

2025 第二届信息技术应用创新大赛

基于大模型教育管理应用创新赛

参赛指南

大赛组委会

2025 年 5 月

目录

1. 易思捷 AI 智能体管理平台简介	1
2. 易思捷 AI 智能体管理平台登录	1
3. 创建团队账号	2
4. 知识库	5
(1) 简介	5
(2) 知识库创建	7
(3) 知识库分段说明	8
(4) 检索设置	10
(5) 召回测试	13
(6) 新增知识库文档	13
5. 项目开发	14
(1) 创建智能体说明	15
(2) 多模型调试	18
(3) 使用模板创建	20
6. 工作流	21
(1) 定义介绍	21
(2) 适用场景	21
(3) workflow 节点说明	22
(4) 使用说明	25
7. 插件（工具）	27
(1) 定义	27
(2) 作用	27
(3) 如何创建自定义工具	27
(4) 如何使用工具	28

基于大模型教育管理应用创新赛参赛指南

——平台使用手册

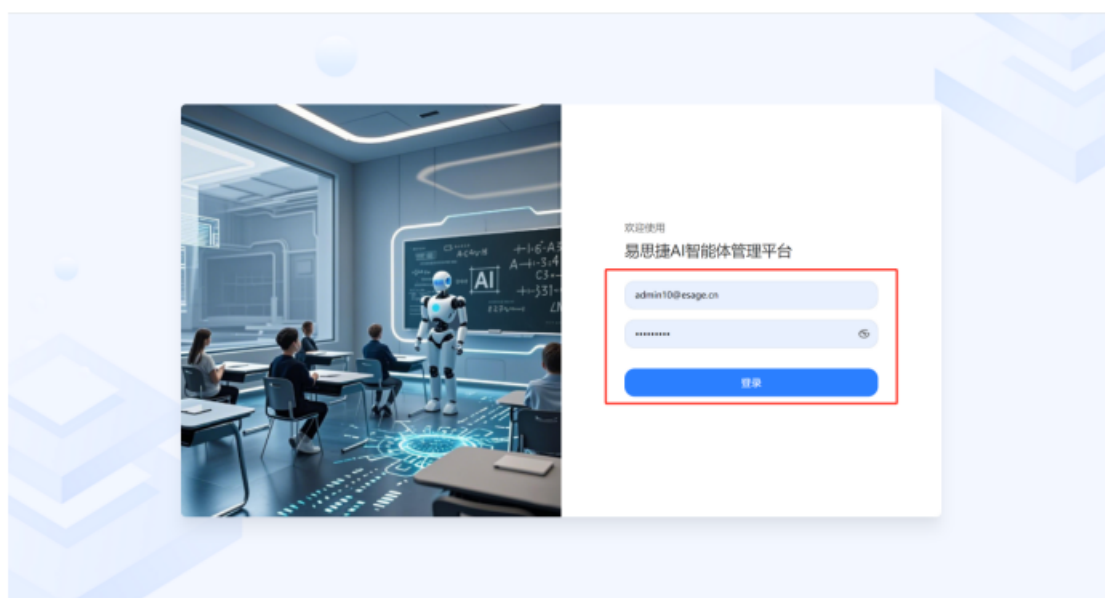
1. 易思捷 AI 智能体管理平台简介

易思捷 AI 智能体管理平台是一款大语言模型(LLM) 应用开发平台。它融合了后端即服务和 LLMOps 的理念，使开发者可以快速搭建生产级的生成式 AI 应用。即使你是非技术人员，也能参与到 AI 应用的定义和数据运营过程中。

由于平台内置了构建 LLM 应用所需的关键技术栈，包括对数百个模型的支持、直观的 Prompt 编排界面、高质量的 RAG 引擎、稳健的 Agent 框架、灵活的流程编排，并同时提供了一套易用的界面和 API。这为开发者节省了许多重复造轮子的时间，使其可以专注在创新和业务需求上。

2. 易思捷 AI 智能体管理平台登录

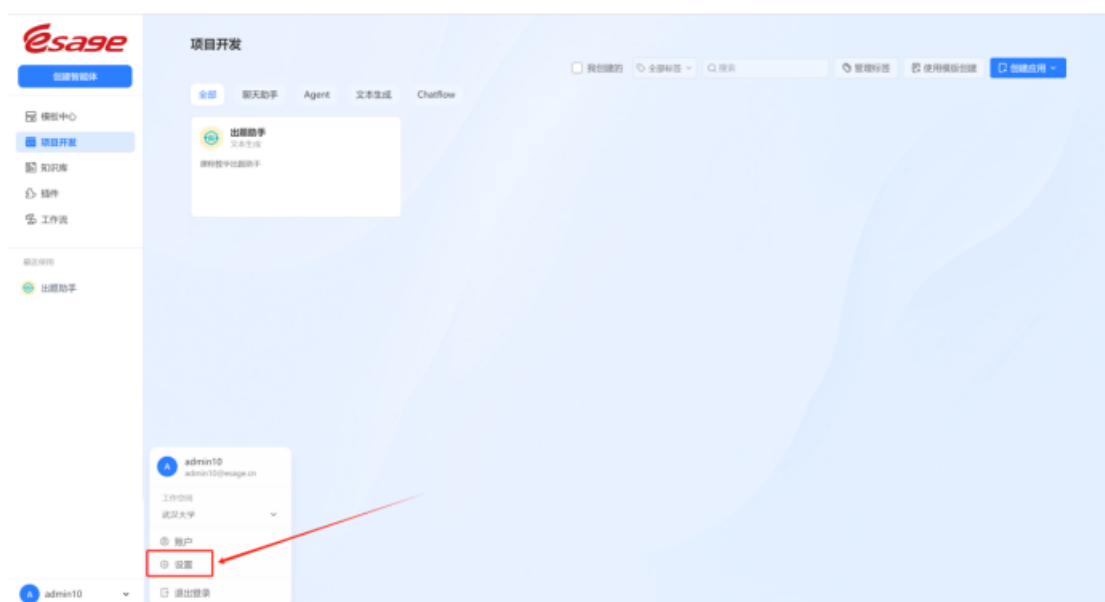
1、在浏览器地址栏输入：<https://ip>（竞赛环境暂未部署，后续补充地址），进入易思捷 AI 智能体开发平台登录界面。



