

DeepSeek + DeepResearch

让科研像聊天一样简单

北京航空航天大学 高研院 助理教授
清华大学新闻学院与人工智能学院双聘教授 沈阳团队博士后
何静

团队自研DeepResearch： 软件免费公测

通过夸克网盘分享的文件： AI学术工具公测版.exe
链接:<https://pan.quark.cn/s/e7df0f68ac97>



AI综述全版本
免费使用



DeepSeek资源合集

❤️ 优质资源推荐

👍 DeepSeek 资源

🌟 DeepSeek 微信群

💻 DeepSeek 工具 (23)

110+ 国内部署 DeepSeek 模型第...

DeepSeek 官方推荐: 70+ 集成...

DeepSeek 官方提示词库

DeepSeek 资料库飞书链接:

<https://acnawnoo7ng6.feishu.cn/wiki/FzozwGlZ2i6RTwkEJNZcYGMOnib>

☐ 文件名 ↕

大小



DeepSeek 15天指导手册——从入门到精通.pdf

1.3M



科学网—DeepSeek-R1的100问 - 王雄的博文.pdf

768.5KB



清华大学第一弹-DeepSeek 从入门到精通.pdf

4.8M



清华大学第二弹: DeepSeek赋能职场.pdf

9.8M



清华大学第三弹-普通人如何抓住DeepSeek红利.pdf

4.8M



清华大学第四弹-DeepSeek+DeepResearch: 让科研像聊天一样简单pdf_restore.pdf

6.8M



厦大团队: 大模型概念、技术与应用实践 (140页PPT读懂大模型) .pptx

60M



目录

一

能做什么？

二

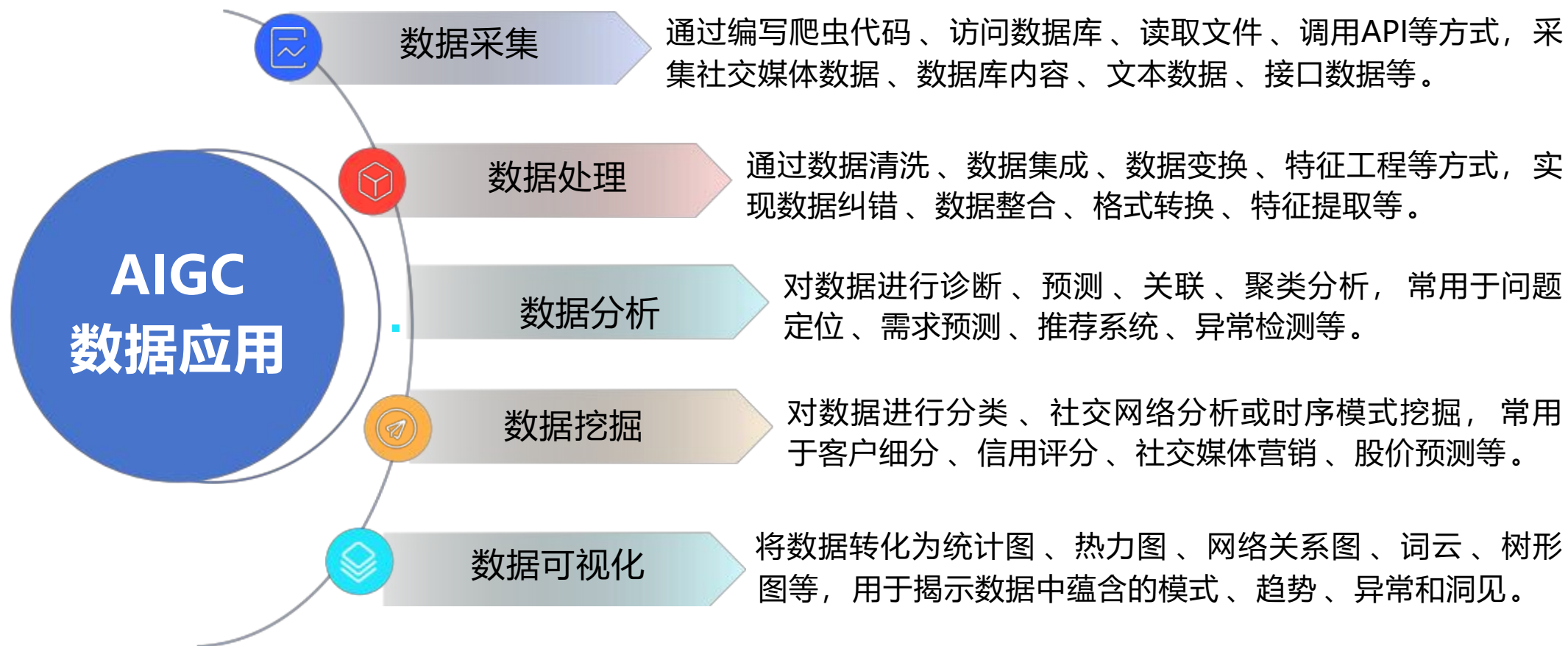
要怎么做？

三

效果如何？

能做什么？

本质：以多agent实现从数据采集到可视全流程



模型特点



DeepSeek
R1

- **高效推理**：专注于低延迟和高吞吐量，适合实时应用。
- **轻量化设计**：模型结构优化，资源占用少，适合边缘设备和移动端。
- **多任务支持**：支持多种任务，如文本生成、分类和问答。



Open AI
o3 mini

- **小型化设计**：轻量级模型，适合资源有限的环境。
- **快速响应**：优化推理速度，适合实时交互场景。
- **通用性强**：适用于多种自然语言处理任务，如对话生成和文本理解。



Claude
3.5 sonnet

- **平衡性能**：在模型大小和性能之间取得平衡，适合中等规模任务。
- **多模态支持**：支持文本和图像处理，扩展应用场景。
- **可解释性**：注重模型输出的可解释性和透明性。



Kimi
k1.5

- **垂直领域优化**：针对特定领域（如医疗、法律）进行优化，提供高精度结果。
- **长文本处理**：擅长处理长文本和复杂文档，适合专业场景。
- **定制化能力**：支持用户自定义训练和微调，适应特定需求。

爬虫数据采集

任务

- 1、阅读网页源代码，提取特定网页内容；
- 2、撰写python脚本；
- 3、提取并合并网址；
- 4、提取网址内容；
- 5、写入文件。

提示词

你需要完成以下两个任务：

1. 阅读网页【网址】源代码【对应网页源代码】。提取所有包含“春运2025 | X月X日，全社会跨区域人员流动量完成X万人次”的网址进行去重、筛选，合并成网址列表
2. 撰写python脚本，基于步骤1输出的网址列表提取所有网址内容“截至目前 2025 年春运（2025年1月14日到2月8日）相关数据（如日期、全社会跨区域人员流动量、铁路客运量、公路人员流动量、水路客运量、民航客运量等）”完成数据提取并写入文件“2025春运数据.txt”

帮我写python爬虫，爬取<https://news.buaa.edu.cn/zhxw.htm>网站中的标题和链接

帮我写python爬虫，爬取https://news.buaa.edu.cn/xxgg_new.htm网站中class="listlefttop auto"下面的标题和链接

帮我写python爬虫，爬取https://news.buaa.edu.cn/xxgg_new.htm网站中class="listlefttop auto"下面的所有标题和链接

(如果乱码？ UTF-8)

在桌面新建一个excel文档，并把结果放在这个文档中

对于每一个链接，请爬取每一个链接的内容，并保存到excel的新列“text”中

对于每一个链接，请爬取每一个链接中class="v_news_content"的内容，并保存到excel的新列“text”中

DeepSeek R1

能够提取所有网址并进行筛选、去重，所撰写代码运行后完成数据爬虫任务，所获取数据准确，少量数据有所遗漏。

2025 春运数据-r1.txt
日期,全社会跨区域人员流动量(万人次),铁路客运量(万人次),公路人员流动量(万人次),水路客运量(万人次),民航客运量(万人次)\n2025-02-08,280567.8,1318.4,78.235.35\n2025-02-07,21895.2,1425.1,78.3,237.82\n2025-02-06,21875.1,1514.,81.4,241.68\n2025-02-05,23153.8,1663.3,382.3,245.21\n2025-02-04,32084.6,3784.,122.244.57\n2025-02-03,33876.2,1645.4,122.8,248\n2025-02-02,32194.3,1469.6,141.9,242.83\n2025-02-01,31438.1,1238.,149.9,236.2\n2025-01-31,18484.1183.2,143.233.81\n2025-01-30,29435.8,967.6,131.6,236.55\n2025-01-29,28588.6,551.2,74.4,283.96\n2025-01-28,28438.9,776.,56.7,108.23\n2025-01-27,25888.2,1337.4,66.3,223.49\n2025-01-26,26581.5,1453.1,89.9,232.45\n2025-01-25,27318.8,1519.7,86.6,240.46\n2025-01-24,26084.5,1432.1,77.6,236.84\n2025-01-23,25651.9,1316.8,74.6,235.51\n2025-01-22,23385.5,1332.2,78.4,233.91\n2025-01-21,23854.1,1312.3,72.4,226.41\n2025-01-20,12226.8,1559.1,86.2,234.48\n2025-01-19,21411.8,1289.3,72.224.45\n2025-01-18,28873.4,1344.4,68.9,224.89\n2025-01-17,28183.1,1238.7,39.7,223.86\n2025-01-16,18785.1,1282.8,54.7,228.82\n2025-01-15,18464.9,1114.1,48.9,219.9\n2025-01-14,16626.4,1839.,53.2,288.18

```
return url_info\n\n# 抓取详情页数据\ndef fetch_detail_data(url_info):\n    for url, data in url_info.items():\n        try:\n            response = requests.get(url, headers=headers, timeout=10)\n            response.encoding = "utf-8"\n            content = response.text\n            soup = BeautifulSoup(content, "html.parser")\n            text = soup.get_text()\n\n            # 使用正则表达式提取数据\n            patterns = {\n                "railway": r"铁路客运量\\s+([\\d\\.]+)\\s+万人次",\n                "highway": r"公路人员流动量\\s+([\\d\\.]+)\\s+万人次",\n                "waterway": r"水路客运量\\s+([\\d\\.]+)\\s+万人次",\n                "aviation": r"民航客运量\\s+([\\d\\.]+)\\s+万人次"\n            }\n\n            for key, pattern in patterns.items():\n                match = re.search(pattern, text)\n                if match:\n                    data[key] = match.group(1).replace(".", "")\n            except Exception as e:\n                print(f"抓取 {url} 失败: {str(e)}")\n\n    return url_info\n\n# 保存数据到文件\ndef save_to_file(data, filename):\n    with open(filename, "w", encoding="utf-8") as f:\n        f.write(f"日期,全社会跨区域人员流动量(万人次),铁路客运量(万人次),公路人员流动量(万人次),水路客运量(万人次),民航客运量(万人次)\\n")\n        for item in data.values():\n            line = f"{item['date']},{item['total']},{item['railway']},{item['highway']},{item['waterway']},{item['aviation']}\\n"\n            f.write(line)
```

Open AI o3mini

响应速度快，能够高效提取所有需求链接，输出完整可运行python脚本，代码运行后生成文件，但数据采集结果为空。

2025 春运数据-o3-mini.txt
日期,全社会跨区域人员流动量,铁路客运量,公路人员流动量,水路客运量\n春运2025 2月8日,全社会跨区域人员流动量,无,无,无,无\n春运2025 2月7日,全社会跨区域人员流动量,无,无,无,无\n春运2025 2月6日,全社会跨区域人员流动量,无,无,无,无\n春运2025 2月5日,全社会跨区域人员流动量,无,无,无,无\n春运2025 2月4日,全社会跨区域人员流动量,无,无,无,无\n春运2025 2月3日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月31日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月30日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月29日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月27日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月26日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月24日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月22日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月21日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月20日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月18日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月17日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月16日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月15日,全社会跨区域人员流动量,无,无,无,无\n春运2025 1月14日,全社会跨区域人员流动量,无,无,无,无\n春运2025 春运数据,预计全社会跨区域人员流动量达17.26亿人次-春运数据-中华人民共和

```
"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202502/120250" 1 复制\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202502/120250266_4163764"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202502/120250265_4163716"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202502/120250264_4163641"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202502/120250263_4163612"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202502/120250262_4163571"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202502/120250261_4163533"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202501/120250131_4163492"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202501/120250130_4163456"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202501/120250129_4163441"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202501/120250128_4163398"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202501/120250127_4163352"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202501/120250126_4163284"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202501/120250125_4163210"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202501/120250124_4163156"\n"https://www.mot.gov.cn/zhuanli/2025chunyun/chunyunshuju/202501/120250123_4163189"\n]
```

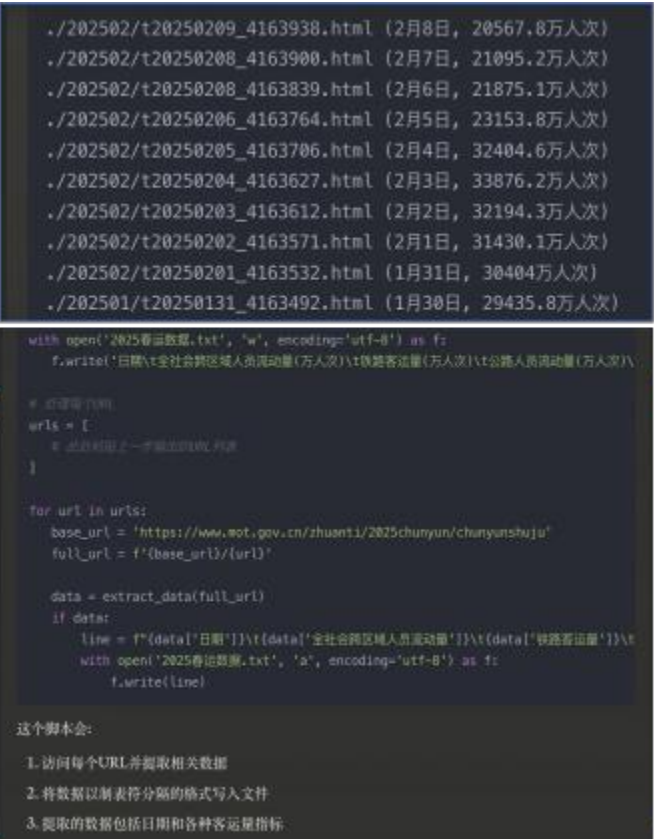
```
# 定义一个函数来提取网页中的数据\ndef extract_data(url):\n    response = requests.get(url)\n    response.encoding = "utf-8" # 设置编码为utf-8\n    soup = BeautifulSoup(response.text, "html.parser")\n    content = soup.find('div', class_='TRS_editor')\n    if content:\n        return content.get_text(strip=True)\n    else:\n        return None\n\n# 打开文件以写入数据\nwith open('2025春运数据.txt', 'w', encoding='utf-8') as file:\n    for url in urls:\n        data = extract_data(url)\n        if data:\n            file.write(data + "\\n")\n        else:\n            file.write(f"未能从 {url} 提取数据。\\n")\n\nprint("数据提取完成，已保存至 '2025春运数据.txt'。")
```

测试结果受到数据样本、测试环境、AI抽卡、提示词模板等因素影响，仅供参考，无法作为决策制定、质量评估或产品验证的最终依据。

爬虫数据采集

Claude 3.5 sonnet

可以提取所有网址，调整后可输出正确代码，运行代码能生成本地文件，但提取数据结果为空。



Kimi k1.5

能够提取所有网址，代码运行后生成本地文件，但提取数据结果为空。



结论

- 目前DeepSeek R1、Open AI o3mini、Kimi k1.5支持联网查询网址，Claude 3.5 sonnet暂不支持；
- 四个模型均能根据上传的网页代码，对多个网址链接进行筛选、去重，完全提取出符合指令要求的所有网址链接并形成列表；
- 在复杂爬虫任务上，DeepSeek R1与Open AI o3min生成的代码均能正常执行数据采集任务，o3响应速度更快，R1数据采集结果更加完整准确；其他2个模型都存在多次调试但代码仍然运行不成功的问题，如代码中罗列URL不全、输出文本中提取数据为空等。

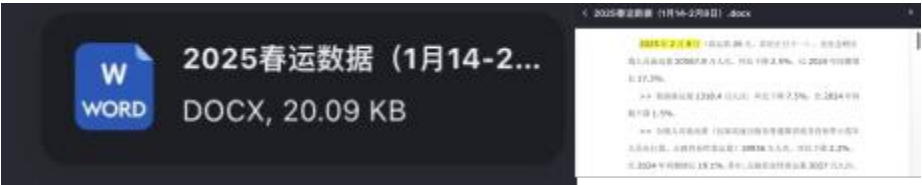
文件数据读取

任务

- 1、读取文件；
- 2、根据指定内容整理成表格。

提示词

所上传的“2025春运数据（1月14-2月8日）.txt”包含了从2025年1月14日至2025年2月8日每天各种交通方式的春运客运量信息，请从中读取每一天的信息，并整理成一张表格，要求包括以下几项信息：1.当天日期；2.当天的铁路客运量、比2024年同期多或者少的百分比、环比的百分比。3.当天的公路客运量、比2024年同期多或者少的百分比、环比的百分比。4.当天的民航客运量、比2024年同期多或者少的百分比、环比的百分比。



DeepSeek R1

能够详细全面地提取文件中的数据，并整理成可视化数据表格，逻辑性强、指标清晰。

2025-01-29	551.2	-0.9%	-29.0%	19671.0	+5.5%	+1.3%	203.96	+5.5%	+9.5%
2025-01-30	967.6	+7.9%	+75.5%	28110.0	+7.4%	+42.8%	226.55	+4.2%	+11.08%
2025-01-31	1183.2	+5.3%	+22.3%	28844.0	+6.9%	+2.6%	233.81	+3.6%	+3.2%
2025-02-01	1338.0	+3.5%	+13.1%	29707.0	+4.3%	+3.0%	238.20	+2.8%	+1.8%
2025-02-02	1460.6	+2.5%	+9.2%	30349.0	+5.0%	+2.2%	242.83	+1.3%	+1.9%
2025-02-03	1645.4	+8.5%	+12.7%	31860.0	+7.4%	+5.0%	248.00	+2.0%	+2.1%
2025-02-04	1704.0	+11.5%	+3.6%	30334.0	+4.7%	-4.8%	244.57	+0.5%	-1.3%
2025-02-05	1663.3	+3.5%	-2.4%	21143.0	-26.8%	-30.3%	245.21	+2.2%	+0.2%
2025-02-06	1514.0	-3.0%	-9.0%	20038.0	+1.9%	-5.2%	241.68	+2.3%	-1.4%
2025-02-07	1425.1	+1.5%	-5.9%	19362.0	+10.0%	-3.4%	237.82	+3.2%	-1.6%
2025-02-08	1318.4	-1.5%	-7.5%	18936.0	+19.1%	-2.2%	235.35	+8.9%	-1.0%
说明：									
1. 公路客运量为文件中“公路人员流动量”的总数据（含营业性与非营业性）。									
2. 同比为“比2024年同期”的百分比，环比为与前一日期相比的百分比。									
3. 原始数据中2025年1月22日重复，已合并处理。									

