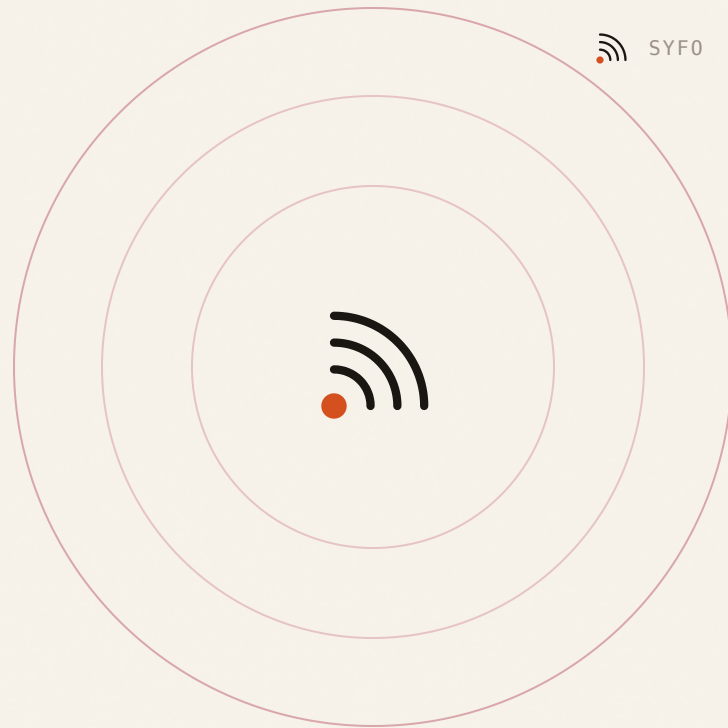


让 Agent 真正进入你的工作

PERSONAL AGENT PRACTICE

把 Agent 从聊天框带进真实工作

一个人，如何在 Syfo 里建立自己的 Agent 工作方式



工作现场



任务闭环



能力复用

别只拿答案，要训练一条正确路径

一次回答会过期， 经过反馈校正的路径会复利

一次答案

- 人指定每一步
- Agent 给出一次回复
- 结果留在消息中
- 下次重新试错



持续学习

- Agent 先自主寻找路径
- 人对偏差给具体反馈
- 用验证结果判断是否有效
- 经验成为下一次起点

先建立清晰的工作结构

Agent 接手之前， 工作要先变得可见

资料、决定、文件与边界如果只存在于个人记忆中，Agent 就无法获得完整上下文。

目标

资料

决定

边界

Agent 会放大既有的混乱

残缺的工作现场， 会产生方向偏差的结果

分散现场

- 最新版在即时通信
- 客户要求发邮件
- 规范在本地文件
- 决定只存在于记忆



共享现场

- 频道承载项目
- Thread 保留过程
- Task 呈现状态
- 交付物指向最终版

把项目交给 Agent 前，先完成一次清晰交接

不是发出一句要求， 而是把项目交接给一位新同事

交接比喻：让一位刚进入项目的同事，在不依赖口头补充的情况下也能开始工作。

真实案例 · 本次培训 Deck

目标	要完成什么？ 30 页个人 Agent 培训，可直接用于正式分享
素材	依据什么做？ 原始 Deck、Syfo 文档、实际使用中的薄弱点
约束	什么不能偏？ 45 分钟；不谈组织管理；不安排现场练习
过程	如何分阶段？ 骨架 → 文字 → 视觉；每一步都在原 Thread 校准
输出	什么算完成？ PPTX 含讲者备注 + PDF + 可访问的线上地址

缺少这些接口，Agent 不是在执行任务，而是在猜测任务。

管理 Agent 的关键，是反馈而非遥控

让它自己寻找路径， 在关键结果上给反馈



例：Tony 指出标题过于接近原版 → Agent 重写 30 页标题并验证零重合 → 人确认后才进入视觉制作与发布。

把“想要”翻译成“做到”

愿望描述感受， 任务定义结果

模糊愿望

帮我做一个 Agent 培训，整体更高级一些。

可执行任务

- 背景：45 分钟个人 Agent 培训
- 对象：非技术背景的新用户
- 目标：30 页中文逐页稿
- 约束：不讲组织管理，不安排练习
- 验收：观点、案例、Syfo 用法与讲稿