2、输入易思捷为学校提供的账号密码登录平台。

3. 创建团队账号

1、点击左下角账号-设置



添加学校团队成员：



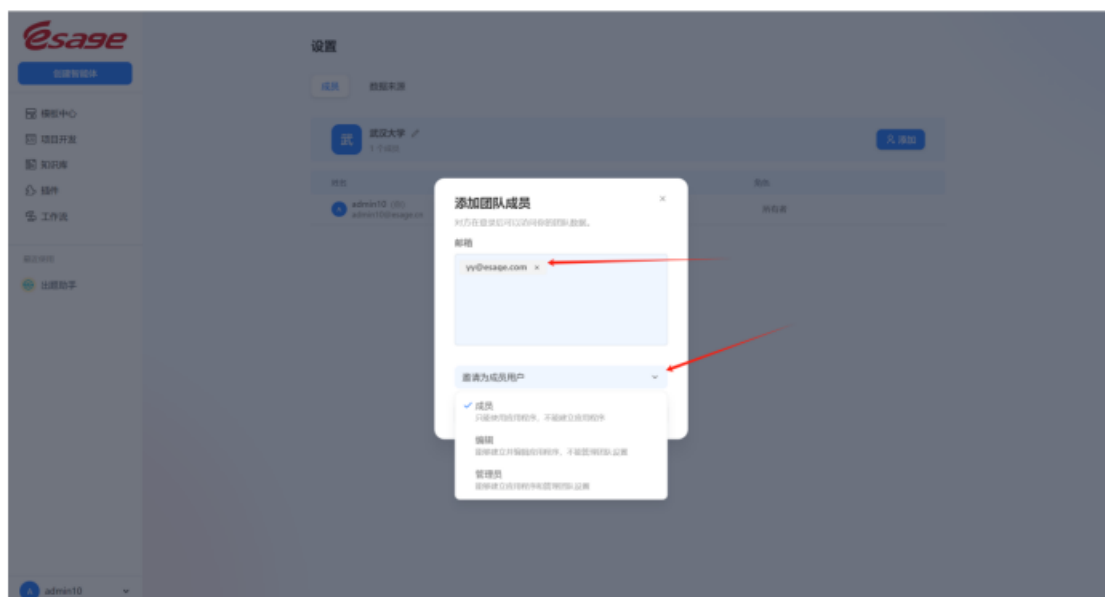
填写邀请团队成员邮箱并设置成员权限

权限说明：

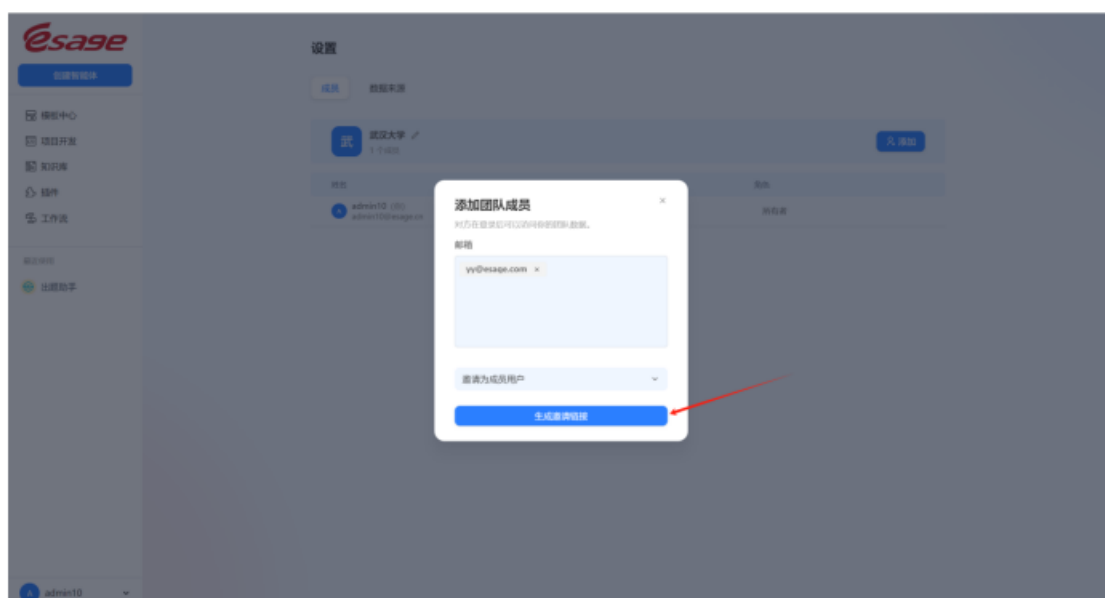
管理员：具备团队管理权限和智能体创建权限。

编辑：具备智能体创建/编辑权限，不能管理团队。

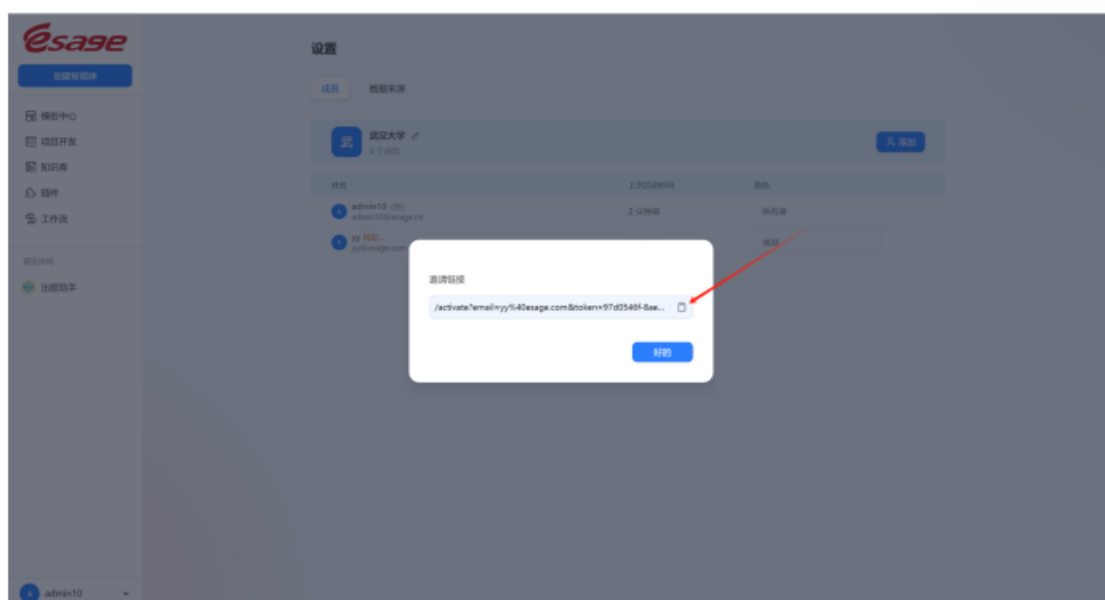
成员：仅能使用团队内的智能体应用，不具备创建权限。



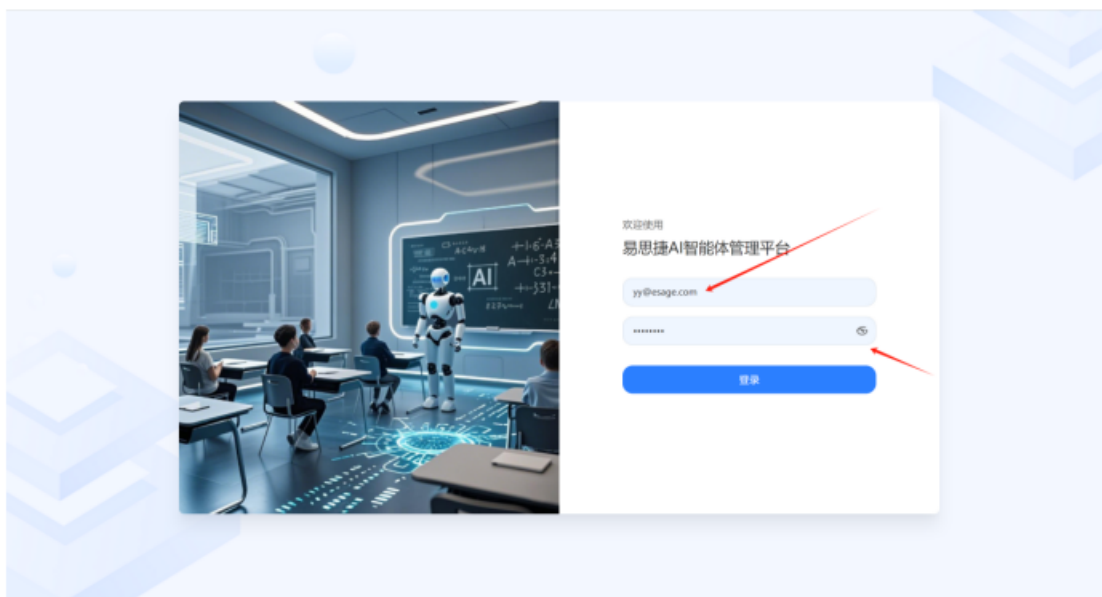
点击“生成邀请链接”按钮，创建平台邀请链接。



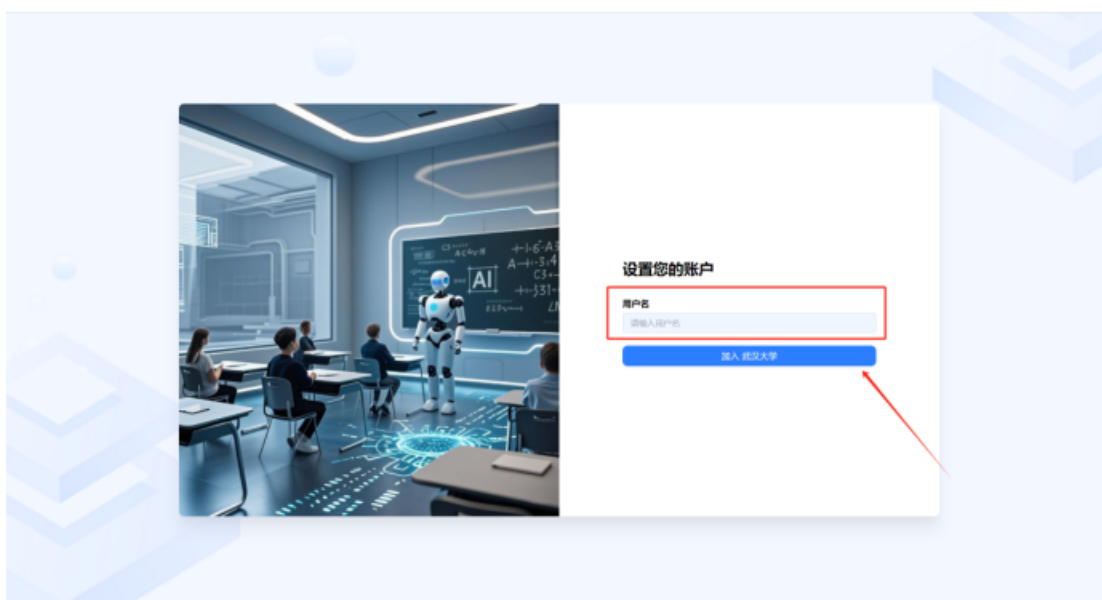
点击复制链接，将链接发送给团队成员，成员即可使用链接注册登录。



团队成员通过链接访问，输入自己账号密码即可登录



被要求的用户初次登录时，需要填写自己的用户名，填写后点击加入团队即可。



4. 知识库

(1) 简介

知识库功能将 RAG 管线上的各环节可视化，提供了一套简单易用的用户界面来方便应用构建者管理个人或者团队的

知识库，并能够快速集成至 AI 应用中。

开发者可以将企业内部文档、FAQ、规范信息等内容上传至知识库进行结构化处理，供后续 LLM 查询。