Open AI o3mini

暂不支持附件上传，响应速度快，能够快速读取粘贴数据，输出结果格式工整、简洁。

日期	铁路客 运量 (万人 次)	铁路客 运量与 2024 年同比 (%)	铁路客 运量与 2019年 同比 (%)	公路客 运量 (万人 次)	公路客 运量与 2024年 同比 (%)	公路客 运量与 2019年 同比 (%)	民航客 运量 (万人 次)	民航客 运量与 2024年 同比 (%)	民航客 运量与 2019年 同比 (%)
2025 年2月 8日	1318.4	-1.5%	13.7%	18936	19.1%	23.5%	235.35	8.95%	28.4%
2025 年2月 7日	1425.1	1.5%	6.4%	19362	10%	13.2%	237.82	3.21%	15.6%
2025 年2月 6日	1514	-3%	5.1%	20038	1.9%	10.6%	241.68	2.35%	7.5%
2025 年2月 5日	1663.3	3.5%	17.3%	21143	-26.8%	-18.4%	245.21	2.27%	16.2%
2025 年2月 4日	1704	11.5%	26.5%	30334	4.7%	13.5%	244.57	0.55%	11.4%
2025 年2月 3日	1645.4	8.5%	15.4%	31860	7.4%	18.3%	248	2.03%	18.6%
2025 年2月 2日	1460.6	2.5%	8.5%	30349	5%	14.3%	242.83	1.38%	14.1%
2025 年2月 1日	1338	3.5%	10.2%	29707	4.3%	11.1%	238.2	2.84%	12.9%
2025 年1月 31日	1183.2	5.3%	9.4%	28844	↓ 7%	10.4%	233.81	3.6%	17.8%

文件数据读取

Claude 3.5 sonnet

很好地完成了数据读取及提取任务，没有漏数据指标，数据逻辑性很好

Spring Festival Transport Data Summary					
Date	Railway Transport (10k)	YoY Change	Road Transport (10k)	YoY Change	Air Transport (10k)
2025-02-08	1318.4	-1.5%	18936	+19.1%	235.1
2025-02-07	1425.1	+1.5%	19362	+10.0%	235.1
2025-02-06	1514.0	-3.0%	20038	+1.9%	241.2
2025-02-05	1663.3	+3.5%	21143	-26.6%	245.1
2025-02-04	1704.0	+11.5%	30334	+4.7%	245.1
2025-02-03	1645.4	+8.5%	31860	+7.4%	241.2
2025-02-02	1460.6	+2.5%	30349	+5.0%	241.2
2025-02-01	1338.0	+3.5%	29707	+4.3%	231.1
2025-01-31	1183.2	+5.3%	28844	+6.9%	231.1
2025-01-30	967.6	+7.9%	28110	+7.4%	221.1
2025-01-29	551.2	-0.9%	19671	+5.5%	201.1
2025-01-28	776.0	-6.0%	19418	+5.2%	186.1
2025-01-27	1337.4	+7.5%	24253	-5.1%	221.1
2025-01-26	1453.1	+12.1%	24826	+8.0%	231.1
2025-01-25	1519.7	+17.9%	25472	+17.3%	241.2
2025-01-24	1432.1	+17.4%	24858	+23.7%	231.1
2025-01-23	1316.8	+7.4%	24025	+34.3%	231.1
2025-01-22	1332.2	+3.9%	21749	+15.7%	231.1
2025-01-21	1312.3	+8.1%	21443	+15.8%	221.1
2025-01-20	1359.1	+12.5%	20605	+13.1%	221.1
2025-01-19	1289.3	+12.2%	19826	+10.8%	221.1
2025-01-18	1344.4	+22.4%	19234	+4.9%	221.1
2025-01-17	1238.7	+2.7%	18581	+5.1%	221.1
2025-01-16	1202.6	+2.6%	17307	-3.1%	221.1
2025-01-15	1114.1	-8.2%	17061	-2.4%	211.1
2025-01-14	1039.0	-6.1%	15326	-12.7%	201.1

Kimi k1.5

能够快速读取文件数据，并整理成可视化数据表格，但填入数据有所缺失。

日期	铁路客流量 (万人次)	铁路客流量同比	铁路客流量环比	公路客流量 (万人次)
2025/1/14	-	-	-	-
2025/1/15	1114.1	+8.2%	+7.2%	17061
2025/1/16	1202.6	+2.6%	+7.9%	17307
2025/1/17	1238.7	+2.7%	+3.0%	18581
2025/1/18	1344.4	+22.4%	+8.5%	19234
2025/1/19	1289.3	+12.2%	+4.1%	19826
2025/1/20	1359.1	+12.5%	+5.4%	20605
2025/1/21	1312.3	+8.1%	+3.4%	21443
2025/1/22	1332.2	+3.9%	+1.5%	21749
2025/1/23	1316.8	+7.4%	+1.2%	24025
2025/1/24	1432.1	+17.4%	+8.8%	24858
2025/1/25	1519.7	+17.9%	+6.1%	25472
2025/1/26	1453.1	+12.1%	+4.4%	24826
2025/1/27	1337.4	+7.5%	+8%	24263
2025/1/28	776	+6%	+42%	19418
2025/1/29	551.2	+0.9%	+29%	19671

结论

- DeepSeek R1与Claude 3.5 sonnet均能很好的完成文件数据读取任务，生成的表格逻辑性强、数据指标清晰，Claude 3.5 sonnet一次性完成表格生成后支持直接复制和表格文件下载。
- Open AI o3mini能够更加高效地完成数据提取任务，输出表格准确，但暂不支持附件上传，文件读取依靠粘贴稍显不便；Kimi k1.5文件数据提取有明显空缺。

文本数据集成

任务

- 1、分别阅读约7000token和15000token的文本内容，测试模型对中、长文本处理效果
- 2、整理集成可视化的数据表格
- 3、按照日期规范排序

提示词

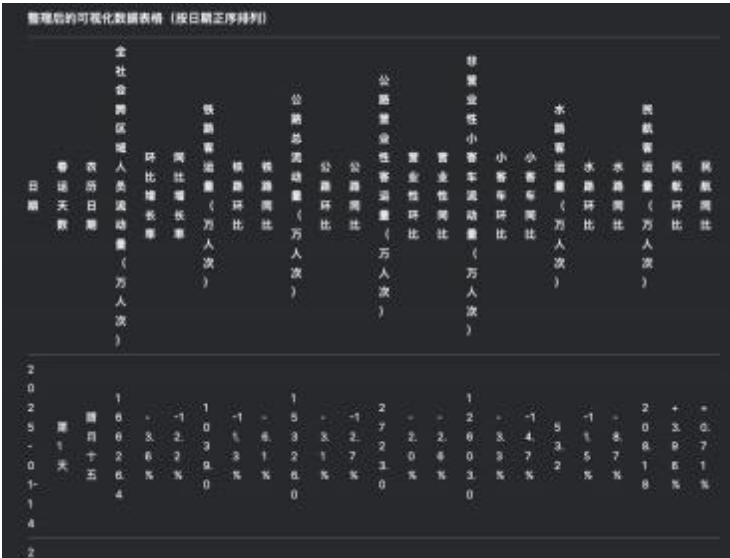
【所需阅读文本】请根据以上文本完成以下三个任务：1、阅读文本内容；2、整理集成可视化的数据表格；3、按照日期规范排序（直接分析并输出结果，不使用python）



DeepSeek R1

一般文本（7000token）：能详细全面地提取文本数据，并集成可视化表格，但受大样本或模型稳定性影响，输出表格末尾缺失，需要重复尝试生成。

长文本（15000token）：暂时无法给出答复。



体上应对有序有效，总体运行平稳，经受住了复杂恶劣天气的严峻考验。个任务：1、阅读文本内容；2、整理集成可视化的数据表格；3、按照日期规范排序（直接分析并输出结果，不使用python）

你好，这个问题我暂时无法回答，让我们换个话题再聊聊吧。

Open AI o3mini

一般文本（7000token）：能够高效提取文本中的数据，并整理成可视化数据表格，格式工整、简洁，数据准确但数据维度有所缺失。

长文本（15000token）：反应迅速，高效提取文本中数据，输出格式规范的数据表格，但集成数据维度仍然不够全面。

日期	跨区域人员流动量 (万人次)	铁路客运量 (万人次)	公路人员流动量 (万人次)	水路客运量 (万人次)	民航客运量 (万人次)
2025年2月8日	20567.8	1318.4	18936	78	235.35
2025年2月7日	21095.2	1425.1	19362	70.3	237.82
2025年2月6日	21875.1	1514	20038	81.4	241.68
2025年2月5日	23153.8	1663.3	21143	102.3	245.21
2025年2月4日	32404.6	1704	30334	122	244.57
2025年2月3日	33876.2	1645.4	31860	122.8	248
2025年2月2日	32194.3	1460.6	30349	141.9	242.83

根据您的要求，下面是整理后的数据：

日期	跨区域人员流动量 (万人次)	铁路客运量 (万人次)	公路人员流动量 (万人次)	水路客运量 (万人次)	民航客运量 (万人次)
2025年2月8日	20567.8	1318.4	18936	78	235.35
2025年2月7日	21095.2	1425.1	19362	70.3	237.82
2025年2月6日	21875.1	1514	20038	81.4	241.68
2025年2月5日	23153.8	1663.3	21143	102.3	245.21
2025年2月4日	32404.6	1704	30334	122	244.57
2025年2月3日	33876.2	1645.4	31860	122.8	248

文本数据集成

Claude 3.5 sonnet

一般文本（7000token）：

能够准确提取文本数据，并整理成可视化图表，表格文字简洁，没有提及文本中的环比和同比数据。

长文本（15000token）：

长文本粘贴后会自动形成文件，能够准确集成文本数据表格，但数据维度有限。

已完成文本阅读和数据整理，按日期排序的2025年春运数据如下：

日期	春运天数	总客流量(万人次)	铁路(万人次)	公路(万人次)	水路(万人次)	民航(万人次)
1.14	1	16626.4	1039.0	15326.0	53.2	206.18
1.15	2	18444.9	1114.1	17061.0	49.9	219.90
1.16	3	18785.1	1202.6	17307.0	54.7	220.82
1.17	4	20103.1	1238.7	18581.0	59.7	223.66
1.18	5	20871.4	1344.4	19234.0	68.9	224.09
1.19	6	21411.8	1289.3	19826.0	72.0	224.45
1.20	7	22256.8	1359.1	20605.0	68.2	224.48
1.21	8	23054.1	1312.3	21443.0	72.4	226.41
1.22	9	23385.5	1332.2	21749.0	70.4	233.91
1.23	10	25651.9	1316.8	24025.0	74.6	235.51
1.24	11	26604.5	1432.1	24858.0	77.6	236.84
1.25	12	27318.8	1519.7	25472.0	86.6	240.46
1.26	13	26581.5	1453.1	24826.0	69.9	232.45

2025 Spring Festival Transportation Data List (1-14th)

Date	Line Name	Total Flow	Bus	Road	Waterway	Aviation
2025-02-08	京沪线	25,087.8	1,278.9	23,808.9	7.9	330.26
2025-02-07	京沪线	21,066.2	1,055.1	19,982.1	15.4	213.62
2025-02-06	京沪线	19,751.1	1,016	18,735.1	81.4	188.66
2025-02-05	京沪线	18,764.4	988.6	17,775.8	66.0	169.91
2025-02-04	京沪线	18,404.8	1,194	17,210.8	121	144.91
2025-02-03	京沪线	18,876.2	1,093.4	17,782.8	122.8	159.6
2025-02-02	京沪线	18,164.2	1,060.8	17,103.4	147.8	243.89
2025-02-01	京沪线	18,434.1	1,068	17,366.1	188.9	218.2
2025-01-31	京沪线	18,484	1,043.0	17,441	141	199.81
2025-01-30	京沪线	18,498.8	987.9	17,510.9	111.6	128.95
2025-01-29	京沪线	18,800.8	1,017.2	17,783.6	115.4	129.88
2025-01-28	京沪线	18,138.8	1,176	16,962.8	185.7	188.23
2025-01-27	京沪线	18,486.2	1,017.4	17,468.8	186.0	225.88
2025-01-26	京沪线	18,588.8	1,031.1	17,557.7	189.9	233.84
2025-01-25	京沪线	18,768.8	1,076.7	17,692.1	188.0	288.88

Kimi k1.5

一般文本（7000token）：

能够快速提取文本数据，并整理成可视化数据表格，数据准确，所提取数据维度不够全面。

长文本（15000token）：

同样能够高效准确地提取文本中数据，较一般长度文本所集成数据维度反而更加全面。

日期	全社会跨区域人员流动量 (万人次)	铁路客运量 (万人次)	公路人员流动量 (万人次)
2025年1月14日	16626.4	1039	15326
2025年1月15日	18444.9	1114.1	17061
2025年1月16日	18785.1	1202.6	17307
2025年1月17日	20103.1	1238.7	18581
2025年1月18日	20871.4	1344.4	19234
2025年1月19日	21411.8	1289.3	19826
2025年1月20日	22256.8	1359.1	20605

环比	同比2024年	铁路客运量 (万人次)	环比	同比2024年	公路人员流动量 (万人次)
-3.6%	-12.2%	1039	-11.3%	-6.1%	15326
10.9%	-2.7%	1114.1	7.2%	-8.2%	17061
1.8%	-2.7%	1202.6	7.9%	2.6%	17307
7.0%	5.0%	1238.7	3.0%	2.7%	18581
3.8%	6.0%	1344.4	8.5%	22.4%	19234
2.6%	10.9%	1289.3	-4.1%	12.2%	19826
3.9%	13.0%	1359.1	5.4%	12.5%	20605
3.6%	15.2%	1312.3	-3.4%	8.1%	21443
1.4%	14.9%	1332.2	1.5%	-3.9%	21749

结论

- 一般文本处理任务中，DeepSeek R1所提取的文本数据维度最为全面，但容易受文本长度或模型稳定性影响出现失误；其他三个模型在文本数据提取过程中，都存在对部分数据的忽略问题，没有完整集成到可视化表格中；
- 长文本处理任务中，Kimi k1.5相较短文本处理表现更加突出，提取准确的同时数据维度更加全面；由于文本过长DeepSeek R1无法完成任务；
- 综合来看，Open AI o3mini响应最快效率最高，但在数据集成维度上稍显不足，同时与Claude 3.5 sonnet所输出的表格更为工整、简洁。

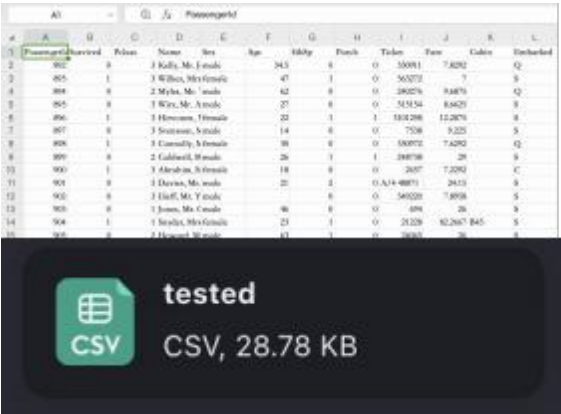
数据分析

任务

- 1、读取titanic遇难者名单excel
- 2、找出其中规律

提示词

请读取所上传的titanic遇难者名单数据文件，并分析找出其中规律



DeepSeek R1

详细展示长思维链，精准提取关键指标“生存率”，分析多个因素特征对生存率的影响，结合历史背景对数据及规律进行验证，并敏锐察觉数据异常，提出了异常处理建议。

2. 关键规律分析

(1) 舱位等级 (Pclass) 与生存率

- 头等舱 (Pclass=1): 生存率最高 (约62%)，因其优先使用救生艇且靠近甲板。
- 二等舱 (Pclass=2): 生存率次之 (约47%)。
- 三等舱 (Pclass=3): 生存率最低 (约24%)，乘客资源分配受限且位置偏远。

(2) 性别 (Sex) 与生存率

- 女性: 生存率显著高于男性 (约74% vs 19%)，符合“妇女儿童优先”原则。
- 男性: 尤其在低舱位中存活率极低 (如三等舱男性生存率不足10%)。

(3) 年龄 (Age) 与生存率

- 儿童 (Age < 18): 生存率较高 (约50%以上)，部分因优先救援政策。
- 成年人: 年龄越大生存率可能下降 (如60岁以上乘客多未存活)。
- 缺失年龄数据: 需进一步处理 (如用中位数填充)。

3. 其他观察

- 客舱号 (Cabin): 缺失值较多，但头等舱客舱记录较完整 (如B、C、D区)，可能与生存率相关。
- 异常值: 部分高票价乘客未幸存 (如ID 916, 票价262.375但未幸存)，可能因特殊情况 (如放弃逃生)。

4. 总结

- 生存关键因素: 女性、儿童、高舱位、小型家庭。
- 数据验证: 符合历史背景 (如阶级差异和救援优先级)。
- 改进方向: 处理缺失年龄数据，深入分析客舱位置影响。

Open AI o3mini

响应速度快，高效输出数据分析结果，分析各因素对关键指标生存率的影响，语言表达自然，重点突出结合历史背景对数据规律进行验证，但没有察觉数据异常。

已推理，持续 9 秒 >

下面给出一种对该 Titanic 数据 (包含乘客ID、是否生还、舱位等级、姓名、性别、年龄、亲友数量、票价、船舱、登船港口等信息) 的探索性分析和归纳规律的思路，供参考:

2. 探索性分析与规律

(1) 性别与生还率

- 观察:
 - 从数据中可以直观看出，女性乘客的生还率明显高于男性。
 - 例如，在许多记录中，若 Sex 为 female，对应的 Survived 多为 1；而男性 (male) 中很多记录显示 Survived 为 0。
- 可能原因:
 - 这与当时的“妇女和儿童优先”救生政策一致。

3. 总结

综合上述分析，可以归纳出 Titanic 遇难者名单数据中的主要规律:

- 性别: 女性乘客生还率显著高于男性，符合“妇女和儿童优先”的救生策略。
- 舱位等级: 头等舱乘客的生还率远高于三等舱，反映了船上社会经济地位的差异对逃生机会的影响。
- 年龄: 儿童和年轻乘客 (尤其是女性) 生还较多，虽然数据中存在缺失，但整体趋势与历史记录一致。
- 经济因素: 票价较高 (对应头等舱) 的乘客更容易生还，间接体现了经济条件和社会地位的重要性。
- 家庭/同伴因素: 与家人或同伴同行可能在一定程度上影响生还概率，但需结合更多统计量来深入探讨。
- 登船港口: 不同港口的乘客背景可能有所不同，进而影响生还率。

这种探索性的数据分析为后续构建更精确的统计模型 (如逻辑回归、决策树、随机森林等) 提供了直观的启发，也印证了历史记载中“贵客优先、妇女儿童优先”的实际情况。

数据分析

Claude 3.5 sonnet

提供数据分析程序代码，能够提取大部分特征并对其进行与生存率的关联进行分析，但最终没有形成明确的结论。

基于分析结果，我发现了以下主要规律：

1. 性别差异：

- 女性的生存率显著高于男性
- 这可能是因为“妇女和儿童优先”的撤离政策

2. 舱位等级影响：

- 一等舱乘客的生存率最高
- 三等舱乘客的生存率最低
- 这反映了社会阶层对生存机会的影响

3. 年龄特征：

- 儿童的生存率较高
- 中年乘客的生存率相对较低
- 老年乘客的生存率最低

4. 登船港口影响：

- 从不同港口登船的乘客生存率有所不同
- 这可能与乘客的社会经济地位有关

5. 家庭因素：

- 有家庭成员同行的乘客生存率略高
- 这可能是因为家人之间相互照应

这些发现揭示了社会地位（舱位等级）、性别、年龄等因素对生存机会的重要影响。救生设施的分配和使用可能存在偏向性，反映了当时的社会价值观和规范。

Kimi k1.5

能够精准分析关键指标生存率，但对特征提取不完整，仅能识别较为浅层的数据关联，分析能力相对较弱。

模式识别

1. 女性和儿童优先：

- 数据显示，女性和儿童的生存率显著高于男性和成年人。这符合泰坦尼克号沉船事件中“女性和儿童优先”的救援原则。

2. 舱位等级对生存率的影响：

- 头等舱和二等舱乘客的生存率显著高于三等舱乘客，这可能与舱位位置和救援资源的分配有关。

3. 家庭关系对生存率的影响：

- 有亲属同行的乘客生存率较高，尤其是女性和儿童；这可能与家庭成员之间的互助有关。

4. 票价与生存率的关系：

- 票价较高的乘客（头等舱和二等舱）生存率较高，这与舱位等级的影响一致。

可视化建议

为了更直观地展示分析结果，可以使用以下可视化图表：

- 生存率柱状图：按性别、舱票等级和年龄分组展示生存率。
- 年龄分布直方图：展示乘客年龄的分布情况。
- 票价分布箱线图：展示不同舱票等级的票价分布。
- 家庭关系与生存率散点图：展示家庭关系（SibSp + Parch）与生存率的关系。

