

金融 AI 体验 审查清单

用于产品、体验、业务、风险与运营团队共同评审一个金融 AI 场景。30 项问题覆盖任务、信任、人工接管、数据授权与效果测量。

v0.1

2026-07-28

6 页 · 30 项检查

建议使用方式

选择一个具体场景，由跨团队成员独立勾选并记录证据，再共同确认最重要的风险、下一步验证动作和停止条件。本清单用于体验与产品评审，不替代监管、法律或合规判断。

任务与风险

先确认 AI 帮助谁完成什么任务，以及错误会造成什么后果。

☐

01 角色已经明确

客户、员工、审核人员和系统分别承担什么任务与责任？

☐

02 任务边界已经明确

AI 可以完成什么，必须停止在什么位置？

☐

03 风险等级已经说明

错误、遗漏或延迟会造成怎样的业务和用户后果？

☐

04 高风险判断保留人工责任

关键决策是否有明确的人工确认与责任主体？

☐

05 非 AI 替代路径存在

模型不可用时，用户能否继续完成核心任务？

☐

06 上线条件可验证

团队是否定义了可测量的准入指标与停止条件？

本节最重要的风险	
已有证据	
下一步验证动作	

信息、解释与信任

让用户知道正在发生什么、依据是什么，以及自己还能做什么。

- ☐ 07 AI 身份和能力被清楚告知
用户能否判断当前是在与 AI、人工还是混合服务交互？
- ☐ 08 结果依据可理解
重要建议是否说明使用了哪些信息及其适用边界？
- ☐ 09 不确定性被诚实表达
系统是否避免把推测包装成确定事实？
- ☐ 10 关键确认突出真正风险
金额、对象、期限、费用和后果是否容易核对？
- ☐ 11 状态持续可见
等待、审核、转交和失败状态是否说明下一步与时间预期？
- ☐ 12 文案没有转移责任
提示是否帮助用户理解，而不是只用免责声明保护系统？

本节最重要的风险	
已有证据	
下一步验证动作	

人工接管与失败恢复

把转人工和恢复设计成正常路径，而不是上线后的临时补丁。

- ☐ 13 转人工条件明确
- 哪些意图、置信度或风险条件会触发人工接管？
- ☐ 14 用户可以主动求助
- 用户是否随时能找到人工入口，而不必反复改写问题？
- ☐ 15 上下文随服务链传递
- 人工是否能看到必要的对话、材料、状态和失败原因？
- ☐ 16 已完成步骤得到保留
- 转接或失败后，用户是否需要重复提交相同信息？
- ☐ 17 错误信息可行动
- 提示是否说明发生了什么、如何修复以及何时再试？
- ☐ 18 投诉与补救路径存在
- 影响资金或权益时，是否有清晰的申诉、补偿和追踪机制？

本节最重要的风险	
已有证据	
下一步验证动作	

数据、授权与治理

让数据使用、权限控制与责任链在界面和服务流程中保持一致。

- ☐ 19 只采集任务所需数据
- 是否能够删除与当前任务无关的字段或材料？
- ☐ 20 授权具体且可撤回
- 用户是否知道授权对象、用途、期限与关闭方式？
- ☐ 21 敏感信息得到保护
- 前端、日志、通知和导出是否避免暴露敏感数据？
- ☐ 22 权限与角色一致
- 员工助手是否只显示岗位完成任务所需的信息？
- ☐ 23 第三方责任可说明
- 模型、数据、平台和机构之间的责任是否有明确边界？
- ☐ 24 变更与版本可追踪
- 模型、提示、规则和界面变化是否留下审核与发布记录？

本节最重要的风险	
已有证据	
下一步验证动作	

效果测量与持续改进

不要只看调用量，持续判断完整服务是否变得更好。

- ☐ 25 任务完成得到测量
是否记录成功完成、放弃、重复操作和人工补救？
- ☐ 26 质量按场景拆分
是否区分角色、任务、风险等级和渠道，而不是只看平均值？
- ☐ 27 客户与员工体验同时评估
前台效率提升是否给后台员工增加了返工或判断负担？
- ☐ 28 异常样本进入复盘
错误回答、失败转接和投诉是否形成可追踪的问题清单？
- ☐ 29 用户反馈能够闭环
纠错和建议是否有负责人、状态、处理期限与回访方式？
- ☐ 30 停止与回滚可执行
达到什么条件时降级、停用或回滚，谁有权限决定？

本节最重要的风险	
已有证据	
下一步验证动作	

方法边界与参考：本清单为无界创新原创设计评审工具，不替代监管、法律或合规意见。参考公开资料包括国家金融监督管理总局《银行业保险业人工智能应用管理办法》、NIST Human-Centered AI 研究框架及中国农业银行智能助手服务协议。