相比于 AI 大模型内置的静态预训练数据，知识库中的内容能够实时更新，确保 LLM 可以访问到最新的信息，避免因信息过时或遗漏而产生的问题。

LLM 接收到用户的问题后，将首先基于关键词在知识库内检索内容。知识库将根据关键词，召回相关度排名较高的内容区块，向 LLM 提供关键上下文以辅助其生成更加精准的回答。

开发者可以通过此方式确保 LLM 不仅仅依赖于训练数据中的知识，还能够处理来自实时文档和数据库的动态数据，从而提高回答的准确性和相关性。

核心优势：

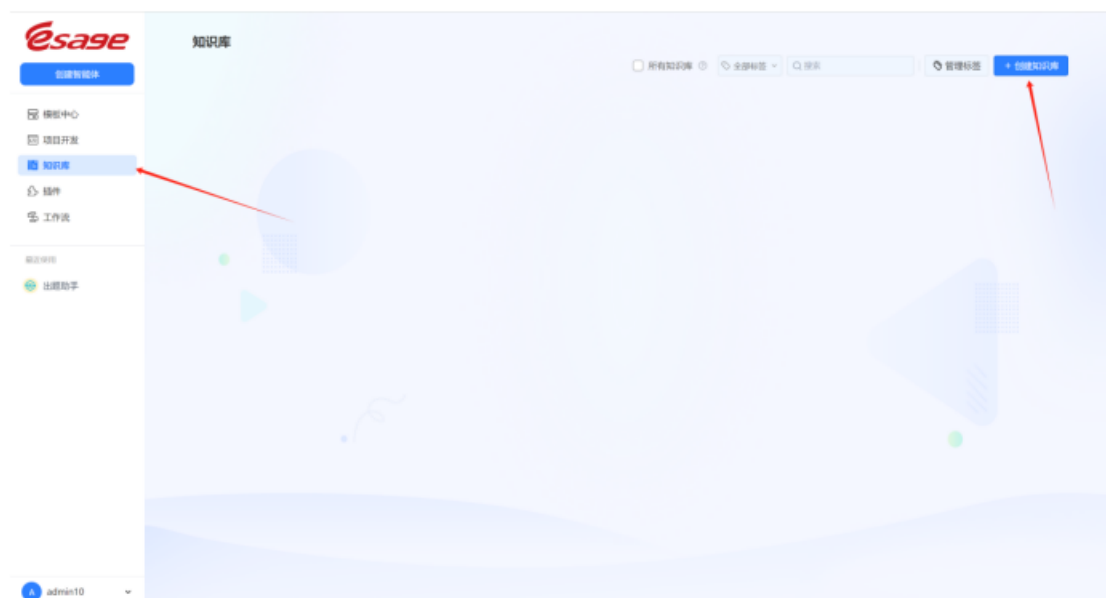
实时性：知识库中的数据可随时更新，确保模型获得最新的上下文。

精准性：通过检索相关文档，LLM 能够基于实际内容生成高质量的回答，减少幻觉现象。

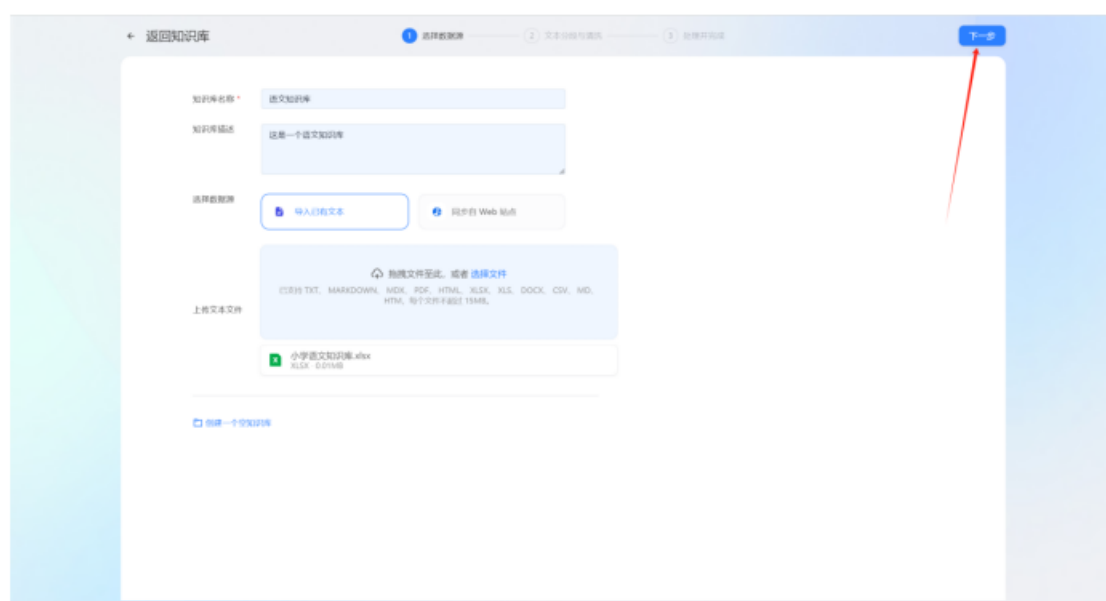
灵活性：开发者可自定义知识库内容，根据实际需求调整知识的覆盖范围。

(2) 知识库创建

登录平台后，选择知识库菜单，点击创建知识库即可进行知识库创建。



创建知识库时，填写知识库名称、描述、选择知识库中的知识来源并导入知识文本，内容填写完成后进入下一步。



(3) 知识库分段说明

分段说明：

将内容上传至知识库后，接下来需要对内容进行分段与数据清洗。该阶段是内容的预处理与数据结构化过程，长文本将会被划分为多个内容分段。

什么是分段与清晰策略：

LLM 收到用户问题后，能否精准地回答出知识库中的内容，取决于知识库对内容块的检索和召回效果。匹配与问题相关度高的文本分段对 AI 应用生成准确且全面的回应至关重要。