结论

通过对泰坦尼克号乘客数据的分析，可以得出以下结论：

- 女性和儿童的生存率显著高于男性和成年人。
- 头等舱和二等舱乘客的生存率显著高于三等舱乘客。
- 有亲属同行的乘客生存率较高。
- 票价较高的乘客生存率较高。

这些结论与历史记录和常识相符，进一步验证了泰坦尼克号沉船事件中的救援原则和乘客特征对生存率的影响。

结论

- DeepSeek R1与Open AI o3mini的数据分析能力相当，且领先其他两个模型，均能够精准抓取数据核心指标并做统计，找到各特征与核心指标的关联，其中R1分析逻辑更加清晰严谨，而o3推理更加高效；
- Kimi k1.5推理逻辑清晰但分析能力相对较弱， Claude 3.5 sonnet能够提供分析思路但没有明确结论。

数据挖掘

任务

- 1、读取即将上映的2025年电影数据集
- 2、对数据集进行深入分析和数据挖掘

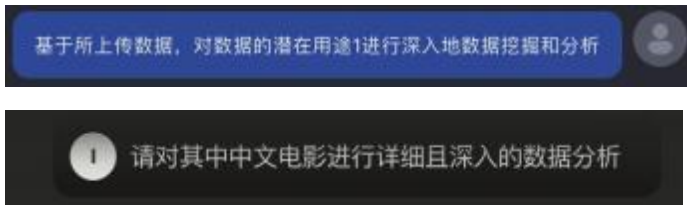
提示词

第一轮对话:



第二轮对话:

(基于初步分析结果, 选择其中一部分或某个方向进行深入的数据挖掘)



DeepSeek R1

能够准确对数据进行分类, 从多个维度进行梳理和分析, 借助可视化图表进行数据挖掘, 基于分析结果提供可行建议, 但整体数据挖掘深度较浅, 缺少对不同类型数据直接关联性的探究。

2. 关键发现

语言分布

- 英语电影最多: 占比约 60% (如《The Wild Robot》《Captain America: Brave New World》)。
- 其他语言:
 - 日语 (如《劇場版「オーバーロード」 聖王國編》)、西班牙语 (如《La habitación de al lado》)、法语、韩语等。
 - 部分小众语言电影 (如拉脱维亚语《Straume/Flow》)。

高分电影 (vote_average ≥ 8)

- 《The Wild Robot》: 评分 8.365, 投票数 4058 (英语, 2024-09-12)。
- 《Flow》: 评分 8.401 (拉脱维亚语, 2024-08-29)。
- 《Attack on Titan: THE LAST ATTACK》: 评分 9.8 (日语, 2024-11-08, 但投票数仅 7)。

低评分电影 (vote_average < 5)

- 《Peter Pan's Neverland Nightmare》: 评分 3.4 (英语, 2025-01-13)。
- 《The Mouse Trap》: 评分 4.325 (英语, 2024-08-23)。

未评分电影 (vote_count = 0)

- 多部 2025 年上映的电影尚未开放评分, 如《Captain America: Brave New World》《Bridget Jones: Mad About the Boy》。

3. 时间分布

- 2024 年上映: 多数电影集中在 2024 年 8-12 月。
- 2025 年上映:
 - 1 月: 多部非英语电影 (如韩语《검은 수녀들/Dark Nuns》)。
 - 2 月: 英语电影密集 (如《The Witcher: Sirens of the Deep》)。

4. 类型与题材

通过标题和译名推测主要类型:

1. 动作/冒险:
 - 《蛟龙行动》(Operation Leviathan)
 - 《唐探1900》(Detective Chinatown 1900)
2. 动画/家庭:
 - 《熊出没·重返未来》(Boonie Bears: Future Reborn)
 - 《哪吒之魔童闹海》(Ne Zha 2)
3. 爱情/剧情:
 - 《真爱营业》(Liar, Liar, Love is on Fire)
 - 《多幸运遇见你》(How Lucky to Meet You)
4. 惊悚/悬疑:
 - 《破·地獄》(The Last Dance)

5. 异常数据与备注

1. 历史电影混入:
 - 《花样年华》(2000 年上映) 可能为数据误录, 需核实是否为重映或误作。
2. 重复语言代码:
 - 语言代码重复出现 (如 "ja") 以标准化分类。
3. 未评分电影潜力:
 - 《唐探1900》《哪吒之魔童闹海》等 IP 续作可能成为票房热点。

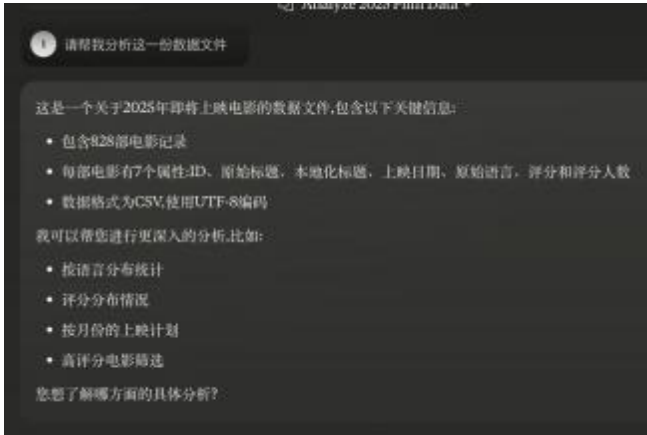
6. 市场建议

1. 重点宣发:
 - 关注 2025 年 1-2 月密集上映期, 提前布局春节档营销。
 - 动画电影 (如《熊出没》《哪吒》) 适合家庭观众, 可联合亲子平台推广。
2. 口碑监控:
 - 跟踪《破·地獄》等高分电影的长尾效应。
 - 对未评分 IP 电影 (如《唐探1900》) 进行预售热度分析。
3. 类型差异化:
 - 加强动作片与悬疑片的国际市场发行 (如东南亚、北美华语圈)。

数据挖掘

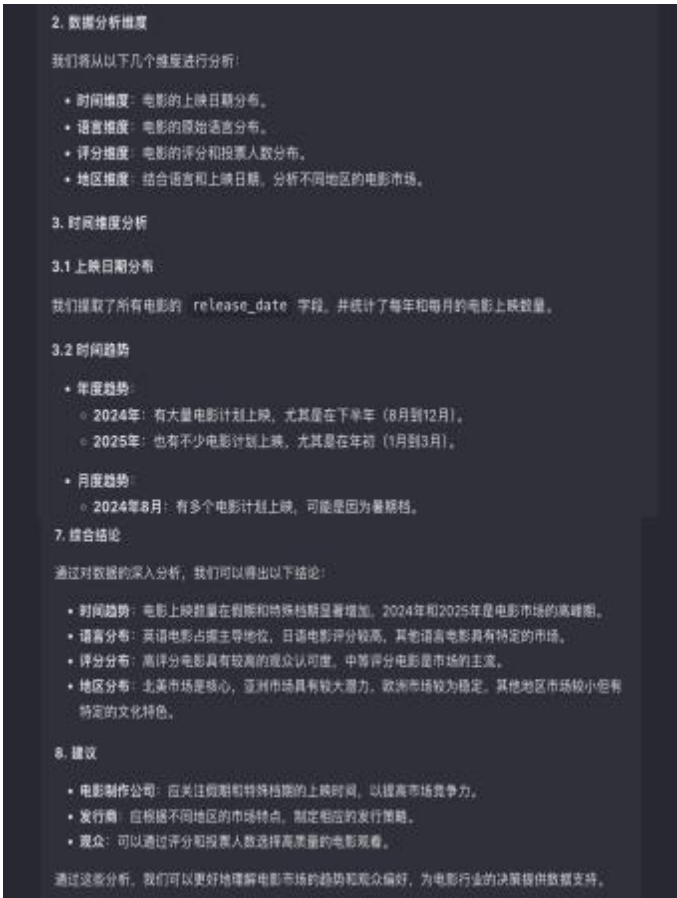
Claude 3.5 sonnet

基于数据集，在整体数据概括后提供多个深入数据挖掘方向，根据需求输入研究倾向，高效生成多个维度的数据分析，语言简洁，挖掘深度较浅。



Kimi k1.5

提供数据的潜在用途方向，深入分析过程中，从多个维度（如时间、语言、地区）深入挖掘数据意义和关联性，进一步总结趋势结论并提出相关建议。



结论

- ❑ Kimi k1.5该任务中表现最为出色，对特征进行精准分类，从多维度深入挖掘指定数据的深层内涵和关联性；
- ❑ DeepSeek R1与Claude 3.5 sonnet同样能准确完成数据的分类任务，但数据之间的关联挖掘程度相对较浅；
- ❑ Open AI o3mini受附件上传限制影响，由于数据集较大，暂不能完成该任务。

数据可视化

任务

基于titanic遇难者数据分析结果绘制可视化图表

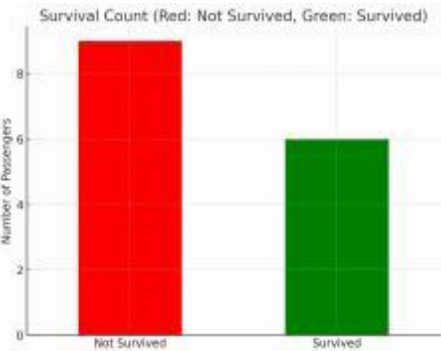
结论

- Open AI o3mini的数据可视化能力突出，能够直接高效地生成多种类型可视化图表，准确度高；
- DeepSeek R1、Kimi k1.5均能基于分析结果提供多种可视化图表绘制方案，但都需要依靠运行Python代码才能完成绘图任务，部分代码会出现错误

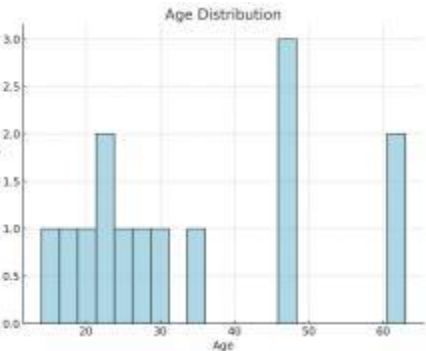
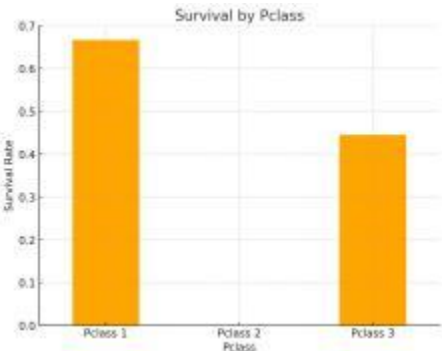
请根据以上分析结果和数据样本绘制对应的可视化图表

Open AI o3mini

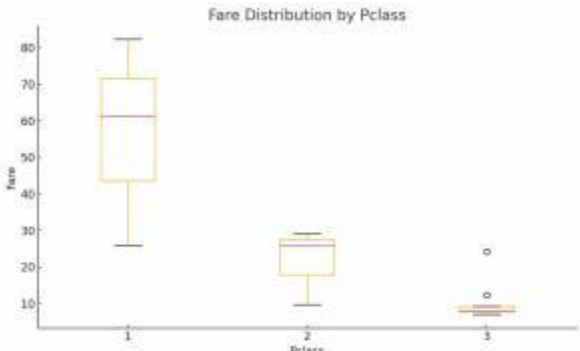
能够直接调用DALLE，根据分析结果和任务需求高效绘制各类可视化图表，部分较为复杂的图表可能出现数据错误或无法生成的情况。



柱状图（生还者和遇难者的比例、按船舱等级分类的生还情况）



年龄分布直方图、票价分布箱线图（展示不同船票等级的票价分布）



DeepSeek R1

能够结合数据样本和分析结果，提供多种可视化图表绘制方案，但暂时不能直接绘制出可视图表，需要将对应的绘图代码复制到本地运行制作图表。

3. **绘制图表**:

- **生存率按舱位**：用'groupby'按舱位分组，计算生存率，绘制柱状图。
- **性别与生存率**：用交叉表(crosstab)统计性别与幸存情况，绘制堆叠条形图。
- **年龄分布**：分箱处理年龄，绘制幸存与非幸存的比例直方图。
- **家庭规模**：分组计算生存率，绘制折线图或柱状图。
- **票价与舱位**：用箱线图展示不同舱位的票价分布。
- **登船港口**：类似舱位的处理方式，分组计算生存率。

5. 票价与舱位的关系

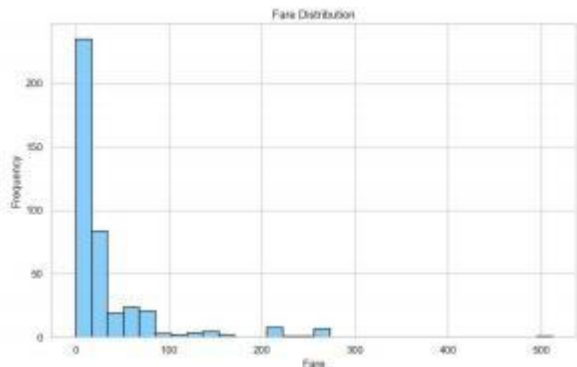
```
python
# 票价按舱位的箱线图
sns.boxplot(x='Pclass', y='Fare', data=df, palette='Set3')
plt.title('Fare Distribution by Passenger Class')
plt.xlabel('Passenger Class')
plt.ylabel('Fare')
plt.yscale('log') # 对数尺度处理票价
plt.show()
```

结论:

- 头等舱票价显著高于其他舱位（中位数约84 vs 13）。
- 高票价与高生存率正相关。

Kimi k1.5

结合数据样本和分析结果，提供多种可视化图表绘制方案及对应代码，需采用Python代码完成绘图任务。大样本会省略数据；小样本不省略数据。



Claude 3.5 sonnet

暂时不能直接绘制出可视图表，需要将绘图代码复制到本地运行。

数据应用情况总结

DeepSeek R1

- ❑ 文件数据读取完整无缺失
- ❑ 数据分析全面、逻辑清晰严谨
- ❑ 网络爬虫任务数据爬取完整、准确
- ❑ 数据挖掘能够准确分类并提供建议
- ❑ 数据可视化任务能力有待完善

Open AI o3 mini

- ❑ 数据分析高效、全面、准确
- ❑ 数据可视化能力突出、直接生成
- ❑ 网络爬虫任务爬取数据结果为空
- ❑ 暂不支持上传数据附件
- ❑ 数据挖掘深度较浅

Claude 3.5 sonnet

- ❑ 数据读取输出逻辑性强、指标清晰
- ❑ 数据分析任务完成得较为简单
- ❑ 爬虫数据采集未形成明确结论
- ❑ 数据挖掘深度较浅
- ❑ 暂不能直接绘制出可视图表

Kimi k1.5

- ❑ 数据挖掘能力出色
- ❑ 快速读取文件数据，提取网址链接
- ❑ 长文本数据处理能力突出
- ❑ 爬虫数据采集存在代码错误问题
- ❑ 数据分析能力相对较弱

新思路：优势互补，协同应用

DeepSeek+Open AI

数据采集的“天罗地网”

DeepSeek R1 负责精准爬取和筛选数据，Open AI o3mini 提供额外的数据补充



DeepSeek+Kimi

数据分析的“慧眼破局”

DeepSeek R1 负责深入分析和异常检测，Kimi k1.5 提供深度挖掘的思路，助于更精准发现数据规律



数据采集

数据预处理

数据分析

可视化呈现

Claude+DeepSeek

数据处理的“洗髓易筋”

Claude 3.5 Sonnet 在文本提取上较稳定，可用于数据清洗，DeepSeek R1 可确保数据完整性



Open AI+Kimi+Claude

数据呈现的“画龙点睛”

Open AI o3mini 直接调用 DALLE 生成图表，Kimi k1.5 提供 Python 代码支持，Claude 3.5 Sonnet 负责图表逻辑优化



新思路： DeepSeek R1的数据应用

- **中小企业AI定制化服务：**为中小企业提供定制化的AI解决方案，如智能客服、营销和办公工具，提升企业竞争力。
- **开源AI教育平台：**借助DeepSeek R1的低成本特性，创建开源AI教育平台，提供免费课程和实验资源，促进AI教育普及。
- **智能编程教育助手：**为编程学生提供实时编程指导，自动生成代码示例，帮助解决编程难题。
- **自动化代码审查工具：**自动审查代码，发现潜在问题并提供优化建议，提升开发效率与代码质量。

中文数据处理优势

低成本高性能优势

数据读取分析能力

编程代码生成能力

创意写作生成能力

- **智能中文古籍修复与注释：**利用DeepSeek R1强大的中文理解能力，自动识别并修复古籍中的破损文字，同时生成准确的注释和解释，帮助修复难以辨认的古籍内容。
- **中文法律文本分析与生成：**基于DeepSeek R1的中文数据处理能力，快速分析法律文本，提取关键信息，自动生成合同草案、法律意见书等，提高律师工作效率。
- **智能医疗数据分析与诊断：**构建智能医疗平台，分析病历、检查报告和基因数据，帮助医生提供更准确的诊断与治疗方案。
- **金融风险预测与管理：**开发金融风险分析工具，收集并分析市场数据，预测风险并为金融机构提供管理建议。
- **智能文学创作辅助：**为作家提供创作灵感和文本构思，生成符合中文文学传统的情节和诗句，助力突破创作瓶颈。
- **智能广告创意生成：**根据产品特点和目标受众自动生成创意广告文案和宣传语，提高广告创作效率。

新思路：Open AI o3mini的数据应用

- **复杂数据模式识别**：借助o3mini高效分析复杂数据，帮助科学研究和工程领域发现模式和规律，如天文学中的星系演化或地质学中的地震数据分析。
- **多源数据融合分析**：在智能交通和城市规划中，o3mini有助于将不同来源的数据（如交通流量、气象数据等）进行融合分析，预测交通拥堵，为城市规划提供决策支持。
- **交互式数据可视化**：在商业智能和数据分析领域，o3mini可以将多维数据以可视化的方式呈现，并支持用户进行交互式分析。
- **实时数据可视化与预警**：在实时监控和数据分析中，o3mini可以实时将数据以可视化的方式展示，并支持用户与数据进行交互。

推理响应速度快

数据分析效率高

格式化输出能力

数据可视化优势

写作情感表达能力

- **实时数据流处理与决策**：利用o3mini在物联网和工业自动化领域，快速处理来自传感器和设备的实时数据，进行即时分析和决策，减少停机时间，提高生产效率。
- **高频交易数据分析**：利用o3mini快速处理高频交易数据，识别市场趋势和交易模式，为交易者提供实时决策支持。
- **数据报告自动化生成**：基于o3mini自动生成格式化的数据报告，包括图表、表格和文字说明，帮助管理者快速理解分析结果。
- **数据接口标准化**：根据标准格式输出数据，利用o3mini方便不同系统和平台之间的数据共享，提升跨机构协作效率。
- **情感分析与数据解读**：利用o3mini结合情感分析，对数据进行深入解读，帮助市场调研等领域理解消费者情感，优化产品和策略。
- **故事化数据呈现**：借助o3mini将数据以故事的形式呈现，增强数据的可读性和吸引力，帮助公众理解复杂的科学和技术知识。

二

要怎么做？

撰写文章标题指令

指令： 我想让您担任学术期刊编辑，我将向您提供一份手稿摘要，您将向我提供 5 个好的研究论文英文标题，并解释为什么这个标题是好的。请将输出结果以 Markdown 表格的形式提供，表格有两列，标题为中文。第一列给出英文标题，第二列给出中文解释。以下文本为摘要：【指令后加上文章的摘要】。

我想让您担任学术期刊编辑，我将向您提供一份手稿摘要，您将向我提供 5 个好的研究论文英文标题，并解释为什么这个标题是好的。请将输出结果以 Markdown 表格的形式提供，表格有两列，标题为中文。第一列给出英文标题，第二列给出中文解释。以下文本为摘要：As human activities intensify, the structures of ecosystems and their food webs often reorganize. Through the study of mesocosms harboring a diverse benthic coastal community, we reveal that food web architecture can be inflexible under ocean warming and acidification and unable to compensate for the decline or proliferation of taxa. Key stabilizing processes, including functional redundancy, trophic compensation, and species substitution, were largely absent under future climate conditions. A trophic pyramid emerged in which biomass expanded at the base and top but contracted in the center. This structure may characterize a transitional state before collapse into shortened, bottom-heavy food webs that characterize ecosystems subject to persistent abiotic stress. We show that where food web architecture lacks adjustability, the adaptive capacity of ecosystems to global change is weak and ecosystem degradation likely.