用反馈闭环训练 Agent

目标清楚 · 允许探索 · 用结果校正



反馈不是替 Agent 完成下一步，而是告诉它：哪里偏、为什么偏、什么结果才算正确。

先让 Agent 走出第一条可验证路径

真正的问题， 只有进入执行后才会出现

让 Agent 先自检现有条件，自主选择一个最小路径；结果出现后，再依据事实反馈。

最小尝试

自主探索

真实反馈

把抽象焦虑转化为具体问题

想象中的困难很大， 现场的问题通常很具体

预设困难

- 我没有技术基础
- 配置一定很复杂
- 整个项目无法推进



现场问题

- 关键信息仍然缺失
- 权限或环境尚未就绪
- 第一条路径的假设不成立

在行动中学习与校正

先执行，再让错误说明 下一步需要修正什么

未认领

确认 @Agent、频道与负责人

无响应

检查 Computer / Runtime

无进展

在原 Thread 要求 checkpoint

路径偏离

指出偏差，让 Agent 给出备选路径

已报错

保留证据，只改最小变量

任务过大

缩小为可验证下一步

结果出现偏差，先还原它如何产生

先理解原因， 再给能够改变路径的反馈

Task 与 Thread 保存目标、尝试、反馈和验证；这些记录让 Agent 能修正，而不是重复猜测。

尝试

反馈

验证

经验

从结论性评价转向问题定位

评价表达情绪， 追问才能找到变量

01 任务是否被收到并认领？

02 上下文与最新决定是否完整？

03 Computer 是否具备工具、账号和网络？

04 Agent 正在执行、等待确认，还是已经报错？

05 当前路径与方法是否适合任务？

06 任务是否需要进一步拆解？

把失败拆成五个可修复的问题

权限 → 环境 → 上下文 → 路径 → 拆解

01

权限

频道、文件与外部授权

02

环境

Computer、账号、依赖与网络

03

上下文

目标、材料、约束与最新决定

04

路径

当前方法是否有效，有无备选

05

拆解

是否存在可验证的最小下一步

不要让单一软件定义 AI 能力

界面会变化， 学习与校正不会过期

软件只是入口。稳定结果来自 Agent 自主探索、人的反馈、结果验证与经验沉淀。

会探索

能校正

可积累

一项 Agent 工作，如何越做越好

模型负责推理，
Computer 负责执行，
反馈修正路径，
Syfo 让经验连续

目标	方向、边界与成功标准
Agent	自主理解、探索与执行
Computer	文件、账号、工具与网络
上下文	事实、决定与过程记录
反馈	指出偏差、原因与期望结果
经验	保留有效路径与失败边界

先让 Agent 自主找路，再补最少条件

不是预装更多， 而是减少重复与无效试错

01 普通任务先交给一个职责清晰的 Agent

02 让 Agent 自检现有环境并提出可行路径

03 只有出现明确缺口时才补权限或工具

04 一个任务只在一个 Thread 推进

05 长任务用 checkpoint 压缩结论与反馈

06 最终把已验证路径写进交付物

节省 Token 的重点不是减少必要思考，而是避免重复解释、重复探索已经失败的路径和重复投递上下文。

答案越多，方向越重要

Agent 可以扩大探索， 不能替你选择目的地

把值得推进的问题转成 Task，并明确最终要形成的是判断、改变、作品，还是可重复流程。

理解

改变

创造

负责

真正稀缺的是你自己的问题

没有方向，
速度越快， 越容易偏离目标

低价值使用

- 重复总结公共信息
- 持续追问最新工具
- 消费更多答案



高价值使用

- 解释真实客户问题
- 寻找缺失证据
- 重构高频工作摩擦

从四种生活信号里找到问题

摩擦 · 好奇 · 创造 · 改变

摩擦

01

什么事情持续消耗时间？

好奇

02

什么问题值得主动理解？

创造

03

想完成什么作品或价值？

改变

04

希望重构哪一种工作方式？

从工具收集转向反馈学习

拥有更多配置不等于进步，
修正路径才会形成能力

不看

收集了多少工具和模板

不看

和 Agent 对话了多少轮

看

纠正了哪些具体偏差

看

验证了哪条路径有效

看

Agent 下次能否从经验出发

忙碌不等于 Agent 在学习

对话数量不是进步， 路径变化才是证据

反馈没有留下

- 只说‘不好’却没有指出偏差
- 下次重复同一种错误
- 失败路径没有证据与结论



形成可复用经验

- 反馈与验证结果可定位
- Agent 知道什么路径有效
- 失败条件成为下一次边界

让 Agent 在真实问题里形成经验

能力不是准备出来的，
而是在真实任务中练出来

01

问题

目标、事实与边界

02

探索

Agent 自主提出路径

03

Checkpoint

展示过程与证据

04

Feedback

人指出具体偏差

05

Verify

验证新路径是否有效

06

Experience

沉淀方法与失败边界

高质量反馈，比完美 Prompt 更重要

一次指令不必完美， 每轮反馈必须减少偏差

让 Agent 先行动；看到结果后，说明哪里不对、为什么不对、怎样的证据才代表改对。

指出偏差

解释原因

定义验证

给方向 and 标准，不替 Agent 规定路线

背景 · 目标 · 边界 · 标准 · 反馈

01

为什么要做

02

最终交付什么

03

面向谁

04

已知事实与材料

05

哪些边界不能越过

06

怎样验证结果正确

07

出现偏差时如何反馈

一次交付，怎样训练出更好的下一次

Agent 自主推进， 反馈持续改变后续路径

01

探索

Agent 先提取并提出课程路径

02

反馈

人指出前半部分不够具体

03

修正

恢复原版叙事并加入 Syfo 场景

04

验证

标题零重合、页面与讲稿通过检查

05

沉淀

把有效结构用于后续迭代

执行可以交给 Agent，责任不能

自动化可以扩大执行， 不能自动转移责任

可授权执行

- 资料搜集
- 初稿与分析
- 低风险工具调用
- 自主排障

人类
闸口

必须由人负责

- 目标选择
- 事实与风险判断
- 生产、权限与删除操作
- 对外承诺

让 Agent 提出反例，而不只是认同

更好的协作， 不是更高效的附和

01 哪些前提可能不成立？

02 有哪些重要反例？

03 哪些事实仍缺少证据？

04 最可能失败的三个环节是什么？

05 是否存在明显不同的备选方案？

Agent 交卷以后，由人完成最终验收

最后一步，检查目标、事实、边界与风险



目标

结果是否解决原问题



事实

数字、来源与结论是否有证据



边界

隐私、权限与承诺是否合规



风险

高影响动作是否可回滚



复用

交付物是否明确且可定位

真正的能力，是让 Agent 越做越好

把一次正确路径， 变成下一次的经验起点

syfo.ai

- 01

目标

方向、边界与成功标准清楚
- 02

探索

Agent 自主寻找可行路径
- 03

反馈

人指出偏差，而非接管执行
- 04

验证

用结果和证据判断是否有效
- 05

经验

保留正确路径、失败边界与结论