好比在智能客服场景下，仅需帮助 LLM 定位至工具手册的关键章节内容块即可快速得到用户问题的答案，而无需重复分析整个文档。在节省分析过程中所耗费的 Tokens 的同时，提高 AI 应用的问答质量。

分段模式：

知识库支持两种分段模式：通用模式与父子模式。如果你是首次创建知识库，建议选择父子模式。

- 通用模式：

系统按照用户自定义的规则将内容拆分为独立的分段。当用户输入问题后，系统自动分析问题中的关键词，并计算关键词与知识库中各内容分段的相关度。根据相关度排序，选取最相关的内容分段并发送给 LLM，辅助其处理与更有效地回答。

在该模式下，你需要根据不同的文档格式或场景要求，参考以下设置项，手动设置文本的分段规则。

- 父子模式：

与通用模式相比，父子模式采用双层分段结构来平衡检索的精确度和上下文信息,让精准匹配与全面的上下文信息二者兼得

其中，父区块（**Parent-chunk**）保持较大的文本单位（如段落），提供丰富的上下文信息；子区块（**Child-chunk**）则是较小的文本单位（如句子），用于精确检索。系统首先通过子区块进行精确检索以确保相关性，然后获取对应的父区块来补充上下文信息，从而在生成响应时既保证准确性又能提供完整的背景信息。你可以通过设置分隔符和最大长度来自定义父子区块的分段方式。

例如在 AI 智能客服场景下，用户输入的问题将定位至解决方案文档内某个具体的句子，随后将该句子所在的段落或章节，联同发送至 LLM，补全该问题的完整背景信息，给出更加精准的回答。

分段操作说明：

对上传的知识库文档进行分段设置，选择通用即可，设置分段后可以点击预览，查看分段效果。



(4) 检索设置

知识库在接收到用户查询问题后，按照预设的检索方式在已有的文档内查找相关内容，提取出高度相关的信息片段供语言模型生成高质量答案。这将决定 LLM 所能获取的背景信息，从而影响生成结果的准确性和可信度。