已深度思考 (用时 21 秒)

```
markdown
| 英文标题 | 中文解释 |
|-----|-----|
| Rigid Trophic Architectures Under Climate Stress: Experimental Evidence of Ecosystem Inflexibility in Coastal Benthic Communities | 标题通过“Rigid”强调食物网刚性特征，使用冒号分隔明确研究场景，包含关键术语“Climate Stress”和“Ecosystem Inflexibility”，突出实验证据的创新性，符合海洋生态学期刊的严谨风格。 |
| The Collapse of Compensatory Mechanisms: How Ocean Warming Disables Food Web Adjustability in Coastal Ecosystems | 标题以核心发现“Compensatory Mechanisms Collapse”开篇，疑问式结构引导读者思考机制失效过程，动词“Disables”生动描述气候影响，明确研究主体为海岸带生态系统，具有学术警示价值。 |
| Climate-Driven Trophic Squeeze: Experimental Validation of Mid-Trophic Contraction in Future Marine Food Webs | 创造新术语“Trophic Squeeze”增强记忆点，“Mid-Trophic Contraction”准确概括营养级收缩现象，“Experimental Validation”强调方法学贡献，冒号结构平衡创新性与科学性，适合高影响力期刊。 |
| From Pyramids to Collapse Trajectories: Thermal Stress Induces Non-Linear Structural Transitions in Coastal Food Webs | 比喻手法(“Pyramids to Collapse”)增强可读性，使用“Non-Linear”突出非线性相变特征，“Induces”明确因果关系，双名词结构兼顾理论深度与实证发现，适用于跨学科期刊。 |
| Structural Inertia of Food Webs Under Global Change: Experimental Demonstration of Lost Adaptive Capacity in Benthic Ecosystems | “Structural Inertia”概念移植体现理论创新，“Lost Adaptive Capacity”直指研究核心结论，“Experimental Demonstration”强化证据等级，冒号后限定研究范围，符合保护生物学刊物的政策相关性需求。 |
```

中-英、英-中互译指令

指令： 我想让你充当一名科研类的**英汉**翻译， 我会向你提供一种语言的一些段落， 你的任务是把这些段落准确地、学术性地翻译成另一种语言。 翻译后不要重复原文提供的段落。 您应使用人工智能工具（如自然语言处理） 以及有关有效写作技巧的修辞知识和经验进行回复。 我会给你如下段落， 请告诉我是用什么语言写的， 然后翻译。 我希望你能以标记表的形式给出输出结果， 其中第一列是原文， 第二列是翻译后的句子， 每行只给出一个句子

所提供段落的语言是中文， 以下是按要求的标记表格式翻译成英文的译文：

Original (Chinese)	Translation (English)
捕食是一个基本的生态过程， 捕食的定义为： 一种生物（捕食者） 捕食了另一种生物（猎物） （Begon等， 1997）。	Predation is a fundamental ecological process,defined as one organism (predator) preying onanother organism (prey) (Begon et al., 1997).
在群落范围内， 捕食能够影响某一营养等级的动态， 也能够影响整个群落结构的动态。	Within a community, predation can affect thedynamics of a specific trophic level as well as thedynamics of the entire community structure.
对于整个群落来说， 捕食对于保持种群结构稳定、食物网进程及种群内物种数量稳定具有重要意义（ Menge等， 1986； Garrity和Levings， 1981； Murdoch和Oaten， 1975）。	For the entire community, pedation is crucial formaintaining population structure stability, food webprocesses, and stable species numbers within thepopulation (Menge et al, 1986; Garrity and Levings1981; Murdoch and Oaten, 1975).
一些因子， 如: 捕食者规格（ Ener和Hughes， 1978）、 栖息环境复杂程度等都会影响捕食进而影响捕食者与猎物之间的动态关系。	Some factors, such as predator size (Elner andHughes, 1978) and habitat complexity, can affectpredation and subsequently influence the dynamicrelationship between predators and prey.

中文学术写作润色指令

指令： 作为中文学术论文写作优化助手，您的任务是改进所提供文本的拼写、语法、清晰度、简洁性和整体可读性，同时分解长句，减少重复，并提供改进建议。请仅提供文本的更正版本，并附上解释。以 markdown 表格的形式提供输出结果，每个句子单独成行。第一列为原句，第二列为修改后的句子，第三列为中文解释。请编辑以下文本：

原始句子	优化后句子	中文释义
捕食是一个基本的生态过程，捕食的定义为：一种生物(捕食者)捕食了另一种生物(猎物)(Begon等，1997)。	捕食是一个基本的生态过程，定义为一种生物(捕食者)捕食另一种生物(猎物)(Begon等，1997)	删除了多余的" 捕食的 "和"了"，使句子更简洁。
在群落范围内，捕食能够影响某一营养等级的动态，也能够影响整个群落结构的动态。	在群落范围内，捕食能影响某一营养等级及整个群落结构的动态。	删除了多余的"能够"，并将两个动态合并在一起，使句子更简洁。
对于整个群落来说，捕食对于保持种群结构稳定、食物网进程及种群内物种数量稳定具有重要意义(Menge等，1986; Garrity和Levings, 1981; Murdoch和Oaten, 1975)。	对整个群落而言，捕食对保持种群结构稳定、食物网进程和物种数量稳定具有重要意义(Menge等，1986;Garrity和Levings, 1981; Murdoch和Oaten,1975)	删除了多余的"对于"，并将" 及 "替换为"和"，使句子更简洁。
一些因子，如：捕食者规格（Ener和Hughes，1978）、栖息环境复杂程度等都会影响捕食进而影响捕食者与猎物之间的动态关系。	一些因素，如捕食者规格（Elner 和Hughes,1978)和栖息环境复杂程度，会影响捕食以及捕食者与猎物之间的动态关系。	将"因子"替换为"因素"，删除了多余的" 如:"和"等"

英文学术写作润色指令

指令： 下面是一篇学术论文中的一个段落。润色文字以符合学术风格，改进拼写、语法、清晰度、简洁性和整体可读性。必要时，重写整个句子。此外，用标记符表格列出所有修改，并解释修改原因。

原始文本		修正后文本	
Numerous marine molluscs, including both gastropods and bivalves, show phenotypic plasticity in their shell morphology in response to predation risk (Appleton & Palmer1988, Trussell & Smith 2000, Bourdeau 2010). Predation can therefore act as a directional selection stressor that produces particular prey phenotypes, such as thicker and heavier shells, round shell shape, and lower growth rate (Leonard et al. 1999,Trussell & Nicklin 2002, Hollander & Butlin 2010, Bourdeau 2010) that can decrease the susceptibility to predation.		Numerous marine mollusks, including both gastropods and bivalves, exhibit phenotypic plasticity in their shell morphology in response to predation risk. Predation can act as a directional selection pressure, resulting in specific prey phenotypes, such as thicker and heavier shells, rounder shell shapes, and lower growth rates that can reduce susceptibility to predation (Appleton & Palmer, 1988; Trussell & Smith, 2000; Leonard et al., 1999; Trussell & Nicklin, 2002; Hollander & Butlin, 2010; Bourdeau, 2010).	
修改	解释	修改	解释
将 “molluscs” 替换为 “mollusks”	将 “molluscs” 替换为 “mollusks”	用 “可以降低捕食的易感性” 代替 “可以降低捕食的易感性” 将 “产生 ” 改为 “导致” 用 “特定的猎物表型 ” 代替 “特定的猎物表型”	使用更简洁的动词，更好地表达原意。
将 “show” 替换为 “exhibit”	将 “show” 替换为 “exhibit”		使用更准确、更简洁的动词，更好地表达捕食与贝壳形态之间的因果关系。
重新组织句子结构	重新组织句子结构		
在每个引用年份后添加逗号	在每个引用年份后添加逗号		

提升段落间逻辑和连贯性指令

指令： 请分析以下文本中每个段落中句子之间的逻辑性和连贯性，指出句子之间的流畅性或关联性有哪些地方可以改进，并提出具体建议，以提高内容的整体质量和可读性。请只提供改进后的文本，然后用中文列出改进之处。请改进以下文字：

原始文本	修正后文本
Over the past several decades, with the explosive growth of renewable energy, large-scale energy storage technologies allow intermittent renewable energy to replace traditional energy. High-performance energy storage technologies are the most promising candidates for large-scale energy storage. Since commercialization, lithium-ion batteries have become mainstream energy storage devices with their high energy density, and long cycle life. In order to meet the demand for even better electrochemical performance, researchers are exploring sustainable anode materials that provide lithium-ion battery performance, while providing high capacity and long cycle life. As a promising option, alloy-based anodes such as silicon (Si, 4200 mAh g⁻¹ theoretical capacity, nearly 10 times higher than graphite anodes (372 mAh g⁻¹)). Unfortunately, silicon-based anodes suffer from large volume expansion/shrinkage during lithiation/delithiation, leading to the formation of cracks, separation of the active material from the current collector, and disruption of the electronic contact within the electrode, ultimately resulting in a sharp attenuation of cycle life. In order to overcome these problems, previous research has put a lot of effort into improving electrode durability.	Over the past several decades, renewable energy has seen explosive growth, and large-scale energy storage technologies have played a vital role in enabling intermittent renewable energy sources to replace traditional energy. Among these technologies, high-performance secondary battery solutions. Lithium-ion batteries (LIBs), in particular, have become mainstream energy storage devices since their commercialization due to their high energy density, and long cycle life. In order to meet the demand for even better electrochemical performance, researchers are exploring sustainable anode materials. The goal is to find materials that not only have a safe and stable cyclic performance, but also offer high capacity. Among various options, alloy-based anodes, especially silicon (Si), have attracted in-depth research and discussion. This is primarily due to their high theoretical capacity, which is nearly 10 times that of commercial graphite anodes. Unfortunately, silicon-based anodes have significant drawbacks. The main issue lies in the large volume expansion/shrinkage during the lithiation and delithiation process. This can lead to the active material separating from the current collector, and a loss of electronic contact within the electrode. All of these issues can cause a sharp attenuation of cycle life. Therefore, much of the previous research has focused on improving the durability of the electrode to overcome these problems.

Improvements in Chinese

- 将原本含糊不清的句子进行细化，增加了文本的流畅度和连贯性。如，对第一段的前两句进行了明确的表达。
- 对一些过于复杂的句子进行了简化，使之更易于理解。如，对第二段的前两句进行了分解和简化。
- 对原文中含有的技术性术语进行了更明确的解释，使读者能更好地理解文本内容。如，对第三段的硅(Si)和商业石墨阳极的理论容量进行了详细的解释。
- 对存在的问题进行了分段讨论，使其更加清晰明了。例如，将第四段的问题和解决方法进行了分段讨论。
- 增加了转折语气，使文本在描述问题和解决方案时更具说服力。如，对第四段的问题和研究聚焦进行了明确的表达。

标点符号错误一键修改指令

指令： 请你帮我确保以下文字中的标点符号正确无误，如果发现标点符号错误，请将你发现的错误列在一个两栏的标记表中，将原文放在第一栏，将更正后的文本放在第二栏，请只列出错误，并在最后提供更正后的文本。

原始文本	修正后文本
In addition. theleft valve of each mussel was examined for compressive force (shell strength)following Burnett and Belk (2018). A universal material+testing machine(MTS System Corporation, Eden Prairie, MIN, USA, Model 661; Fig1,)was used to determine the shell strength. Each shell valve was placed horizontally with the shell edge on a flat surface, while a compressive force was applied at a constant loading rate of 10 mm-min until the real-time force curve on the monitor screen fast drop indicating failure occurred.	In addition, the left valve of each mussel was examined for compressive force (shell strength)following Burnett and Belk (2018). A universal material-testing machine (MTS System Corporation, Eden Prairie, MN, USA, Model 661; Fig. 1) was used to determine the shell strength. Each shell valve was placed horizontally with the shell edge on a flat surface, while a compressive force was applied at a constant loading rate of 10 mm/min until the real-time force curve on the monitor screen fast drop indicating failure occurred.

改写降重指令

指令：我想让你充当科研写作专家，并提供一些英文或中文段落，你的任务是用原文改写段落。你应该使用人工智能工具（如自然语言处理）、修辞知识和你在有效科学写作技巧方面的专业知识来回答。请只提供改写后的文本，不作任何解释，请用科研语气风格重写下面的文字：

I want you to act as an expert in scientific writing, I will provide you with some paragraphs in English or Chinese and your task is rewrite the paragraph in the original language. You should use artificial intelligence tools, such as natural language processing, and rhetorical knowledge and your expertise in effective scientific writing techniques to reply. Please only provide the text after rewriting without any explanation. Please rewrite the following text in a scientific tone:

Our study demonstrates substantial benefits of multispecies plantations, which corroborates and generalizes findings from forest BEF experiments (14, 16, 21) and helps to show their implications for real-world forestry applications. Our analyses offer predictions that go beyond

 已深度思考 (用时 28 秒) ▾

This investigation establishes significant ecological advantages associated with multispecies plantation systems, empirically validating and extending foundational biodiversity-ecosystem functioning (BEF) principles derived from controlled forest experiments (14, 16, 21) while demonstrating translational relevance for silvicultural practices. Our predictive modeling framework reveals critical limitations inherent in conventional BEF experimental designs characterized by randomized species composition and temporal constraints, thereby underscoring the necessity for implementing extended experimental frameworks employing deliberately curated species assemblages to elucidate biodiversity-mediated temporal dynamics in ecosystem productivity.

我想让你充当科研写作专家，并提供一些英文或中文段落，你的任务是用原文改写段落。你应该使用人工智能工具（如自然语言处理）、修辞知识和你在有效科学写作技巧方面的专业知识来回答。请只提供改写后的文本，不作任何解释，请用科研语言风格重写下面的文字：

为了更加清晰的了解捕食者的捕食模式，研究者们引入了捕食周期理论描述捕食者捕食过程中的捕食行为。这一理论可以清晰的表述捕食过程中捕食者是如何搜寻、攻击、捕获和摄食猎物的（Holling, 1966; O'Brien, 1979; Wong, 2005）。基于这一理论，捕食率可以被拆分为一系列行为概率：捕食者和猎物的相遇概率、相遇后攻击概率、攻击后捕获概率和捕获后摄食概率（Holling, 1966；O'Brien, 1979）。这些相遇行为可以与捕食选择（主动选择与被动选择）联系起来。

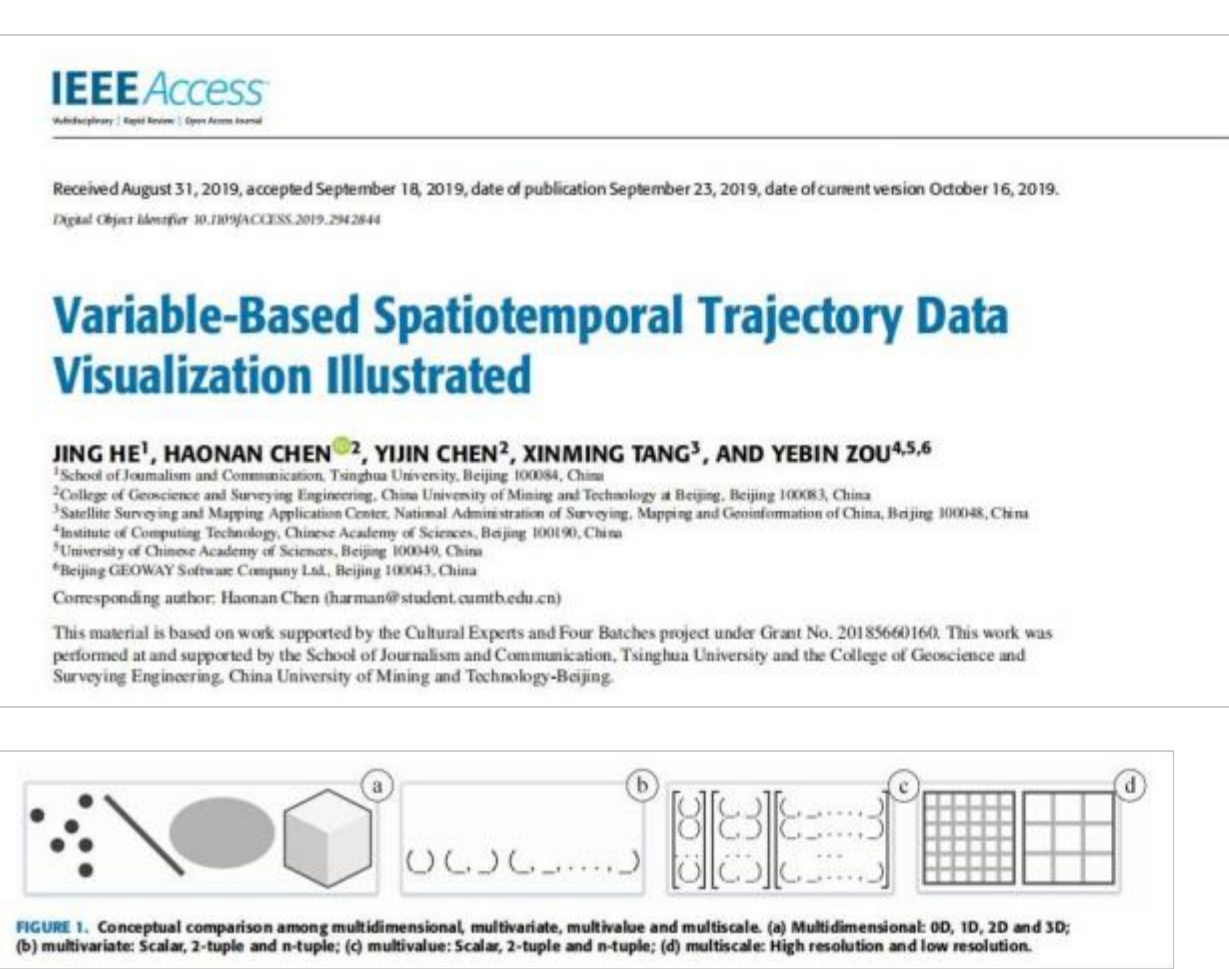
 已深度思考 (用时 26 秒) ▾

为系统阐释捕食者的猎物获取机制，学界通过捕食周期理论对捕食行为序列进行建模分析（Holling, 1966; O'Brien, 1979; Wong, 2005）。该理论框架将捕食过程解构为四个连续阶段：猎物搜寻定位、攻击行为启动、猎物捕获实施及能量摄取完成。通过概率链建模方法，总捕食率可分解为若干连续行为事件的概率乘积，具体包括捕食者-猎物空间相遇概率、遭遇后攻击决策概率、攻击成功概率及最终摄食概率（Holling, 1966; O'Brien, 1979）。此类量化行为参数为解析捕食策略选择提供了数学模型基础，特别是主动选择偏好与被动机遇捕食的机制区分。

解读文献配图指令

指令：这是发表在【杂志名称】期刊上的一篇文章中的一幅图，标题为【文章标题】，图例为【图的标题】。作为该领域的专家，请帮助我解读该图。



新对话

This is a figure from a paper published on the journal of [IEEE Access] with the title[Variable-Based Spatiotemporal Trajectory Data Visualization Illustrated].The legend of the figure is [FIGURE 1. Conceptual comparison among multidimensional, multivariate, multivalued and multiscale. (a) Multidimensional: 0D, 1D, 2D and 3D; (b) multivariate: Scalar, 2-tuple and n-tuple; (c) multivalued: Scalar, 2-tuple and n-tuple; (d) multiscale: High resolution and low resolution.]. Please help me to interpret the figure as an expert in that field.

