响的人却不知道怎样请求复查。	低风险下一步：先确认具体批准人，记录采用理由，并向接收者说明怎样提出异议、由谁人工复查和怎样纠正。

12 保存与分享

上下文交接卡：输入之前先补齐六类信息

好的提示词不只写任务，还要写使用者、事实来源、约束、反例、验收标准和最终责任人。

可委派的任务

把任务交给我或工具前，请先说明真正使用者、用途和出错影响。

必须补足的上下文

背景或来源不全时，我需要把未知项列出来，不用更流畅的文字填补空白。

证据怎么核对

未经确认的个人、客户、员工和公司敏感资料不进入当前工具；能用摘要或占位符时只给最少信息。

哪些数据不能输入

速度与核验冲突时，高影响结果先停止放行，由具备知识和权限的人决定。

失败怎样留下记录

允许资料、质量标准、原始来源、反例、停止条件、复核人与批准人写清后，协作才可检查。

最终由谁批准

只说“有人看过”、让另一个 AI 再看一遍，或把流畅表达当成可靠，会遮住责任空白。

这张工具卡要帮助你完成

你现在可以按任务风险决定哪些部分可委派、哪些资料可输入、谁必须复核批准，以及出现什么情况要停止。

13 保存与分享

证据复核与停止条件：流畅输出不能自动进入下一步

把来源核对、反例检查、专业复核、人工批准和失败升级变成能被团队重复执行的控制点。

01

标记任务风险

记录结果会影响谁、出错后能否发现和撤回，再区分低、中、高影响任务。

停止或升级：若影响范围或出错后果无法说明，停止扩大委派并由实际负责人确认。

02

限定用途与允许资料

写下产出用途、禁止用途、允许输入的资料和可以替换成占位符的字段。

停止或升级：如果资料授权、敏感程度或允许工具不清楚，停止输入并联系资料所有者或合规负责人。

03

补足上下文和质量标准

说明使用者、场景、事实限制、时限，并提供一个可接受和一个不可接受例子。

停止或升级：若不可接受错误仍无法定义，中断生成并先与使用者确认。

04

核对来源与反例

打开关键原始来源，记录适用日期、范围和一个可能失败的条件。

停止或升级：如果关键事实无法独立核验或需要专业判断，停止自动放行并转交具备相应知识的人。

05

指定停止、复核和批准

写下必须停止的异常、负责复核的人、最终批准人和可以否决的条件。

停止或升级：若复核人没有足够资料、时间、能力或否决权，停止进入正式使用。

06

记录采用、舍弃与纠正

记录采用和舍弃了哪些建议、理由、版本，以及受影响的人怎样请求人工复查。

停止或升级：出现法律、医疗、财务、招聘、安全或其他高影响决定时，停止仅凭自评与生成内容，使用当地规则要求的专业程序。

14 专项工具

这些结论从哪些作答来，以及为什么它不能构成安全或合规认证

先核对测量边界和本次原始分，再看看每条结论引用了哪些具体作答。

本次维度

委派边界7 能否按出错影响划分 AI 可做部分、人必须保留部分和输出用途。 原始分 · 仅在当前量表内查看	上下文交接6 能否交代使用者、场景、事实、限制、质量标准与反例。 原始分 · 仅在当前量表内查看	证据复核3 能否追到原始来源、主动找反例并让具备领域知识的人检查。 原始分 · 仅在当前量表内查看
数据边界5 能否识别敏感资料、最小化输入并确认工具与资料授权。 原始分 · 仅在当前量表内查看	迭代控制6 能否定位失败原因、保留版本差异、失败样例与人工接手点。 原始分 · 仅在当前量表内查看	人工定责4 能否指定复核与批准人、记录采纳理由，并保留异议和纠正路径。 原始分 · 仅在当前量表内查看

怎样使用这份结果

依据最近 30 天的 24 道工作情境自述，按六组维度的原始分整理交接、复核、资料与责任线索。

用于为具体任务建立可检查的人机分工、资料边界、核验步骤、停止条件和人工责任。

- 本报告基于虚构的自述答案，未检验真实工作流程、输出质量或合规性。
- 不同任务、工具、组织和当地法规需要不同控制，报告中的做法不能直接复制。
- 报告未检查实际输入、提示、输出日志、权限或政策，不能替代安全、隐私、法律或专业审查。
- 中文版本为试行版，不能用于跨语言、跨国家比较，也不适用于员工考核或晋升。
- 自报的原始分仅反映答题者的自述，不能证明真实工作表现或效率。

不得据此评定 AI 能力、生产力、合规或安全，也不得替代法律、医疗、财务、招聘或安全等高影响领域的专业判断。

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

15 专项工具

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。

	你会根据出错影响和现场知识保留部分人工工作，但输出用途还不总在开始前说清。	第 1 题 第 2 题 第 4 题
也要看到：	第 5 题	
	使用者和场景通常会被交代，但事实限制、质量反例与缺失背景还可能被一轮轮问答代替。	第 5 题 第 7 题 第 8 题
也要看到：	第 6 题	
	原始来源、反例和领域人员检查较少稳定发生，结构完整的输出更容易提前获得信任。	第 9 题 第 10 题 第 12 题
也要看到：	第 13 题	
	你会主动减少部分输入，但敏感识别、资料所有者授权和权限不清时的停止动作仍不稳定。	第 13 题 第 15 题 第 16 题
也要看到：	第 14 题	
	重要任务有时能说清复核和批准人，但采纳理由、异议纠正路径和出错后的个人责任记录较弱。	第 21 题 第 22 题 第 24 题
也要看到：	第 17 题	