在第二步中同步设置知识库的索引方式和设置

- 索引方式：高质量、经济（推荐高质量，索引更精准）
- 检索设置：

向量检索：向量化用户输入的问题并生成查询文本的数学向量，比较查询向量与知识库内对应的文本向量间的距离，寻找相邻的分段内容。

全文检索：关键词检索，即索引文档中的所有词汇。用户输入问题后，通过明文关键词匹配知识库内对应的文本片段，返回符合关键词的文本片段；类似搜索引擎

中的明文检索。

混合检索:同时执行全文检索和向量检索,或 Rerank 模型,从查询结果中选择匹配用户问题的最佳结果。

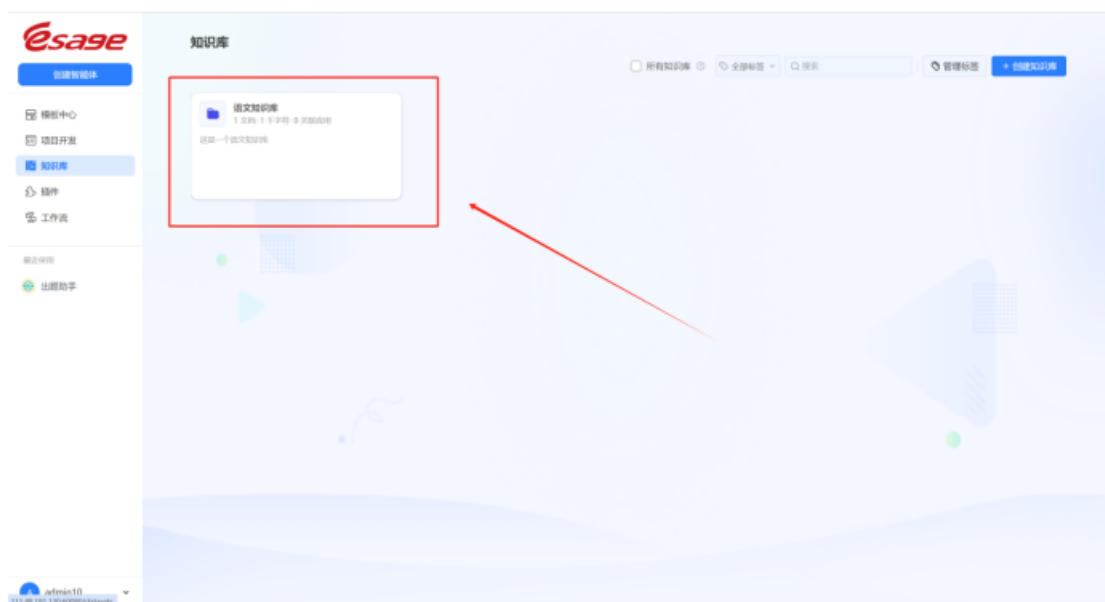


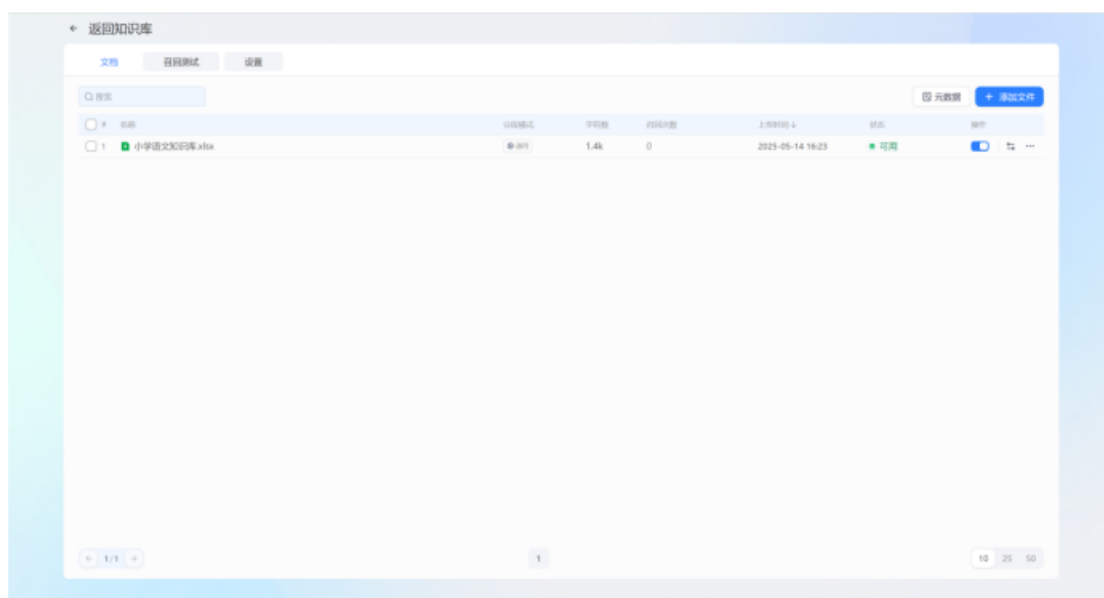
点击保存并处理即可完成知识库创建





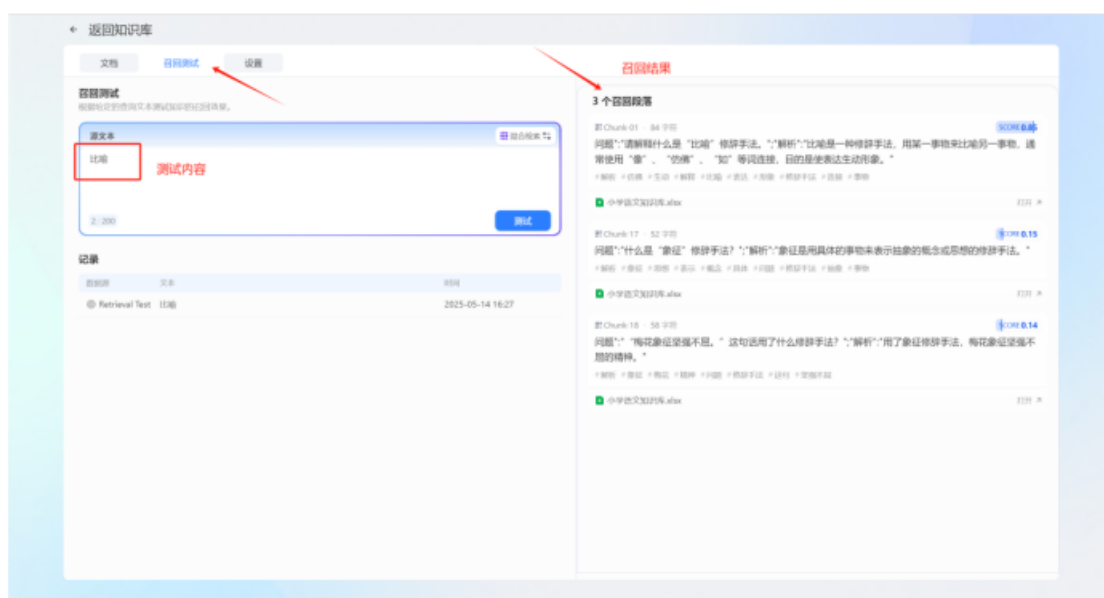
点击知识库卡片即可进入知识库，查看知识库详情信息。





(5) 召回测试

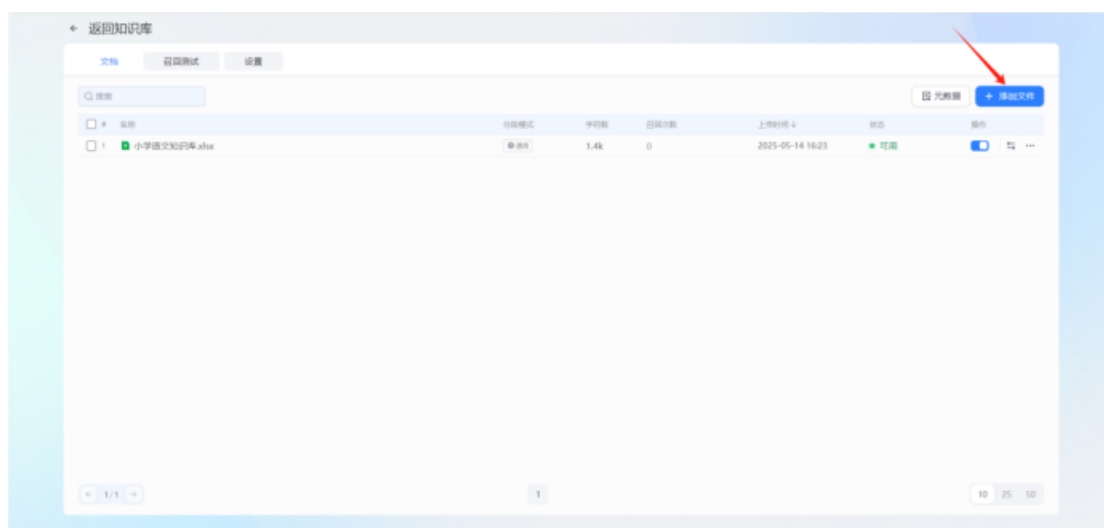
在知识库详情中，点击召回测试，输入召回测试内容，即可查看召回测试结果



(6) 新增知识库文档

在知识库详情中，点击添加文件，即可执行为知识库添加

文件。



5. 项目开发

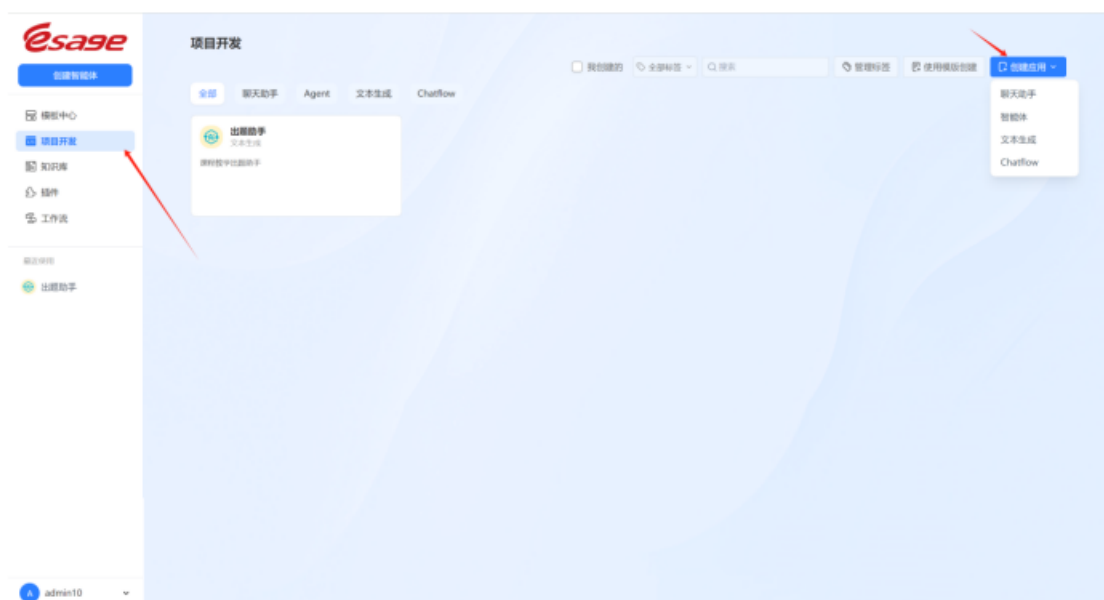
登录平台后，进入项目开发页面，点击创建应用，就可以开始创建智能体。

聊天助手：对话型应用采用一问一答模式与用户持续对话。通常应用在客户服务、在线教育、医疗保健、金融服务等领域。这些应用可以帮助组织提高工作效率、减少人工成本和提供更好的用户体验。