图 已深度思考 (用时 56 秒) ^

Okay, I need to interpret Figure 1 from this IEEE Access paper titled "Variable-Based Spatiotemporal Trajectory Data Visualization Illustrated." The figure is a conceptual comparison of multidimensional, multivariate, multivalued, and multiscale. Let me start by recalling what each of these terms generally means in data visualization and spatiotemporal contexts.

First, multidimensional (a). The figure shows 0D, 1D, 2D, and 3D. Multidimensional data usually refers to data with multiple dimensions, like spatial dimensions (x, y, z) or other axes. In spatiotemporal data, time is often a dimension. So 0D might be a single point, 1D a line (like time series), 2D a plane (maybe latitude and longitude), and 3D adding height or depth. This part is probably about how data can exist in different dimensional spaces.

Next, multivariate (b). The figure mentions scalar, 2-tuple, and n-tuple. Multivariate data involves multiple variables. A scalar is a single value, 2-tuple could be two variables (e.g., temperature and pressure), and n-tuple

给 DeepSeek 发送消息

深度思考 (R1)

联网搜索

论文参考文献格式指令

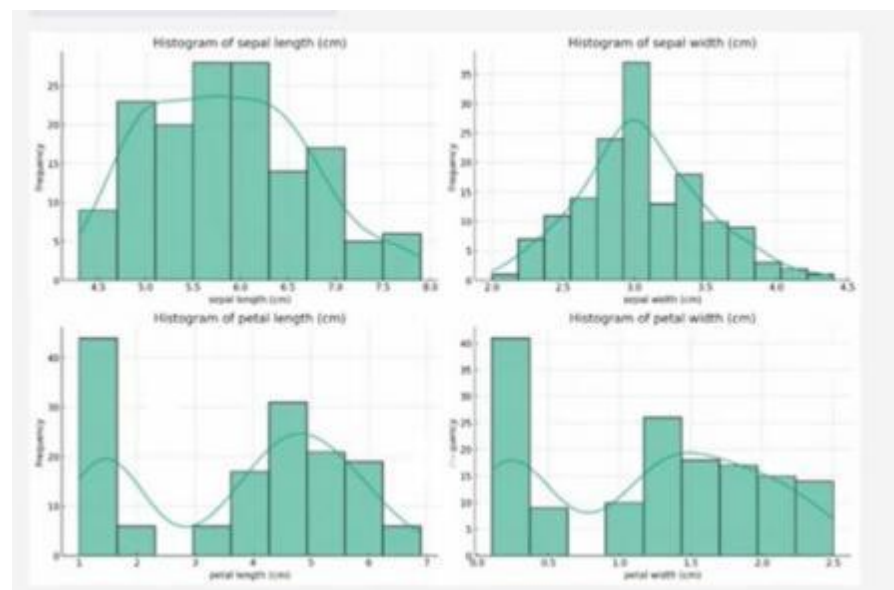
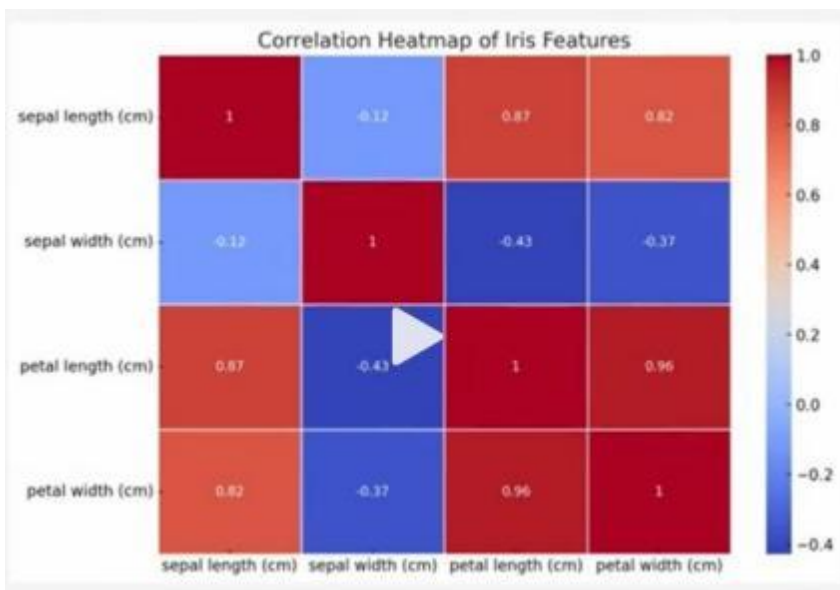
指令： 我想请你担任一份研究手稿的参考文献编辑。我将为你提供五个参考文献模板， 你应将其作为指南。之后，我会提供更多参考文献， 你需要检查这些参考文献的格式问题， 如标点符号的位置和间距。 给出一个包含三列的标记表， 第一列是原文， 第二列是固定文本， 第三列是解释， 然后提供所有固定的参考文献。 以下是需要修正的五个示例模板和参考文献：

Template:
1. Appleton RD, Palmer AR (1988) Water-borne stimuli released by predatory crabs and damaged prey induce more predator-resistant shells in a marine gastropod. *Proc Natl Acad Sci USA* 85:4387–4391.
2. Bibby R, Cleall-Harding P, Rundle S, Widdicombe S, Spicer J (2007) Ocean acidification disrupts induced defences in the intertidal gastropod *Littorina littorea*. *Biol Lett* 3:699–701.
3. Bible JM, Griffith KR, Sanford E (2017) Inducible defenses in *Olympia* oysters in response to an invasive predator. *Oecologia* 183:809–819.
4. Blumstein DT, Daniel JC, Springett BP (2004) A test of the multi-predator hypothesis: rapid loss of antipredator behavior after 130 years of isolation. *Ethology* 110:919–934.
5. Bonsdorff E (2006) Zoobenthic diversity-gradients in the Baltic Sea: Continuous post-glacial succession in a stressed ecosystem. *Journal of Experimental Marine Biology and Ecology* 330:383–391.
Reference needed improve:
Boullis A, Fassotte B, Sarles L, Lognay G, Heuskin S, Vanderplanck M, Bartram S, Haubruge E, Francis F, Verheggen FJ (2017) Elevated Carbon Dioxide Concentration Reduces Alarm Signaling in Aphids. *J Chem Ecol* 43:164–171.
Bourdeau PE (2010) An inducible morphological defence is a passive by-product of behaviour in a marine snail. *Proc R Soc B* 277:455–462.
Bourdeau PE (2011) Constitutive and inducible defensive traits in co-occurring marine snails distributed across a vertical rocky intertidal gradient. *Functional Ecology* 25:177–185.
Brönmark C, Hansson L-A (eds) (2012) Chemical ecology in aquatic systems. Oxford University Press, Oxford; New York.
Brönmark C, Lakowitz T, Hollander J (2011) Predator-induced morphological plasticity across local populations of a freshwater snail. *PLOS ONE* 6:e21773.
Burnett NP, Belk A (2018) Compressive strength of *Mytilus californianus* shell is time-dependent and can influence the potential foraging strategies of predators. *Mar Biol* 165:42.

原始文本	修正后文本
Boullis A, Fassotte B, Sarles L,LognayG, Heuskin S, Vanderplanck M.Bartram S, Haubruge E, Francis F,Verheggen F(2017) Elevated Carbon Dioxide Concentration Reduces Alarm Signaling in Aphids. <i>J Chem Ecol</i> 43:164-171.	Boullis A, Fassotte B, Sarles L, Lognay G, Heuskin S, Vanderplanck M.Bartram S, Haubruge E, Francis F.Verheggen FJ (2017) Elevated CarbonDioxide Concentration Reduces Alarm Signaling in Aphids. <i>J Chem Ecol</i> 43:164-171.
Bourdeau PE(2011) Constitutive and inducible defensive traits in co-occurring marine snails distributed across a vertical rocky intertidal gradient. <i>Functional Ecology</i> 25:177-185	Bourdeau PE(2011) Constitutive and inducible defensive traits in co-occurring marine snails distributed across a vertical rocky intertidal gradient. Funct Ecol 25:177-185.
Brnmark C, Lakowitz T Hollander J(2011) Predator-induced morphological plasticity across local populations of a freshwater snail. <i>PLOS ONE</i> 6:e21773.	Bronmark C. Lakowitz T ,Hollander J(2011) Predator-induced morphological plasticity across local populations of a freshwater snail. <i>PLOS ONE</i> 6:e21773.

高阶数据分析指令

指令：请你制作 10 个不同的可视化图表，以反映上传数据集的各个方面？[适当加上一些对你的数据的描述，尤其是各个列（变量），这里尽量详细的描述你的数据]



- 用 10 张精美的绘图，展示数据变量之间的关系或潜在的分析思路，包括展示数据分析的直方图和箱线图。
- 呈现关键变量之间的配对散点矩阵图和小提琴图，以及常用的相关性热图等，每张图都有详细的解释。

其他常用英文指令

Prompts （指令）	描述	Prompts （指令）	描述
Can you load and preview the data?	加载，预览一下数据	Can you list the top 10 key points?	最重要的十个要点
What are the trends shown in this data?	找趋势	Can you describe the data?	描述数据
Show me the top trends in a visual format.	以视觉形式显示趋势	Can you clean this dataset?	清洗数据
Can you create a heatmap using this data?	创建一个热力图	Can you segment this data and create a table?	切分数据
Can you create a graph using this data?	制作一个图	Can you create a word cloud?	做一个词云
Can you create a chart using this data?	画一个图表	What are the rows and columns in this dataset?	描述一下行和列
Can you make the graphs more beautiful?	把图美化一下	Can you write a one sentence recap of this data?	快速回顾一下
Create a visual chart, based on this data.	做一个视觉图表	What’s the main takeaway from this dataset?	找出最主要的信息
Can you explain this dataset like I’m 5 years old?	像给五岁小朋友讲故事那样解释一下这个数据集	Can you create a presentation based this dataset?	做一个整体展示
Can you create 10 graphs to present different data?	创作10个不同的图展示数据	Can you write me an article based on this dataset or statistic results?	根据结果写文章
Can you explain this dataset in one paragraph?	用一段话来解释一下这个数据集	What insights do you see here? Give me a numbered list.	提供一些见解
Can you explain this dataset in simple terms?	用简单的话来解释一下这个数据集		

其他常用中文指令

Prompts（指令）
跨学科融合： 将“舆论分析”概念与其他领域的最新具有突破性的理论深度结合， 提出极其具有创新的交叉领域的十个问题。
探索“舆论分析”概念的基础理论、 哲学基础或科学原理等深层次原理， 提出挑战这些基础的前所未有的突破性十个问题。
舆论分析这个概念在最前沿科技或理论中的潜在应用， 列出十个充满想象力和震撼性， 前所未有的应用。
如果要量化研究审美智能概念， 请提出一个合理的， 有效的， 各指标不重叠的， 你自己能提取数据的指数体系框架， 不少于三十个指数。
请大家研究任何问题， 先用这四个提示词进行提问。一是跨学科融合， 二是深层次原理， 三是概念前沿应用， 四是如何量化分析。任何学术概念。
里面会有些冗余信息， 可以删除回复中的冗余信息。另外大家有空还可以对我的提示词进行改进， 围绕四个方面。我们需要建立一套研究提示词集。
AI for research 提示词集。

三

效果如何？

元知AI综述工具

产概况

元知是国内由清华、北航专家团队研发的__个AI学术平台，目前其AI综述生成工具已开放使用，能够帮助用户从海量文献中提取核心信息，通过自然语言处理算法，实现从文献梳理到观点提取到研究评论的一键式全自动生成。

功能亮点

- ❑ **多版本与模块化支持：** 目前提供三个版本（基础版、增强版、专业版），能够灵活应对不同用户的综述需求。工具内包括文献观点梳理、问题提出等功能模块，确保用户在不同科研需求下得到充分支持。
- ❑ **增强版绘图功能：** 增强版具备绘图功能，可通过可视化图示（如文献关键词共现图）直观展示综述内容，帮助用户更好理解和呈现研究成果。
- ❑ **无数据检索：** 以现有真实数据库作为支撑，通过关键词检索，自动搜集相关文献并生成综述报告，目前只支持英文检索。
- ❑ **低重复率：** 结合现有查重机制与AI技术，在内容生成阶段引入重复检测与优化策略，从源头上降低重复率风险，所生成的综述普通重复率与AIGC重复率均在5%以下。
- ❑ **无限双语数据导入：** 支持中文与英文文献的导入，并且文献数据量没有限制，能够轻松处理中文文献的系统性梳理，以及国际文献的跨语言分析。
- ❑ **幻觉克服：** 以现有真实数据库作为支撑，借助由专家设计撰写的提示词，精准规避AI生成中的幻觉问题。
- ❑ **高规范格式输出：** 所生成的综述文档格式规范、结构清晰，符合学术论文标准，用户几乎无需进行二次整理。

中科院PubScholar平台

产品概况

“PubScholar”平台是由中国科学院开发的公益学术平台，整合了国内外多种学术资源。该平台提供文献检索、引用分析、文献推荐等功能，用户可通过平台高效获取科研资源，并生成相关的综述报告。平台的优势在于其广泛的数据源和智能化的文献推荐系统，支持跨学科文献分析。

功能亮点

- ❑ **免费开放使用：** 所有用户均可免费访问，注册后可直接使用。
- ❑ **海量学术资源整合：** 包含约1.8亿条学术元数据，涵盖科技论文、专利文献、科学数据等多个类别。超过8000万篇资源可直接获取全文，包括2122万篇论文全文和5878万篇专利全文。
- ❑ **无数据检索：** 以现有真实数据库作为支撑，通过关键词检索，自动搜集相关文献并生成综述报告，支持中、英文检索。

知网研学平台

产品概况

“PubScholar”平台是由中国科学院开发的公益学术平台，整合了国内外多种学术资源。该平台提供文献检索、引用分析、文献推荐等功能，用户可通过平台高效获取科研资源，并生成相关的综述报告。平台的优势在于其广泛的数据源和智能化的文献推荐系统，支持跨学科的文献分析。

功能亮点

- **较高格式规范输出：**根据学术规范自动排版，生成符合论文要求的文献综述结构。
- **中文内容丰富：**在中文文献的分析上具有优势，能够详细呈现中文领域的研究成果，用户可手动选择想要分析的50篇文献。
- **无数据检索：**以中国知网数据库作为支撑，通过关键词检索，自动搜集相关文献并生成综述报告，仅支持中文检索。

斯坦福STORM

产品概况

斯坦福STORM平台是由斯坦福大学的oval团队开发的的__款AI科研工具，其核心功能是通过多智能体协作，实现从提纲到段落再到文章的迭代式生成，为用户生成内容大纲及高质量长文本。

功能亮点

- ❑ **资料整合与文章生成：**能够浏览网络，搜集大量文献，并通过基于主题的多个智能代理，将这些文献转化为连贯的文章或研究论文，长度可达数万字。
- ❑ **转化文献为连贯文章：**可以将现有的文献资料进行分析和整合，转化为逻辑连贯的新文章，为学者和知识工作者提供了极大的便利。
- ❑ **模拟对话与问题生成：**模拟文章写作前的调研过程，通过发掘话题研究中的多样视角，模拟具有不同视角的作者向话题专家提出问题的对话，并基于这些对话整理收集到的信息来创建文章大纲。
- ❑ **多智能体协作对话：**Co-STORM模式引入了协作对话机制，并采用轮次管理策略，实现流畅的协作式AI学术研究。

用户体验对比：使用步骤

整体来看，**元知AI综述工具**提供了一键式的自动化流程，只需导入数据，即可自动生成高质量且规范的文献综述，适合快速高效的研究需求。

元知AI综述工具



PubScholar平台



知网研学平台



斯坦福STORM

元知AI综述工具官网：<http://118.31.250.10:8081/#/>

- **选择版本：**根据需求选择工具的四个版本，包括基础版、增强版、专业版（单图）、专业版（双图）。
- **文献导入：**用户可从现有文献数据库中下载中英文数据后导入平台，或直接通过实时联网访问免费数据库进行在线分析，操作简单便捷。
- **信息提取与分析：**平台自动运用AI技术对导入的文献进行关键信息提取和深度梳理分析，用户无需进行复杂操作，等待平台处理完成即可。
- **综述生成：**根据智能分析结果，平台自动生成结构化的文献综述文本内容和可视化图表，用户可直接获取完整的综述报告，也可根据需要进行自定义调整，如综述主题、目标、参数等。

PubScholar平台官网：<https://pubscholar.cn/>

- **输入关键词：**进入官网后，在搜索框键入关键词进行文献检索。
- **选取文章：**勾选想要分析的20篇文献。
- **综述生成：**点击生成综述，等待2-3分钟即可下载综述报告。

知网研学平台官网：<https://aiplus.cnki.net/sumup/sumup>

- **输入关键词：**进入官网后，在搜索框键入关键词进行文献检索。
- **选取文章：**勾选想要分析的20篇文献。
- **综述生成：**点击生成综述，等待2-3分钟即可下载综述报告。

斯坦福STORM官网：<https://storm.genie.stanford.edu/>

- **选择模式：**进入主页后，用户可选择STORM或Co-STORM模式。
- **输入主题：**直接输入主题词后，STORM开始进行信息检索和文章生成。
- **查看生成过程：**点击“See BrainSTORMing Process”，可获取不同LLM Role的头脑风暴过程。
- **参考其他文章：**在“发现”栏，可参考其他学者生成的文章及聊天示例。

用户体验对比：可操作性

元知AI综述工具

- ❑ **界面直观**：平台设计简洁、直观，使用户能够方便、快捷地进行文献数据的导入、分析和综述生成，操作路径清晰，交互体验流畅高效。
- ❑ **模块分区**：将功能模块与信息展示分区设计布局，用户可以轻松找到所需功能，提高了操作的便捷性和效率。
- ❑ **多语言支持与定制化设置**：语言支持对于国内研究者更为友好，能够适应综述撰写的国内外研究需求，同时定制化设置满足用户在个性化需求下的使用。



知网研学平台

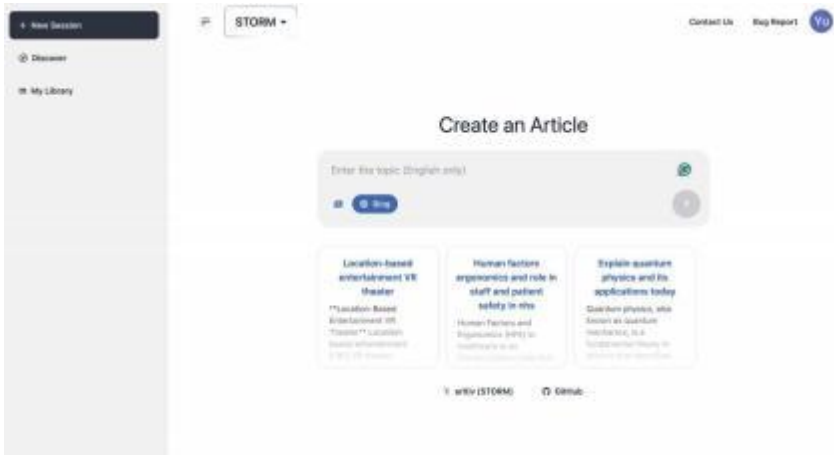
- ❑ **界面直观**：平台设计简洁、直观，使用户能够快捷地进行文献数据的检索、选取和综述生成，操作路径清晰，交互体验流畅高效。
- ❑ **语言支持**：支持英文和中文内容生成。

PubScholar平台

- ❑ **界面直观**：平台设计简洁、直观，使用户能够快捷地进行文献数据的检索、选取和综述生成，操作路径清晰，交互体验流畅高效。
- ❑ **语言支持**：支持英文和中文内容生成。

斯坦福STORM

- ❑ **界面友好**：操作界面简洁明了，用户容易上手，非技术背景用户也能快速学会使用该工具进行文献综述的生成。
- ❑ **灵活定制**：支持使用本地部署的语言模型，为有特定需求的用户提供了更多灵活性。
- ❑ **实时演示**：提供实时演示demo，方便用户了解和体验工具的功能。
- ❑ **语言支持**：仅支持英文输入和英文内容生成。



生成综述对比： 多维度对比

PS：使用感受会因个体差异而有不同，仅作参考

维度	支持篇数	格式评分	文献类别	支持绘图	提炼观点
元知基础版	不限篇数	4分	中文英文	不支持	基础提炼
元知增强版	不限篇数	5分	中文英文	支持	深入提炼
元知专业版 (单图)	不限篇数	5分	中文英文	支持	全面提炼
元知专业版 (双图)	不限篇数	5分	中文英文	支持	全面提炼
PubScholar	20篇	2分	中文英文	不支持	未有提炼
知网研学	50篇	4分	中文	不支持	未有提炼
斯坦福 STORM	不限篇数	3分	英文	不支持	基础提炼

生成综述对比：准确性与专业性

PS：使用感受会因个体差异而有不同，仅作参考

AI综述平台	元知AI综述工具	PubScholar平台	知网研学平台	斯坦福STORM
数据来源	依托真实且可靠的学术数据库，确保文献数据的准确性与可信度，为综述内容的真实性提供坚实保障	涵盖全球科技论文、专利文献、科学数据、学位论文、预印本、图书专著及开放资源	中国知网数据库，涵盖海量的中文文献	通过必应搜索引擎收集数据，确保来源的广泛性，但主要依赖互联网主流来源，可能包含推广内容，需进一步筛选和验证
文本类型	文本更加贴近学术综述，内容涵盖了研究现状、简要评述和主要参考文献，结构完整，生成文本适合辅助学术研究和论文撰写	文本较为学术，内容涵盖引言、各层面的分析，总结与展望、参考文献	文本贴近学术综述，内容涵盖了引言、研究现状、参考文献	文本倾向于事实现状，内容包括历史背景、当前趋势、应用领域、挑战与局限、未来方向等，结构清晰，适合用于行业分析和趋势预测
参考文献管理	参考文献数量相对更多，涵盖国内外学术文献，参考文献标注格式规范，引用的文献来自真实数据源，确保内容的准确性和可靠性	参考文献20篇以内，涵盖国内外学术文献，参考文献标注格式不规范，引用的文献来自真实数据源，确保内容的准确性和可靠性	参考文献50篇以内，只涵盖国内学术文献，参考文献标注格式较为规范，引用的文献来自真实数据源，确保内容的准确性和可靠性	生成内容引用了高质量的学术文献和行业报告，参考文献标注格式规范，支持直接点击标注查看参考来源，便于追溯

生成综述对比：逻辑性与结构性

PS：使用感受会因个体差异而有不同，仅作参考

AI综述平台	元知AI综述工具	PubScholar平台	知网研学平台	斯坦福STORM
语言逻辑	语言逻辑清晰，条理分明，各部分之间过渡自然，逻辑连贯。在研究现状部分，按照不同研究领域和主题进行分类，逻辑性强	报告整体呈现出总分总的逻辑架构，语言描述清晰，避免冗长，使用简短的句子表达复杂的信息	报告整体架构严谨，以引言、技术原理、应用现状、技术挑战、未来展望等部分进行层层递进。语言中多使用中性描述，客观呈现研究进展与问题	语言逻辑严谨，条理清晰，各部分之间逻辑关系明确。在历史背景和当前趋势部分，按照时间顺序和技术创新进行分类，逻辑性强
	通过逻辑排序、层次化分段和观点与事实的清晰区分，确保生成的内容符合学术写作标准。内容结构完整，包括研究现状、简要评述和主要参考文献等板块。同时，研究现状部分围绕研究主题进一步细分为多个研究层次，结构合理	内容结构完整，格式较一般	综述结构较为标准，在中文文献分析上具有优势	在写作前，系统会先生成详细的写作大纲，为文章的结构提供清晰的框架。文本内容结构清晰，包括历史背景、当前趋势、应用领域、挑战与局限、未来方向。每个部分都有详细的子标题，结构合理，层次分明

生成综述对比：完整性与全面性

PS：使用感受会因个体差异而有不同，仅作参考

AI综述平台	元知AI综述工具	PubScholar平台	知网研学平台	斯坦福STORM
文本长度	文本长度较长，内容丰富，涵盖了多个研究领域和研究层次，提供了详细的分析和评述	文本长度中等长度，内容较为丰富，也分了多个层次进行总结	文本长度稍长，内容丰富性在中文文献的分析上具有优势，能够详细呈现中文领域的研究成果	文本长度适中，内容精炼，重点突出，适合快速阅读和理解
研究视角	研究视角多样，从不同领域和研究层面出发，提供了全面的分析和评述	研究视角多样，从不同领域和研究层面出发	研究视角多样，从不同领域和研究层面出发	研究视角集中，从研究主题相关的历史背景、当前趋势、应用领域、挑战与局限、未来方向等方面进行深入分析

生成综述对比：可读性与实用性

PS：使用感受会因个体差异而有不同，仅作参考

AI综述平台	元知AI综述工具	PubScholar平台	知网研学平台	斯坦福STORM
引用格式规范	生成的引用格式标准且规范，能够清晰准确地列出参考文献，符合学术出版的要求，确保文献的格式符合高水平的学术标准	引用格式较为简化，虽然能提供基本的引用信息，但在一些细节上缺少学术规范，适合较为基础的文献综述	引用格式较为标准，尤其在中文文献的引用上符合国内学术出版的常见格式，适合中文领域的文献综述	文章语言类似维基百科风格，具有百科的深度和结构，可读性强，适合不同层次的读者阅读和理解
自动生成参考文献	能够自动生成参考文献列表，确保引用格式的统一性，确保与文中引用——对应	生成的参考文献信息不完全，且格式较为简化，不符合学术引用的标准，在学术规范方面存在一定不足	提供自动生成参考文献的功能，在中文文献的引用格式上比较标准，能够确保格式的规范化	Co-STORM通过多智能体协作对话生成动态思维导图，帮助用户发现信息盲点并组织内容，进一步提升了综述的完整性和全面性

综上所述，在生成综述的**准确性、逻辑性、完整性及可读性**方面，

- **元知AI综述工具**依托于真实的学术数据库，具备较强的学术性和深度，尤其在学术研究领域具有较大优势。其生成内容结构严谨、内容全面，并提供了有益的可视化工具，增强了综述的直观性与理解度。
- **PubScholar平台**在内容多样性和研究视角多样性方面表现良好，但引用格式和参考文献管理方面有待提高。
- **知网研学平台**则在中文文献分析和结构化内容生成方面具有优势，引用格式较为规范。
- **斯坦福STORM**在行业报告的生成中表现突出，能够提供较为简洁且针对性的内容，适合行业趋势分析与快速阅读。

生成综述案例：元知（增强版）AI综述工具

大语言模型传播偏向规制与风险治理：以 ChatGPT 为例

一、研究现状

1. 大语言模型传播规制研究层面

重点关注大语言模型在不同领域的应用与挑战。郑春萍等（2024）提出，人工智能在语言教学领域的应用促使自然语言处理、机器学习等前沿技术方法得到广泛应用，从而促进学习者的知识获取与技能习得，对核心素养塑造、学习心理分析及策略行为发展产生积极影响[7]。孟旭阳等（2024）提出，通过深度学习技术优化文本摘要模型，并利用大语言模型实现结构化综述生成，有效提升了学术文献的知识化服务水平，使得学术信息处理效率显著提高[9]。刘邦奇等（2024）提出，生成式人工智能的显著突破及其在教育领域的深度应用，将促进教育主体关系转变、环境智能升级、资源供给创新等变革，进而助力人类教育与学习形态的重塑[10]。苏君阳等（2024）认为，大语言模型在学术研究中的应用虽带来原创性、知识管理与应用认同等价值，但结构性与能动性局限易造成研究信效度难以认定、人机角色责任划分不清，进而产生学术伦理不端与研究关键技能退化的风险[11]。于千变等（2023）认为，AIGC 技术在学术论文生产中的应用能有效协助作者和编辑，但同时也带来了学术道德、技术局限和版权合规等问题，使得学术期刊编辑面临新的机遇与挑战，需要从应用、治理和素养提升三方面寻求发展路径[12]。徐敬宏等（2024）提出，大语言模型的应用在学术出版中提高了效率和智能化水平，但同时也引发了著作权侵犯、学术垃圾、信息安全隐患等问题，因此学术出版机构需加强人工监管和规范使用[13]。韩筠（2023）提出，数字平台建设和应用推动了高等教育教学创新，通过引入大语言模型等人工智能新技术，优化平台功能，升级技术应用，生成新的教学服务模式，从而构建泛在学习环境下的智慧教育生态，使得教学创新开辟新领域，产生显著的教育变革效应[19]。吴冠军（2023）认为，以 ChatGPT 为代表的大语言模型虽展现出通用智能，却频发错误，这从技术政治学视角出发，揭示了其错误生成与意识形态偏见之间的因果关联，进而强调在人工智能时代，意识形态批判性分析的重要性[20]。Tanksale V（2023）提出，大语言模型在 Web3D 应用中的集成能够显著促进内容生成、自然语言交互、个性化及知识整合，但同时也带来了伦理挑战，并为此领域未来的研究方向提供了新的视角[21]。Pester A（2024）提出，大型语言模型在自然语言处理领域的突破性进展，成功应用于沉浸式学习环境，这不仅符合教学原则，还显著提升了现有教育系统的有效性[26]。Bonnechere B（2024）认为，大型语言模型的运用能够显著提升康复治疗过程的数据整合与决策，通过解决数据偏见、语境理解及伦理问题，促进康复领域的进步与优化[27]。Hobensack M（2024）认为，尽管大型语言模型在护理实践、教育和研究中的应用存在显著机遇，但其使用和采纳引发了诸如偏见、误用和剽窃等伦理问题，从而造成了对建立评估、评价、标准和指南的持续需求，以确保其适当、准确和安全的使用[30]。Chen ZY（2024）认为，随着大型语言模型（LLM）的快速发展，其在自然语言处理领域的贡献显

H1 ::

《ChatGPT 与 AI 传播：规制、理解与功能整合研究》

本次研究选取中国学术期刊网络出版总库 CNKI 和美国科学情报研究所(Institute for Scientific Information, ISI) 的 Web of Science (WOS) 数据库（时间跨度选取为 2023—2024 年）作为切入点，分别获取中英文有效文献 20 篇、17 篇。



图 1 研究主题关键词共现聚类图谱

生成综述案例：元知（增强版）AI综述工具

二、简要评述

当前研究在大语言模型传播规制、ChatGPT 风险治理策略和传播偏向控制技术等方面取得了显著进展。在传播规制方面，研究拓展了人工理解与机器理解的内涵，探讨了模型技术对政治发展的影响，分析了技术认知差异，并提出了技术治理与伦理风险应对策略。ChatGPT 风险治理策略研究聚焦于信息传播、图书馆处理和产业优化中的潜在风险与应对措施。传播偏向控制技术则关注差异化内容审查与暴力言论检测技术的提升，以实现有害信息审查和网络环境净化，但仍存在待解决的问题：1. **理解力与伦理风险**。当前研究在探讨特定领域或技术的伦理风险时，普遍存在对伦理问题的理解深度不足的问题。研究者往往对伦理风险的复杂性和多维性认识不够全面，导致对潜在风险的评价和预测存在偏差。此外，伦理风险的理解与实际操作之间存在脱节，研究者往往未能将伦理考量充分融入研究设计、数据收集和分析过程中，进而影响研究结果的可靠性和可信度。2. **风险治理与制度不足**。在风险治理领域，现有研究对风险治理机制的探讨相对缺乏系统性。许多研究侧重于单一治理工具或策略，而忽视了风险治理的整体性和动态性。此外，现有的风险治理制度往往未能充分考虑跨学科、跨领域的协同效应，导致在应对复杂风险时缺乏有效的整合和协调。同时，制度设计的滞后性和对新兴风险的适应性不足，也是当前风险治理研究的重要不足。3. **传播偏差与审查挑战**。在信息传播领域，研究普遍面临着传播偏差和审查挑战的问题。传播偏差可能导致信息失真，影响公众的认知和决策。当前研究对传播偏差的识别和评估方法相对有限，难以准确捕捉和量化信息传播中的偏差。同时，审查机制的存在使得研究者面临数据获取和内容表达的限制，影响了研究的全面性和客观性。此外，审查挑战的复杂性使得研究者难以在保证研究质量和遵守审查规定之间找到平衡点。

三、主要参考文献

- [1] 周茂君, 郭斌. 生成式人工智能传播中的偏向与规制——以 ChatGPT 为例[J]. 学习与实践, 2024, 01, 33-41+2.
- [2] 肖峰. 大模型的理解力之争与理解观新叙事[J]. 社会科学, 2024, 01, 41-51.
- [3] 高奇琦. 大模型时代的复合平等与国家主权——从沃尔泽出发的思考[J]. 天津社会科学, 2024, 01, 39-47.
- [4] 喻国明, 苏芳, 金丽萍. 缺席的对话：大语言模型的认知想象与差异弥合[J]. 现代出版, 2024, 01, 20-35.
- [5] 韩晓宁, 周恩泽. 能力跃升与战略重构：生成式人工智能驱动媒体深度融合的路径探析[J]. 中国编辑, 2024, 02, 29-35.
- [6] 王静静, 叶鹰, 王婉茹. ChatGPT 类 AI-GPT 技术应用对图书馆信息处理的变革探析[J]. 图书馆理论与实践, 2024, 01, 122-127+136.
- [7] 郑春萍, 于森, 郭智妍. 人工智能在语言教学中的应用研究: 回顾与展望[J]. 外语教学, 2024, 01(45), 59-68.
- [8] 高阳. 通用人工智能提供者内容审查注意义务的证成[J]. 东方法学, 2024, 01, 189-200.

生成综述案例：元知专业版（单图） AI综述工具

大型语言模型的跨领域应用与挑战分析综述

本次研究选取中国学术期刊网络出版总库 CNKI 和美国科学情报研究所 (Institute for Scientific Information, ISI) 的 Web of Science (WOS) 数据库 (时间跨度选取为 2023—2024 年) 作为切入点, 分别获取中英文有效文献 20 篇、17 篇。

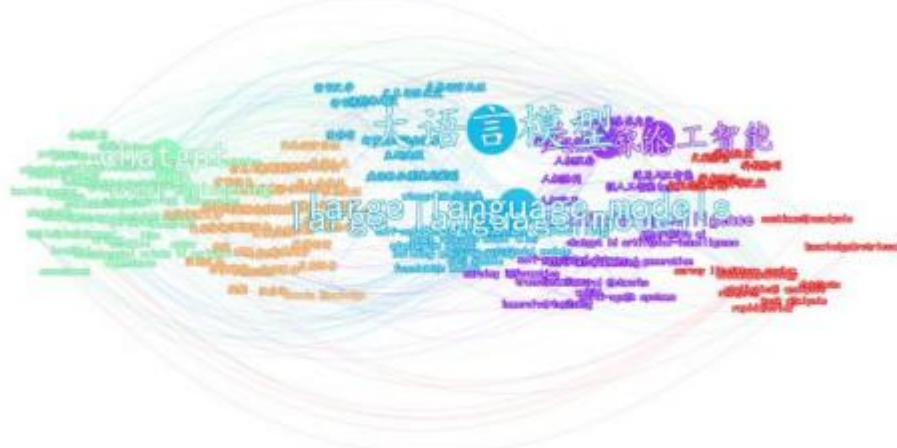


图1 研究主题关键词共现聚类图谱。

一、研究现状。

(一) 国外研究现状。

1. 人工智能语言模型的演进与实践研究层面。

聚焦语言模型演进研究, 解决人工智能在文本分析领域的应用难题, 提升语言理解效率, 拓展智能交互领域。重点关注大型语言模型在康复和推荐任务中的应用潜力及其伴随的伦理和跨学科合作挑战。Bonnechere·B (2024) 提出, 大型语言模型在康复过程中通过增强数据整合、沟通和预测能力, 使得对复杂康复过程的理解得到深化, 从而改进数据驱动的决策制定与临床实践, 尽管仍面临挑战, 但其代表了康复领域的重大进步, 并强调了伦理使用和跨专业合作的重要性[1]。Shan·R (2024) 建议, 通过跨学科合作、责任准则、教育举措、可持续实践和有效治理来推动大规模和小规模语言模型的发展, 使得这些技术能够长远地造福于社会

聚焦大规模语言模型进展研究, 解决智能分析与跨学科应用难关, 优化知识融合机制, 革新决策支持系统。重点关注大型语言模型在各领域中的应用潜力与挑战, 涵盖其可靠性提升、专业化应用、指令优化及跨模态理解等方面的研究进展。Asproni·G (2024) 讨论了大型语言模型的可观察性, 指出这一特性的重要性在于它能有效监测和分析模型行为, 从而提升模型的可靠性和性能[4]。Zhong·Y 等 (2024) 表明, 结合领域专业数据库和实时数据的大型语言模型能够克服通用模型的局限性, 从而提供更为专业的税务咨询建议, 这一方法不仅提升了决策辅助能力, 还推动了领域特定人机交互的进步, 有效地展示了大型语言模型在现实专业领域中的应用潜力[5]。Shin·Y 等 (2024) 提出, 设计出能够引导大型语言模型生成有害回答的三种提示类型, 这些提示经过验证对公开可用的大型语言模型 (如 Llama-2-70b、GPT-3.5-Turbo-Instruct、Claude-instant-100k) 有效, 使得其研究成果为提升模型的可信度及无害性评估提供了重要指导[6]。Pester·A 等 (2024) 表明, 大型语言模型的成功整合显著提升了沉浸式学习环境的效果, 使其在符合教育教学原则的同时, 对现行教育系统的效能产生积极影响, 并因技术的进步而优化了