智能体：智能助手，利用大语言模型的推理能力，能够自主对复杂的人类任务进行目标规划、任务拆解、工具调用、过程迭代，并在没有人类干预的情况下完成任务。

文本生成：适用于定义等复杂流程的多轮对话场景，具有记忆功能的应用编排方式

Chatflow：适用于自动化、批处理等单轮生成类任务的场景的应用编排方式



(1) 创建智能体说明

以创建智能体举例：

关键设置说明：

对话开场白：用于配置智能体和用户交流的第一句开场白对话，同时可以通过此配置内置开场白问题让用户更快了解智能体的能力。

下一步问题建议：设置下一步问题建议可以在每次对话交互后，让 AI 根据之前的对话内容继续生成 3 个提问，引导下一轮对话。

语音转文字：开启后，支持将用户的语音转换成文字进行输入。

引用和归属：开启功能后，当 LLM 引用知识库内容来回答问题时，可以在回复内容下面查看到具体的引用段落信息，包括原始分段文本、分段序号、匹配度等。

内容审查：我们在与 AI 应用交互的过程中，往往在内容安全性，用户体验，法律法规等方面有较为苛刻的要求，此时我们需要“敏感内容审查”功能，来为终端用户创造一个更好的交互环境。

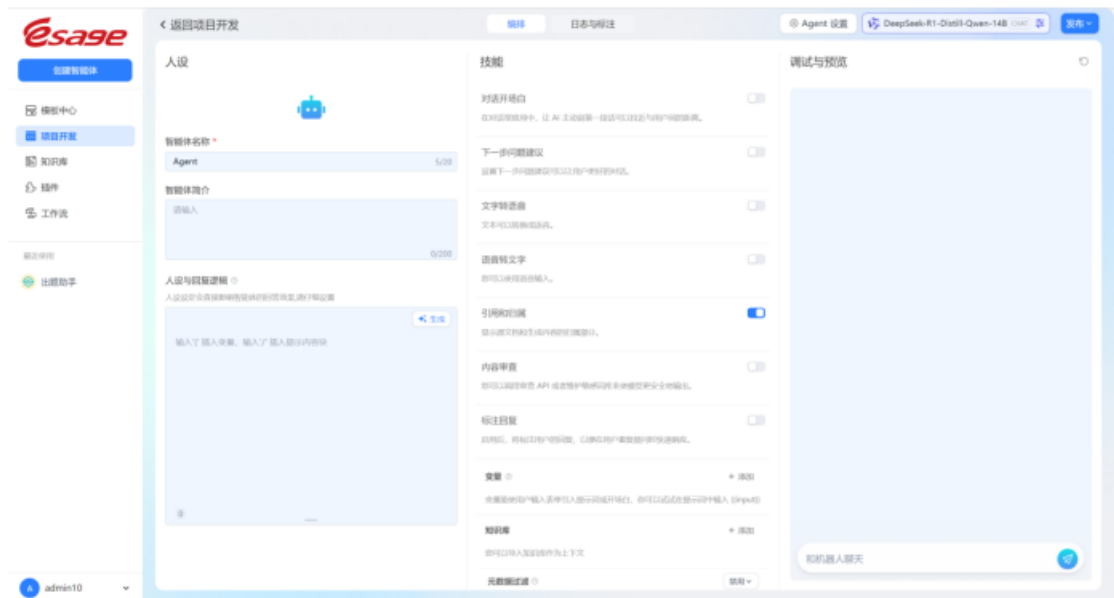
标注回复：可以标注智能体的历史回答记录，标注后，如果用户重复提问，则可以直接使用标注内容快速回答。

变量：变量可以使用户输入表单引入提示词或者开场白

知识库：可以通关关联知识库，当用户和智能体交流时，大模型可以检索知识库内容进行回答。

元数据过滤：元数据过滤式是使用元数据属性（例如：标签、类别、访问权限）来细化和控制系统内相关信息的检索过程。

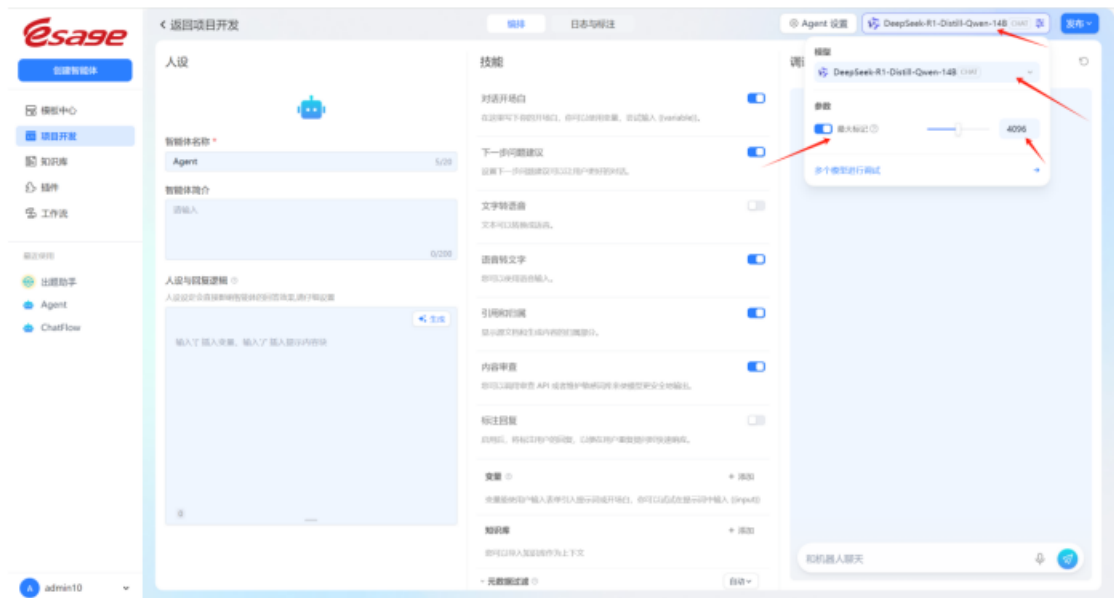
工具：在“工具”中，你可以添加需要使用的工具。工具可以扩展 LLM 的能力，比如联网搜索、科学计算或绘制图片，赋予并增强了 LLM 连接外部世界的能力。比如关联一个文本转换工具，智能体即可实现将文本格式进行转换。也可以关联一个工作流，如关联一个问题分类的工作流，可以将用户提出的问题进行分类，不同分类走向不同的处理节点。平台提供了两种工具类型：内置工具和自定义工具。



智能体编排做完后即可直接进行调试预览



在智能体编排页面中，点击大模型，即可选择不同模型来调试智能体。



(2) 多模型调试

在调试过程中可以选择多个模型同时进行调试，以便直观查看不同模型的使用效果。





点击发-运行即可查看智能体独立运行的效果

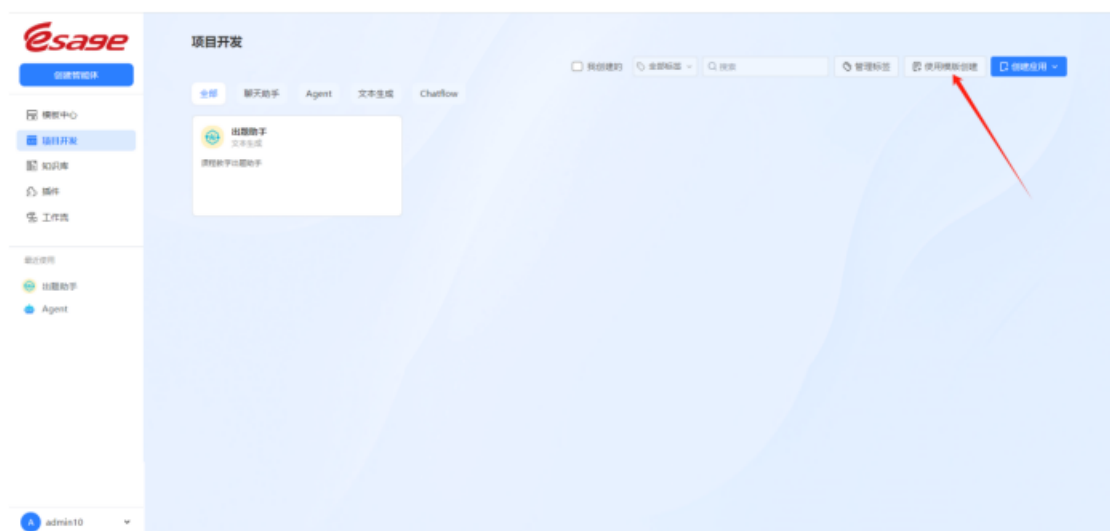


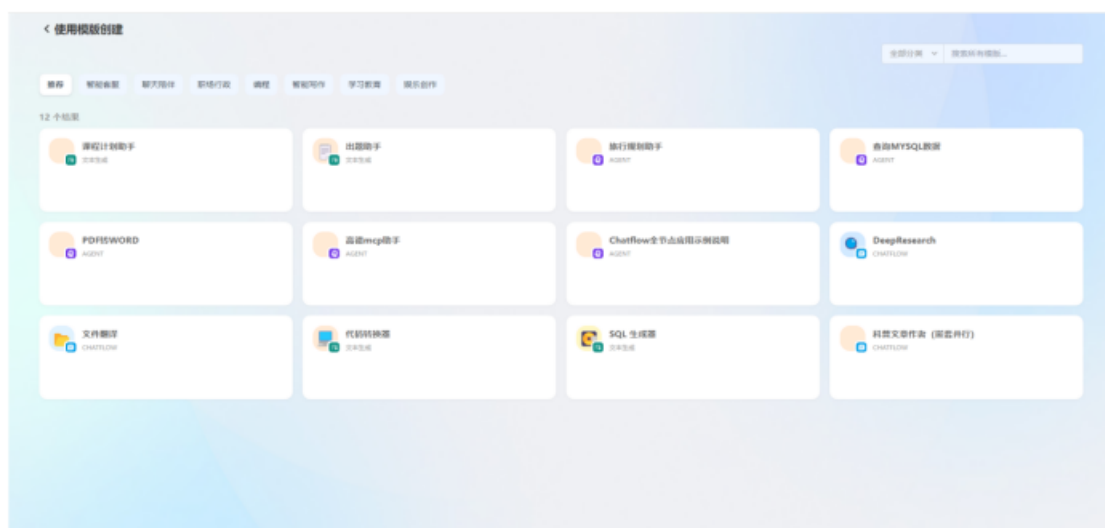


(3) 使用模板创建

当你在不知道如何创建一个智能体时，可以使用平台内置的模板，来快速创建一个智能体，通过模板来发散思维进行创作。

通过点击界面的使用模板创建按钮，进入模板选择界面，通过模板实现快速创建智能体。





6. 工作流

(1) 定义介绍

工作流通过将复杂的任务分解成较小的步骤（节点）降低系统复杂度，减少了对提示词技术和模型推理能力的依赖，提高了 LLM 应用面向复杂任务的性能，提升了系统的可解释性、稳定性和容错性。工作流主要面向自动化和批处理情景，适合高质量翻译、数据分析、内容生成、电子邮件自动化等应用程序。

(2) 适用场景

面向自动化和批处理情景，适合高质量翻译、数据分析、内容生成、电子邮件自动化等应用程序。该类型应用无法对生成的结果进行多轮对话交互。

常见的交互路径：给出指令 → 生成内容 → 结束

(3) workflow 节点说明

开始：开始节点是个工作流必备的预设节点，为后续 workflow 节点以及应用的正常流转提供必要的初始信息，例如应用使用者所输入的内容、以及上传的文件等。

LLM：调用大语言模型的能力，处理用户在“开始”节点中输入的信息（自然语言、上传的文件或图片），给出有效的回应信息。

知识检索：从知识库中检索与用户问题相关的文本内容，可作为下游 LLM 节点的上下文来使用。

问题分类：通过定义分类描述，问题分类器能够根据用户输入，使用 LLM 推理与之相匹配的分类并输出分类结果，向下游节点提供更加精确的信息。

条件分支：根据 `If/else/elif` 条件将流程拆分成多个分支。

代码执行：代码节点支持运行 `Python/NodeJS` 代码以在工作流程中执行数据转换。它可以简化你的工作流程，适用于 `Arithmetic`、`JSON transform`、文本处理等情景。

模板转换：模板节点允许你借助 `Jinja2` 这一强大的 `Python` 模板语言，在工作流内实现轻量、灵活的数据转换，适用于文本处理、`JSON` 转换等情景。例如灵活地格式化并合并来自前面步骤的变量，创建出单一的文本输出。这非常适合于将多个数据源的信息汇总成一个特定格式，满足后续步骤的需求。

文档提取器：LLM 自身无法直接读取或解释文档的内容。因此需要将用户上传的文档，通过文档提取器节点解析并读取文档文件中的信息，转化文本之后再内容传给 LLM 以实现对于文件内容的处理。文档提取器节点可以理解为一个信息处理中心，通过识别并读取输入变量中的文件，提取信息后并转化为 `string` 类型输出变量，供下游节点调用。

列表操作：列表操作节点可以对文件的格式类型、文件名、大小等属性进行过滤与提取，将不同格式的文件传递给对应的处理节点，以实现对不同文件处理流的精确控制。例如在一个应用中，允许用户同时上传文档文件和图片文件两种不同类型的文件。需要使用列表操作节点进行分拣，将不同的文件类型交由不同流程处理。

变量聚合：变量聚合节点（原变量赋值节点）是工作流程中的一个关键节点，它负责整合不同分支的输出结果，确保无论哪个分支被执行，其结果都能通过一个统一的变量来引用和访问。这在多分支的情况下非常有用，可将不同分支下相同作用的变量映射为一个输出变量，避免下游节点重复定义。通过变量聚合，可以将诸如问题分类或条件分支等多路输出聚合为单路，供流程下游的节点使用和操作，简化了数据流的管理。

变量赋值：变量赋值节点用于向可写入变量进行变量赋值，已支持会话变量和循环变量。通过变量赋值节点，你可以将会话过程中的上下文、上传至对话框的文件（即将上线）、用户

所输入的偏好信息等写入至会话变量，并在后续对话中引用已存储的信息导向不同的处理流程或者进行回复。

迭代：对数组中的元素依次执行相同的操作步骤，直至输出所有结果，可以理解为任务批处理器。迭代节点通常配合数组变量使用。使用迭代的条件是确保输入值已格式化为列表对象；迭代节点将依次处理迭代开始节点数组变量内的所有元素，每个元素遵循相同的处理步骤，每轮处理被称为一个迭代，最终输出处理结果。

参数提取：利用 LLM 从自然语言推理并提取结构化参数，用于后置的工具调用或 HTTP 请求。工作流内的部分节点有特定的数据格式传入要求，如迭代节点的输入要求为数组格式，参数提取器可以方便的实现结构化参数的转换。

http 请求：允许通过 HTTP 协议发送服务器请求，适用于获取外部数据、webhook、生成图片、下载文件等情景。它让你能够向指定的网络地址发送定制化的 HTTP 请求，实现与各种外部服务的互联互通。这个节点的一个实用特性是能够根据场景需要，在请求的不同部分动态插入变量。比如在处理客户评价请求时，你可以将用户名或客户 ID、评价内容等变量嵌入到请求中，以定制化自动回复信息或获取特定客户信息并发送相关资源至特定的服务器。