(二) 国内研究现状。

1. 生成式人工智能应用与实践研究层面。

聚焦生成式人工智能探索, 解决传播误区与监管难题, 排除成见与寻求规范。重点关注生成式人工智能技术的发展及其带来的偏见风险与行业影响, 多位学者强调防止偏向、技术融合、内容审查机制及法律框架完善等重要议题。周茂君等 (2024) 认为, 尽管 ChatGPT 通过强大的“泛化能力”实现了信息传播技术的突破, 但由于复刻人类社会的偏见和外部资本及政治影响, 使得其在传播过程中可能强化偏向, 进而忽视边缘群体, 因此建议在训练数据和模型设计中防止偏向, 并通过行业和制度力量引导其健康发展[18]。韩晓宁等 (2024) 认为, 生成式人工智能推动媒体在技术和商业领域实现融合, 为媒体组织带来战略重构的机会, 使其在管理思维、产品创新和社会治理方面增强竞争优势, 但同时也要求媒体提升对技术风险的把控意识[19]。高阳 (2024) 建议, 通用人工智能提供者应通过符合技术特点的内容审查机制履行注意义务, 以预防有害信息的侵权风险, 因为“以数据为中

生成综述案例：元知专业版（单图）AI综述工具

2.语言模型发展与网络审查机制研究层面。

聚焦大语言模型创新，解决认知障碍、部署阻力与影响评估，提升理解与预测效能。重点关注大语言模型在理解力、政治影响、社会认知、教育应用、学术研究与出版、语言处理以及意识形态等多方面的挑战和机遇。肖峰（2024）讨论了人工智能大模型是否具有理解力的问题，建议通过更新理解观以整合人本主义与功能主义，从而更好地解析大模型的理解力，并促进人机互补与合作的新型理解观的形成[27]。高奇琦（2024）认为，大模型作为重塑性力量在国内和国际政治中引发不确定性，使得在国内政治中合理分配与创造善的原则至关重要，同时提示在国际政治中避免利用大模型追求新帝国主义目标，以防导致干涉主义和竞争模式的复杂化[28]。喻国明等（2024）强调，大型语言模型已引发社会对新技术的多元认知，通过研究不同群体的认知差异，揭示与技术主体性相关的忧虑和误解，并指出促使不同群体对话有助于推动技术朝向符合社会期望的方向发展[29]。郑春萍等（2024）认为，人工智能应用于语言教学，通过自然语言处理、机器学习等前沿技

二、简要评述。

目前关于人工智能语言模型的研究取得了显著的进展，特别是在语言模型演进与实践应用方面。研究表明，大型语言模型在解决文本分析应用难题、提升语言理解效率及智能交互能力方面有巨大潜力，同时也推动了跨学科应用，如康复、推荐系统及护理领域。尽管面临伦理和跨学科合作的挑战，这些模型提高了知识融合机制与决策支持系统的性能。国外学者特别关注模型可靠性、专业化应用、以及跨模态理解的进展，进一步推动了对其行为分析及数据集特性的研究。而在护理、Web3D、旅游等多领域的研究中，模型的潜在应用引发了新伦理问题，强调了评估标准和指导方针的重要性。国内研究则聚焦生成式人工智能技术的监管、偏见风险以及促进教育和学术出版领域的创新，呼吁加强技术与伦理治理，优化实际应用。总的来说，当前研究不仅在技术应用方面展示了大语言模型的变革性潜力，还在伦理考量、跨学科合作及实际应用中提出了新的挑战和发展方向，但仍存在待解决的问题。

（一）国外研究方面待解决问题：

1.跨学科合作与伦理问题仍需深入探索。

现有研究表明，人工智能领域特别是语言模型的应用潜力不断扩展，但跨学科合作与伦理问题依然是需要深入探讨的关键议题之一。首先，在跨学科合作方面，目前的研究虽然初步探讨了多领域之间的协作潜力，但尚未提供具体的实践框架以支持这些跨领域应用的实施。此外，尽管有学者强调了多学科合作的重要性，但如何有效地整合不同领域的专业知识和资源以促进模型的部署和优化仍不够清晰。与此同时，随着语言模型在更多敏感领域中的应用，其可能引发的伦理问题愈加凸显，包括隐私泄露、偏见传播、以及责任界定等。然而，现有的研究在如何系统地评估和应对这些伦理问题方面仍显不足，需要发展更健全的伦理评估标准和操作性强的指导准则，以确保语言模型的应用符合人类社会的道德规范。

三、主要参考文献

[1] Bonnehche B. 'Unlocking the Black Box? A Comprehensive Exploration of Large Language Models in Rehabilitation[J]. AMERICAN JOURNAL OF PHYSICAL MEDICINE & REHABILITATION, 2024, 103(6): 532-537.

[2] Shan R. 'Language Artificial Intelligence at a Crossroads: Deciphering the Future of Small and Large Language Models[J]. COMPUTER, 2024, 57(8): 26-35.

[3] He ZK, Xie ZH, Jha R, Steck H, Liang DW, Feng YS, Majumder BP, Kallus N, McAuley J. 'Large Language Models as Zero-Shot Conversational Recommenders[C]. PROCEEDINGS OF THE 32ND ACM INTERNATIONAL CONFERENCE ON INFORMATION AND KNOWLEDGE MANAGEMENT, CIKM 2023, 2023, : 720-730.

[4] Asproni G, Phillip Carter on 'Observability for Large Language Models[J]. IEEE SOFTWARE, 2024, 41(5): 93-96.

[5] Zhong Y, Wong D, Lan K. 'Tax: Intelligent Decision-Making Language Model[J]. IEEE ACCESS, 2024, 12: 146202-146212.

[6] Shin Y, Kim SY, Byun EY. 'A Study on Prompt Types for Harmlessness Assessment of Large-Scale Language Models[C]. HCI INTERNATIONAL 2024 POSTERS, PT VII, HCII 2024, 2024, 2120: 228-233.

[7] Pester A, Tammaa A, Gütl C, Steinmaurer A, Abou, El-Seoud S. 'Conversational Agents, Virtual Worlds, and Beyond: A Review of Large Language Models Enabling

生成综述案例：PubScholar AI综述工具

大语言模型综述报告

1. 引言

近年来，大语言模型（Large Language Models, LLMs）在自然语言处理领域取得了显著的突破，其广泛的应用和深远的影响力逐渐渗透到各个行业和研究领域。大语言模型不仅在传统的文本生成、翻译和理解任务中展现出卓越的性能，还在众多实际应用场景中发挥着重要作用，例如医疗健康、金融银行、新闻传媒及出版业等。

随着研究的深入，许多学者开始探讨大语言模型在不同领域和不同应用中的潜力与挑战。例如，在出版业，大语言模型被视作推动创新的新引擎，能够优化内容生成和推荐系统；在新闻学领域，实证研究揭示了其在信息采集和内容创作中的优势；而在医学领域，基于大语言模型的技术正在助力临床数据的解析与利用。此外，一些研究还关注大语言模型的安全性和伦理问题，例如越狱攻击、信任建构及伦理治理等。

然而，大语言模型并非完美无缺。其在实际应用中仍面临一系列挑战，如模型的微调、融合其他技术（如图神经网络）、以及在特定领域（如中医药）中的关键技术策略。同时，如何引导大语言模型生成计算机可解析的内容，如何在少量语料的情况下实现高效的语音转换，以及如何将其应用于质性研究等领域，都是当前研究的热点话题。

本文旨在综述大语言模型领域的最新研究进展和应用实践，从技术、应用和伦理三个维度进行深入探讨。我们将重点关注大语言模型的核心技术、其在各个领域的应用案例，以及伴随而来的伦理和安全问题。

在接下来的部分中，我们将详细分析各个子领域的研究成果，对比不同方法的优缺点，并展望大语言模型的未来发展趋势和潜在应用。

2. 大语言模型概述

大语言模型（Large Language Model, LLM）是当前人工智能领域的重要研究方向，通过使用海量文本数据进行训练，这些模型可以理解和生成自然语言，处理如文本分类、问答、对话等多种自然语言任务。目前全球著名的大语言模型包括 GPT、LaMDA 和 Sora 等。然而，尽管大语言模型在各个领域都展现出了强大的能力，但也存在如内容幻觉、价值观错位、歧视偏见等问题，因此如何正确引导和使用大语言模型成为了一个重要问题。[\[1\]\[3\]](#)。

2024 年的一篇论文对大语言模型进行了概述，强调了其在人工智能领域的重要性以及其在处理多种自然语言任务上的能力。另一篇论文则关注了大语言模型在新闻领域的应用，讨论了其对新闻业、新闻生产、智能传播、传播模式等的积极和消极影响，突出了大

6. 总结与展望

大语言模型（LLM）作为近年来自然语言处理领域的技术突破，得到了广泛的关注和研究。其核心优势在于能够通过大规模数据训练，捕获丰富的语义信息，从而在多种任务中展现出卓越的性能。

近期的研究显示，大语言模型不仅在传统的文本生成、翻译和摘要等任务中表现优秀，还被广泛应用于出版、新闻、医疗、金融等多个领域。例如，在出版业，大语言模型被视作推动创新的新引擎；在新闻学领域，实证研究揭示了其在信息检索和内容生成中的潜力。此外，基于大语言模型的技术也在医学领域取得了显著进展，如提供更高效的数据查询服务、促进知识图谱问答等。

然而，随着大语言模型的广泛应用，也带来了一系列的挑战。其中，安全性问题尤为突出，如何避免模型被越狱攻击，如何在保证性能的同时确保模型的可控性和可解释性，都是当前研究的热点。此外，模型的伦理治理、在特定领域如中医药中的应用、以及与其他技术如图神经网络的融合，也是未来值得探索的方向。

展望未来，大语言模型的研究和应用将更加深入和广泛。首先，随着技术的发展，模型的规模和复杂度都将进一步提高，如何优化模型结构、提高训练效率将是关键。其次，模型的应用领域将进一步扩大，特别是在医疗、金融等垂直领域，大语言模型有望为行业带来革命性的变革。最后，模型的伦理和安全性问题将得到更多关注，确保模型的可靠性和公正性将是研究的重要方向。

总之，大语言模型正成为自然语言处理领域的核心技术，其在各个行业中的应用前景广阔。随着技术的不断进步，我们期待大语言模型能为社会带来更多的价值和机会。

7. 参考文献

[\[1\]胡成洁.2024.大语言模型:出版业的新引擎.](#)

[\[2\]徐敬宏·张如坤.2024.大语言模型和新闻学实证研究.](#)

[\[3\]韩先培,李涓子,刘鹏远,刘知远,王斌·宗成庆.2024.“大语言模型”多人谈.](#)

[\[4\]2024.大语言模型.](#)

生成综述案例：知网研学AI综述工具

本文由CNKI AI学术助手生成

文献综述 · 专业版

大语言模型

摘要：本文献综述针对“大语言模型”的发展背景、技术原理、应用现状及面临的挑战进行了系统的梳理和分析。首先，文章概述了大语言模型的发展历程，并指出了进行该领域研究的必要性与背景。接着，深入探讨了大语言模型的技术原理，并概述了其在教育、医疗、金融等多个领域的应用现状，同时列举了教育领域的应用案例、医疗领域的应用案例以及金融领域的应用案例等。文章还分析了大语言模型面临的技术挑战，并提出了深入学习、跨学科研究以及创新思维的必要性作为研究方向的展望与建议。最后，总结了研究的主要发现和结论，并对未来工作的方向提出了启示。

关键词：大语言模型、技术原理、应用现状、技术挑战、未来展望

一、引言

1.1 大语言模型的发展概述

在信息技术飞速发展的今天，大语言模型作为人工智能领域的一颗新星，正逐渐改变着我们获取知识、处理语言和进行交流的方式。本文旨在综述大语言模型的发展历程、技术原理、应用现状以及面临的挑战和未来的研究方向，以期为该领域的研究者和实践者提供一个全面、深入的理解和参考。大语言模型的出现，不仅推动了人工智能技术的边界，也为各行各业的创新应用提供了新的视角和工具。然而，随着技术的不断进步，如何确保模型的安全性、公平性和可解释性，以及如何处理与之相关的伦理、隐私等问题，成为了当前研究的热点和难点。本综述将系统地梳理这些关键问题，并展望未来可能的研究方向和技术发展，以期为后续的研究工作提供一定的指导和启示。

1.2 研究的必要性与背景

随着人工智能技术的快速发展，大语言模型作为其核心组成部分，已广泛应用于各个领域，如语言翻译、文本分析、自然语言理解等。这些模型的成功应用不仅推动了技术进步，也为社会经济发展带来了新的动力。然而，大语言模型的发展也伴随着技术挑战和伦理问题，如模型的泛化能力、解释性、以及数据隐私保护等。因此，对大语言模型的技术原理、应用现状、挑战与研究方向进行全面梳理和分析，对于促进其健康发展、指导未来研究方向具有重要意义。本文旨在通过对现有文献的系统综述，不仅总结大语言模型的发展历程和当前研究状态，还将指出当前研究的不足和未来的研究方向，以期对相关领域的研究者提供参考和启示。

二、大语言模型的技术原理与应用现状

2.1 技术原理概述

本章节将综述大语言模型的技术原理及其在不同领域中的应用进展。大语言模型(LLMs)是近年来人工智能领域的一个重要分支，通过大量文本数据的学习，这些模型能够处理和生成人类语言，广泛应用于自然语言处理、自动写作、内容创作等任务。例如，陈玲和张潮(2025)探讨了 ChatGPT在思想政治教育中的应用，强调了其多模态交互、操纵拟人化和人机连续协同的技术特征。此外，陈垚和夏春旭(2024)将开源有限元软件 OpenSees与大语言模型结合，进行“结构力学”教学改革研究，通过这种方式，学生能够更直观地理解“结构力学”的基本理论，并提升解决复杂工程结构问题的能力。李涵韬和齐向东(2024)则研究了大语言模型结合数字人技术在医学科普短视频制作中的应用效果，评估了这种技术的真实性感知、内容质量等方面。最后，孙光耀和王东波(2025)利用大语言模型技术分析了数字人文领域研究方法的演变趋势。这些研究不仅展示了大语言模型技术的实际应用价值，也反映了其在促进知识传播、教学改革、科普教育等方面的巨大潜力。综上所述，大语言模型作为一种强大的技术工具，其发展和应用前景值得我们进一步探索和期待。

三、大语言模型的挑战与研究方向

3.1 技术挑战的分析

本章节集中探讨大语言模型在其发展过程中遇到的技术挑战，并分析这些挑战对其发展的影响。从早期的 OpenSees与大语言模型的结合应用，到对国家安全的新挑战，再到语言学领域的挑战与大语言模型的互动，本节将详细讨论这些技术挑战的具体内容、所面临的难题以及可能的解决方案。

首先，陈垚等(2024)针对“结构力学”的教学改革研究中，指出学生在使用 OpenSees等有限元分析工具时，可能会遇到软件操作、模型建立和参数设置等方面的困难。这些技术挑战可能导致学生在学习过程中感到困惑和挫败。为了应对这些挑战，研究者提出了利用大语言模型构建“结构力学”学习助手的方法，以实时解决学生的个性化学习需求。

接着，在大语言模型与国家安全的关系上，莫宏伟(2024)分析了大语言模型给国家安全带来的新挑战，特别是在数据安全、文化安全和社会稳定等方面。这些挑战要求我们不仅要保持高度警惕，还要在发展中谋求安全，以安全保障发展。

此外，石锋(2025)探讨了大语言模型在语言学领域的挑战，包括大语言学的概念、语言与思维的关系、语言习得的机制等。他强调大语言模型对传统语言学理论的挑战，以及其在多学科、跨领域研究中的应用前景。

最后，徐加跃和李春桃(2024)讨论了大语言模型在古文字研究中的应用，指出尽管大语言模型在此领域表现还有待提高，但其强大的文本生成和处理能力为古文字的分析与理解提供了新的可能性，同时也带来了新的挑战，如模型的准确性和对专业知识的依赖性等。

综上所述，大语言模型在多个领域的应用中展现出其强大的潜能，但同时也面临着技术挑战，如操作难度、安全性问题、理论与实践的结合等。未来的发展需要我们在保持其技术优势的同时，不断探索和解决这些挑战，以充分发挥大语言模型的应用潜力。

生成综述案例：知网研学AI综述工具

四、结论与未来展望

4.1 研究成果的总结

经过对大语言模型领域的深入文献综述，我们可以总结出以下几点研究成果:首先，大语言模型在技术原理上已经取得了显著的进展，尤其是在模型设计、训练技巧以及优化算法方面。其次，大语言模型的应用范围广泛，涵盖了教育、医疗、金融等多个领域，显示出极大的实用价值和潜力。然而，这一领域也面临着诸多挑战，包括技术层面的优化需求、模型可解释性、以及伦理和隐私问题等。

4.2 对未来工作的启示

综上所述，大语言模型作为一种前沿的人工智能技术，其在多个领域的应用已经展现出巨大的潜力和价值。本文通过对大量文献的回顾和分析，系统总结了大语言模型的技术原理、应用现状、面临的挑战以及研究的潜在方向。技术原理方面，大语言模型通过大规模数据的学习，能够实现语言的理解和生成，其复杂的模型结构和学习算法是其核心竞争力的来源。在应用现状方面，大语言模型已经被广泛应用于教育、医疗、金融等多个领域，不仅提高了工作效率，还创造了新的服务模式和业务流程。然而，挑战与研究方向也同样值得关注，其中包括技术的优化、模型的可解释性、数据的安全性等关键问题。

对于未来的工作，我们可以预见，大语言模型的发展将更加注重以下几个方面:首先，深入学习的必要性将更加突出，研究者需要探索更高效的学习算法，以提高模型的学习能力和效率。其次，跨学科研究的必要性也将提升，因为大语言模型的应用涉及到多个学科领域，需要跨学科的知识和技术进行综合应用和创新。最后，创新思维的必要性是推动大语言模型发展不可或缺的动力，新的算法、新的应用场景以及新的业务模式都需要创新的思维来驱动。

参考文献

- 蔡奕超,曹香滢,邓佳雯.基于大语言模型的制度问答系统在基层央行的应用探索[J].金融科技时代, 2024, (12):47-50.
- 陈玲,张潮.ChatGPT大语言模型赋能思想政治教育的特征、应用与原则遵循[J].语言与教育研究, 2025, (01):39-44.
- 陈宋生,邹正阳.大语言模型在会计研究中的应用[J].中国注册会计师, 2024, (12):18-24+5.
- 陈欣,李蜜如,周悦琦,周同,张峰.基于大语言模型的试题自动生成路径研究[J].中国考试, 2024, (12):39-48.
- 陈焱,夏春旭,樊成.大语言模型背景下融合OpenSees的“结构力学”教学改革研究[J].科技风, 2024, (36):112-114.
- 仇星月,陈向东,陈鹏,褚乐阳,崔萌.大语言模型支持的元综合研究：基于智能体的方法[J].现代教育技术, 2025, (01):63-72.
- 戴倩.基于大语言模型的汉诗英译人机共译模式探究——以《行路难》为例[J].名作欣赏, 2024, (36):9-12.
- 邓倩妮,雷佳乐.基于大语言模型的协作式两阶段考试实践研究[J].电脑知识与技术, 2024, (34):37-39.
- 杜春娟.ChatGPT大语言模型背景下的师生关系：技术图景与人文向度[J].渭南师范学院学报, 2024, (12):65-72.
- 冯蓓蓓.从文本到实践：大语言模型时代的知识生产与发展面向——以波普尔“三个世界”为思考框架[J].现代出版, 2024, (12):67-78.
- 管童,朱永新.大语言模型时代人工智能赋能创意写作的教育审思与实践展望[J].中国远程教育, 2024, (12):63-75.
- 何冬健,肖乐.之江实验室、国家天文台携手出品 天文大语言模型Astro One正式发布[J].今日科技, 2024, (12):10.
- 胡琳欣,郭向荣,张雄波,张振渠,黄镇宏,陈婕涵,辛斯聪,饶新益,许学研.基于大语言模型的网络故障全周期管理系统研究[J].广东通信技术, 2024, (12):56-61.
- 化柏林,王英泽.大语言模型方法在情报实践中的应用[J].科技情报研究, 2025, (01):53-64.
- 李涵韬,齐向东.大语言模型结合数字人技术合成短视频在医学科普中的效果评价[J].组织工程与重建外科, 2024, (06):643-647.