Agent：Agent 节点是平台工作流中用于实现自主工具调用的组件。它通过集成不同的 Agent 推理策略，使大语言模型

能够在运行时动态选择并执行工具，从而实现多步推理。

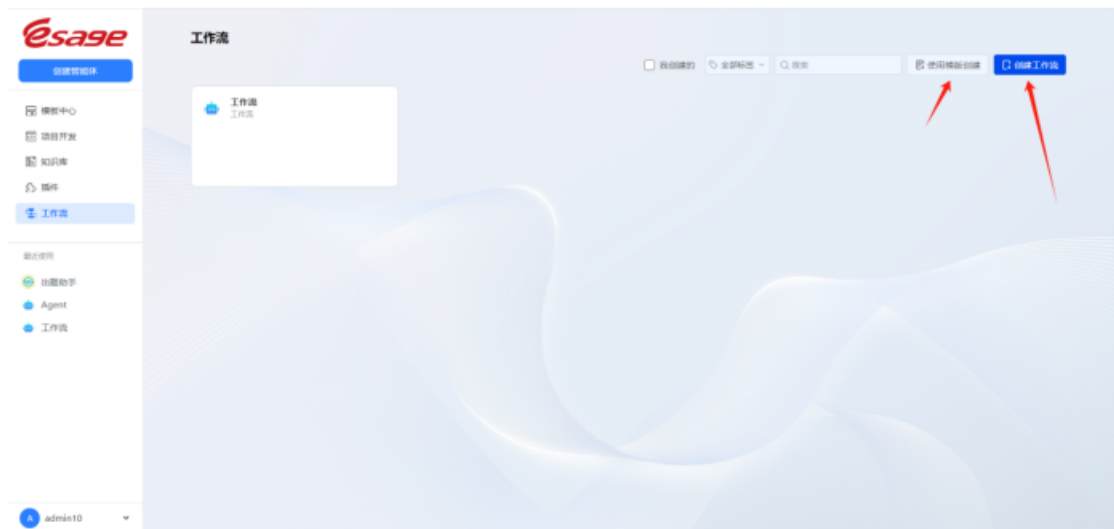
工具：“工具”节点可以为工作流提供强大的第三方能力支持，提供内置工具、自定义工具、工作流三种类型工具

结束：定义一个工作流程结束的最终输出内容。每一个工作流在完整执行后都需要至少一个结束节点，用于输出完整执行的最终结果。结束节点为流程终止节点，后面无法再添加其他节点，工作流应用中只有运行到结束节点才会输出执行结果。若流程中出现条件分叉，则需要定义多个结束节点。结束节点需要声明一个或多个输出变量，声明时可以引用任意上游节点的输出变量。

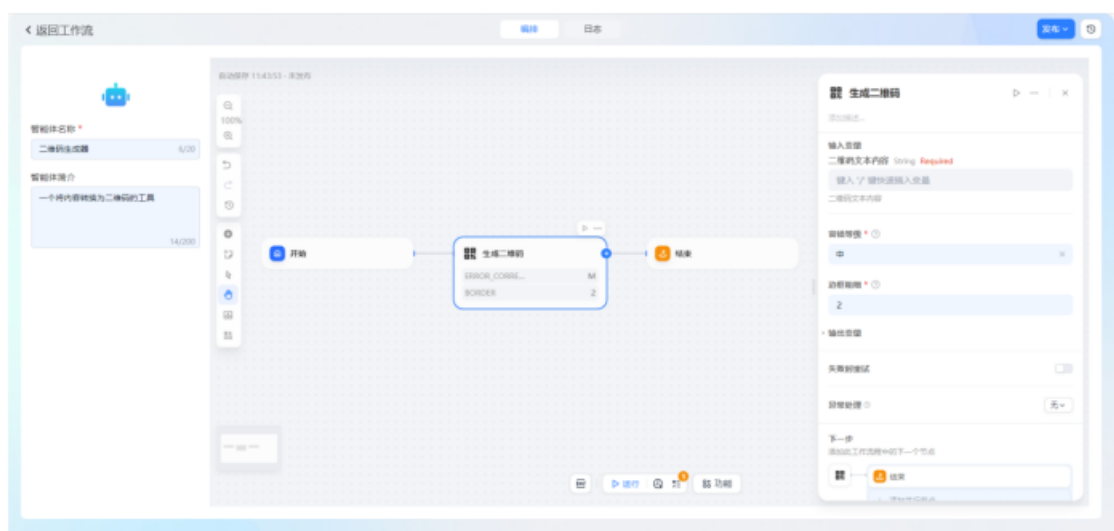
循环：循环（Loop）节点用于执行依赖前一轮结果的重复任务，直到满足退出条件或达到最大循环次数。注意：循环是需要前一轮的计算结果，而迭代每一轮都是独立执行。

（4）使用说明

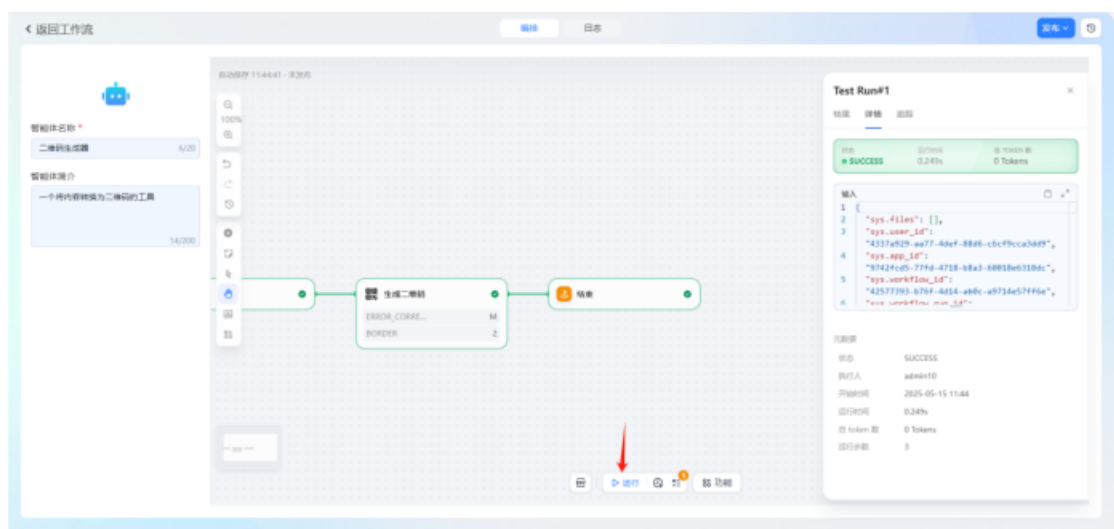
选择工作流菜单页，点击创建工作流/使用模板创建即可实现创建工作流。



可以根据自己的需求添加节点组合成工作流。



工作流通过点击运行，同样支持预览调试。



7. 插件（工具）

（1）定义

工具可以扩展 LLM 的能力，比如联网搜索、科学计算或绘制图片，赋予并增强了 LLM 连接外部世界的能力。平台提供了两种工具类型：内置和自定义。

你可以直接使用平台生态提供的内置工具，或者轻松导入自定义的 API 工具（目前支持 OpenAPI / Swagger 和 OpenAI Plugin 规范）。

（2）作用

工具使用户可以在平台上创建更强大的 AI 应用，如你可以为智能助理型应用（Agent）编排合适的工具，它可以通过任务推理、步骤拆解、调用工具完成复杂任务。

方便将你的应用与其他系统或服务连接，与外部环境交互，如代码执行、对专属信息源的访问等。

（3）如何创建自定义工具

你可以在“插件-自定义”内导入自定义的 API 工具，目前支持 OpenAPI/Swagger 和 ChatGPT Plugin 规范。你可以将 OpenAPI schema 内容直接粘贴或从 URL 内导入。



(4) 如何使用工具

在项目开发菜单的智能体或工作流的编排界面中，都可以使用工具。