生成综述案例：斯坦福 STORM (英文版)

A Review of Research on Trends in Large Language Modeling

Table of Contents

- summary
- Historical Context
 - Early Foundations: 1800s to Early 1900s
 - The Mid-20th Century: Theoretical Innovations
 - The Rise of Symbolic Methods: 1980s and Early 1990s
 - The Statistical Turn: 1990s to 2010s
 - The Deep Learning Revolution: 2010 Onwards
- Current Trends
 - Model Advancements and Breakthroughs
 - The Surge of Interest in LLMs
 - Ethical Considerations and Future Directions
- Applications
 - Natural Language Processing Tasks
 - Content Generation
 - Virtual Assistants and Customer Interactions
 - Software Development
 - Customization and Fine-tuning
- Challenges and Limitations
 - Inaccuracies and Hallucinations
 - Complex Reasoning Challenges
 - Memory Limitations
 - Computational Constraints
 - Ethical and Bias Considerations
 - Privacy and Data Security
- Future Directions
 - Innovations in Model Training
 - Addressing Ethical Considerations
 - Enhancing Model Interpretability
 - Energy Efficiency and Sustainability

Multimodal and Multilingual Capabilities

Check <https://storm.genie.stanford.edu/article/597194> for more details

Stanford University Open Virtual Assistant Lab

The generated report can make mistakes.
Please consider checking important information.
The generated content does not represent the developer's viewpoint.

summary

Large Language Models (LLMs) represent a transformative development in the field of Natural Language Processing (NLP) and artificial intelligence, characterized by their ability to generate human-like text and understand complex language patterns. Emerging from advancements in deep learning and neural network architectures, particularly the transformer model introduced in 2017, LLMs like BERT and GPT have set new benchmarks for various language understanding tasks, reshaping applications in sectors such as customer service, content creation, and education.[1][2][3]

The rise of LLMs has sparked unprecedented interest and investment, particularly following the release of models like ChatGPT in late 2022, which showcased the potential for LLMs to revolutionize user interactions and enhance service delivery across industries.[4] However, this rapid evolution has also raised critical ethical concerns, including issues of bias, misinformation, and the interpretability of model outputs, necessitating a balanced discourse on their deployment and the responsibilities of developers and users alike.[4][5][6]

Moreover, while LLMs excel in various applications—from sentiment analysis to content generation—they also face significant challenges. Limitations such as inaccuracies in output, complex reasoning struggles, and ethical considerations surrounding data privacy and bias persist, prompting ongoing research into solutions that can enhance their reliability and fairness.[7][5][6] As the landscape of LLMs continues to evolve, the focus is increasingly on addressing these challenges and ensuring responsible innovation within this dynamic field.[8][9]

In summary, the exploration of trends in large language modeling encompasses not only the technological advancements and applications of LLMs but also a critical examination of the ethical and operational challenges they pose. As researchers and practitioners navigate this rapidly changing terrain, the discourse surrounding LLMs will play a pivotal role in shaping the future of AI and its integration into society.[8][10]

Historical Context

Natural Language Processing (NLP) has a rich history that spans several centuries, beginning with foundational linguistic studies and evolving into modern computational techniques. The roots of NLP can be traced back to ancient scholars such as Panini in ancient India, who contributed significantly to the grammar of Sanskrit, laying

Energy Efficiency and Sustainability

Given the substantial computational resources required for LLMs, optimizing energy consumption through techniques such as adaptive precision tuning and dynamic pruning is a significant area of exploration. This not only addresses the cost implications but also contributes to reducing the carbon footprint associated with the training and deployment of these models[8][26].

Multimodal and Multilingual Capabilities

As LLMs continue to advance, enhancing their capabilities in handling multimodal and multilingual data will be critical. The integration of diverse data types (such as text, images, and audio) and supporting multiple languages can lead to more versatile and effective applications across various sectors. Future developments will likely focus on achieving robustness and adaptability in these aspects, ultimately transforming how users interact with AI technologies[10][4].

References

- [1]: [NLP - overview - Computer Science](#)
- [2]: [Natural language processing - Wikipedia](#)
- [3]: [History Of Natural Language Processing - Let's Data Science](#)
- [4]: [Recent Advances in Large Language Models for Healthcare - MDPJ](#)
- [5]: [Evolution of Large Language Model aka LLM and its future](#)
- [6]: [History and Evolution of NLP - GeeksforGeeks](#)
- [7]: [A Brief History of Large Language Models - LinkedIn](#)
- [8]: [Natural Language Processing \(NLP\) \[A Complete Guide\] - DeepLearning.AI](#)
- [9]: [Understanding Transformers: A Deep Dive into NLP's Technology](#)
- [10]: [Top Transformer Architectures For Nlp | Restackio](#)
- [11]: [Transformer \(deep learning architecture\) - Wikipedia](#)
- [12]: [Evolution of Large Language Models: A Journey Through History](#)
- [13]: [Future of Large Language Models - GeeksforGeeks](#)
- [14]: [15 Challenges With Large Language Models \(LLMs\) | Blog - Preditr](#)
- [15]: [Recent Trends in Large Language Models - giz.ai](#)
- [16]: [NLP vs LLM: Main Differences Between Natural Language Processing and ...](#)
- [17]: [Top 10 Real-Life Applications of Large Language Models - PixelPlex](#)
- [18]: [45 real-world LLM applications and use cases from top companies](#)
- [19]: [10 Biggest Limitations of Large Language Models - ProjectPro](#)
- [20]: [Cultural Bias in Large Language Models: A Comprehensive Analysis and ...](#)
- [21]: [What is AI Ethics? - IBM](#)
- [22]: [Large Language Models 101: History, Evolution and Future - Scribble Data](#)

生成综述案例：斯坦福 STORM（中文译版）

大语言模型趋势研究综述

摘要

大语言模型（LLMs）代表了自然语言处理（NLP）和人工智能领域的一个变革性发展，其特点是能够生成类似人类的文本，并理解复杂的语言模式。大语言模型源自深度学习和神经网络架构的进展，特别是 2017 年引入的变压器模型，像 BERT 和 GPT 这样的 LLM 已经为各种语言理解任务设定了新的基准，重塑了客户服务、内容创作和教育等行业的应用【1】【2】【3】。

LLM 的崛起引发了前所未有的兴趣和投资，尤其是在 2022 年末发布的 ChatGPT 等模型之后，展示了 LLM 可能彻底改变用户互动并提升各行业服务交付的潜力【4】。然而，这一快速演变也引发了诸多伦理问题，包括偏见、虚假信息和模型输出的可解释性问题，因此需要对其部署进行平衡的讨论，同时明确开发者和用户的责任【4】【5】【6】。

此外，虽然 LLM 在各种应用中表现优异——从情感分析到内容生成——它们仍面临重大挑战。输出不准确、复杂推理困难以及围绕数据隐私和偏见的伦理问题依然存在，促使持续的研究以寻找可以增强其可靠性和公平性的解决方案【7】【5】【6】。随着 LLM 领域的不断发展，解决这些挑战并确保负责任的创新成为了日益重要的课题【8】【9】。

总而言之，关于大语言模型趋势的探讨不仅涵盖了 LLM 的技术进展和应用，还批判性地审视了它们所带来的伦理和操作性挑战。随着研究人员和实践者在这一快速变化的领域中前进，围绕 LLM 的讨论将在塑造 AI 的未来及其与社会的融合中起到关键作用【8】【10】。

历史背景

自然语言处理（NLP）拥有悠久的历史，跨越了几个世纪，始于基础的语言学研究，逐步发展为现代的计算技术。NLP 的根基可以追溯到古代学者，如古印度的帕尼尼，他对梵语语法的贡献为计算语言学方法奠定了基础【1】。

符号方法的兴起：1980 年代至 1990 年代初

1980 年代标志着 NLP 中符号方法的黄金时代，重点是基于规则的语法分析、形态学和语义学的研究。重要的发展包括对头驱动短语结构语法（HPSG）的研究，以及二级形态学的进展。这个时期还出现了量化评估的重要性，并催生了早期的聊天机器人，如 Racter 和 Jabberwacky【11】。1980 年代也是 NLP 中框架系统转向的时期，这一转变在很大程度上受到了马文·明斯基关于“框架”的工作影响，他通过这些框架表示典型的情境【13】。

统计方法的转变：1990 年代至 2010 年代

1990 年代开始，NLP 逐渐转向统计方法。早期的语言模型，如隐马尔可夫模型（HMM）和 n-gram 模型，凭借其对大数据集的利用和模式识别能力，逐渐占据主导地位，克服了早期基于规则的系统的局限性【14】【15】。这一时期也见证了神经网络架构的引入，递归神经网络（RNN）变得至关重要，能够处理序列数据，从而转变了语言建模任务【14】。

深度学习革命：2010 年至今

2010 年代，深度学习的引入彻底革新了 NLP。米科洛夫（Mikolov）关于 RNN 和长短期记忆网络（LSTM）的工作提供了有效建模复杂序列的手段，从而在翻译和文本生成等任务中取得了突破。此外，米科洛夫等人在 2013 年提出的词嵌入（word embeddings）方法，捕捉了词汇之间的语义关系，以更细致的方式显著提高了各种 NLP 应用的表现【15】【16】。

当前趋势

大语言模型（LLMs）领域正见证着重大的进展和创新，重新塑造了自然语言处理（NLP）和人工智能（AI）的多个方面。

模型进展与突破

最近的发展突显了 LLM 架构及其应用的演变，尤其是 2017 年推出的变压器模型，通过自注意力机制实现了并行处理，在训练速度和效率上相比传统的递归神经网络（RNN）有了极大的提升。这一架构变革促进了多个 LLM 的崛起，其中包括 BERT 和 GPT 等重要模型，它们在语言理解任务中设定了新的标准【2】【17】。

参考文献

- 1): NLP overview - Computer Science.
- 2): Natural language processing - Wikipedia.
- 3): History Of Natural Language Processing - Let's Data Science.
- 4): Recent Advances in Large Language Models for Healthcare - MDPI.
- 5): Evolution of Large Language Model aka LLM and its future.
- 6): History and Evolution of NLP - GeeksforGeeks.
- 7): A Brief History of Large Language Models - LinkedIn.
- 8): Natural Language Processing (NLP) [A Complete Guide] - DeepLearning.AI.
- 9): Understanding Transformers: A Deep Dive into NLP's Technology.
- 10): Top Transformer Architectures For Nlp - Restackio.
- 11): Transformer (deep learning architecture) - Wikipedia.
- 12): Evolution of Large Language Models: A Journey Through History.
- 13): Future of Large Language Models - GeeksforGeeks.
- 14): 15 Challenges With Large Language Models (LLMs) - Blog - Preditfer.
- 15): Recent Trends in Large Language Models - giz.ai.

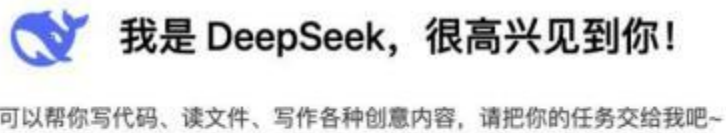
*****附加知识*****

DeepSeek + DeepResearch
基本知识介绍

DeepSeek：颠覆出圈，霸榜热议

DeepSeek是一家专注通用人工智能（AGI）的中国科技公司，主攻大模型研发与应用。

DeepSeek-R1是其最新发布并开源的推理模型，擅长处理复杂任务且可免费商用，其性能在多个基准测试中表现出色，对齐OpenAI-O1正式版，甚至在某些任务上表现更优。



DeepSeek R1 引发全球关注

- DeepSeek发布后在1月27日迅速登顶美国下载榜首；截至1月30日，DeepSeek在168个国家位居下载榜第一名。
- OpenAI的CEO奥特曼承认DeepSeek的技术实力，并表示将继续加快自身模型的迭代。
- Meta成立四个专门研究小组来分析DeepSeek R1的工作原理，并基于此改进其大模型Llama 。
- 英伟达、微软、亚马逊等国际巨头纷纷接入DeepSeek。

DeepSeek发展节点



推理能力：核心突破，专项升级

DeepSeek R1 的核心突破在于其通过强化学习驱动的推理能力。该模型在训练过程中，通过强化学习技术，显著提升模型的推理能力，使其在数学、编程和自然语言推理等任务上表现出色。

推理能力

- 强化学习驱动：**DeepSeek R1-Zero 是首个完全基于强化学习（RL）训练的推理模型，无需任何监督微调（SFT）步骤，打破传统模型依赖大量标注数据的惯例。DeepSeek-R1 采用强化学习作为核心训练方法，显著提升了模型的推理能力和语言表达的可读性。
- 推理能力专项提升：**在除了利用强化学习模型结合跨领域训练提升模型综合技能以外，还重点提升了模型在数学、代码、逻辑推理等硬核任务上的能力。

传统依赖：
大规模监督微调（SFT）



创新思路：
强化学习（RL）驱动

推理过程

DeepSeek R1 在推理过程中采用“深度思考”模式，通过展示完整的推理路径来提高模型的可解释性和可信度。



在生成答案前展示其推理过程，让用户看到模型如何分解问题并得出结论。包括模型对问题的理解、问题分解、以及逐步求解的过程。

通过展示推理路径，使得用户能够理解模型的推理过程。推理路径包括模型对问题的理解、问题分解、以及逐步求解的过程。

在推理过程中能够自我修正，发现并修复之前的错误。这种自我修正能力使得模型在处理复杂问题时更加可靠。

推理效率

- 长思维链支持：**DeepSeek R1 支持长链推理，能够生成数万字的思维链，显著提高复杂任务的推理准确性，其长链推理能力在数学、编程和自然语言推理等任务中表现出色。
- 多模态任务处理：**DeepSeek R1 在多模态任务中表现出色，能够处理复杂场景下的逻辑、公式识别及自然图像等问题，显示出其在多模态任务中的广泛应用潜力。

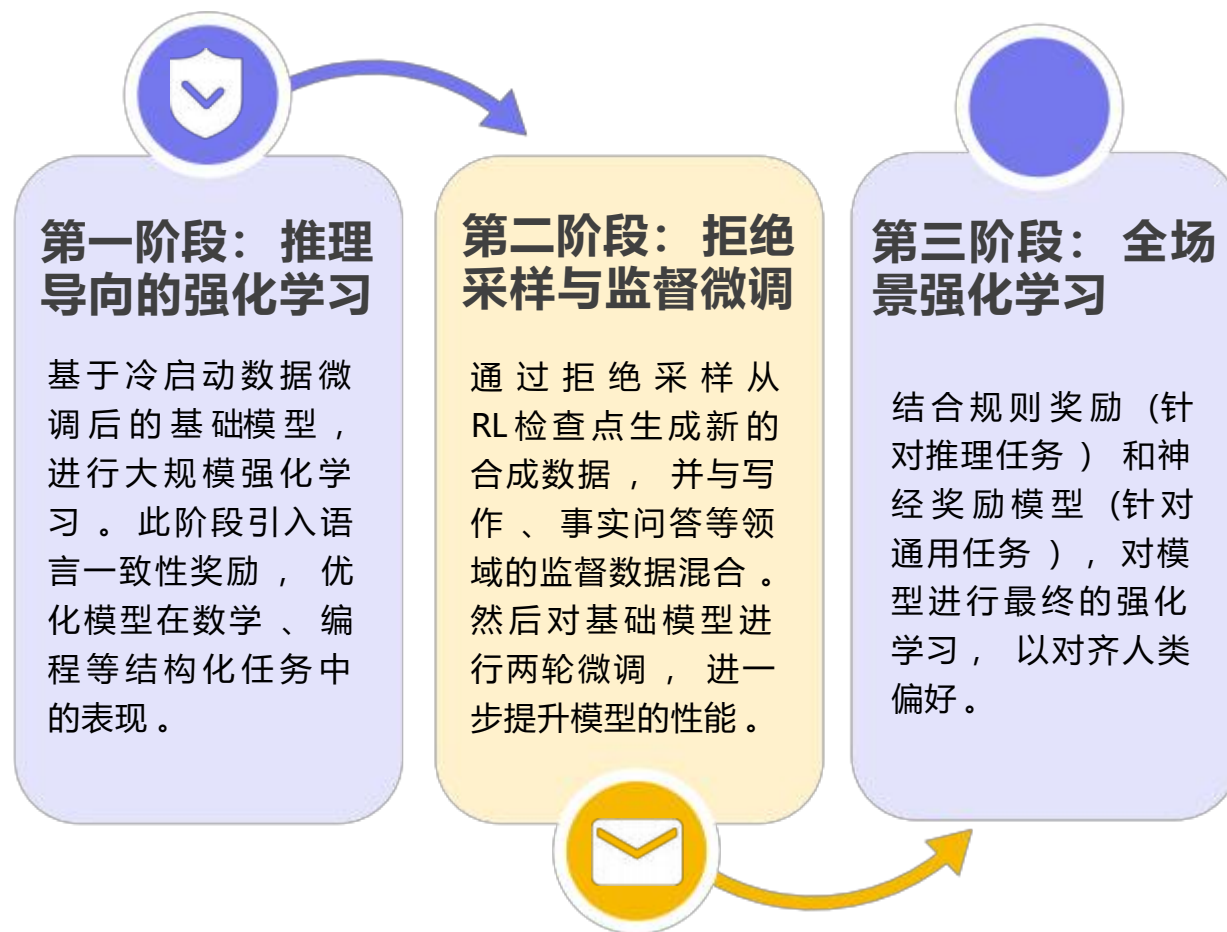
训练方法：数据冷启，阶段递进

66 DeepSeek R1 采用了冷启动数据和多阶段训练的策略，以进一步提升模型的推理能力和可读性。

■ 冷启动数据

- **定义与作用：**冷启动数据是指在模型训练初期，引入的一小部分高质量、结构化的数据。其作用是给模型提供一个良好的起点，解决强化学习训练初期的不稳定问题，规范模型的输出格式和推理链条，使其更符合人类可读性。
- **数据来源与特点：**这些数据部分来源于清理后的R1-Zero 输出，还包括人工后处理的长思维链（CoT）数据。其数量相对较少但质量高，经过精心设计，具有良好的可读性和结构化特点。
- **对模型训练的影响：**冷启动数据为模型训练奠定了坚实的基础，使模型在后续的强化学习阶段能够更稳定地学习和优化。它解决了纯强化学习训练中可能出现的可读性差和语言混杂等问题。

■ 多阶段训练



降本提能：架构创新，技术增效

DeepSeek通过架构创新和模型蒸馏技术，在提升模型性能的同时，显著降低计算成本和内存占用。这些技术不仅在长文本处理、代码生成、数学推理等任务中表现出色，还为大模型的轻量化和实际应用提供了有力支持。

架构创新

混合专家（MoE）架构

通过将模型划分为多个专家模块，实现高效计算和推理。DeepSeek通过无辅助损失的自然负载均衡和共享专家机制，解决了专家模块工作量不平衡的问题。

多头潜在注意力（MLA）机制

通过低秩压缩减少推理时的内存占用，同时保持与传统多头注意力（MHA）相当的性能。MLA在训练中减少了内存和计算开销，在推理中降低了KV缓存占用空间。

多令牌预测（MTP）

通过序列化预测未来多个令牌，增强模型的上下文建模能力，并支持推测解码加速推理。MTP在特定场景下同时预测多个令牌，提高信号密度，减少上下文漂移和逻辑连贯性问题。

FP8混合精度训练

采用FP8混合精度训练，通过在训练过程中使用更适宜的数据精度，减少了计算量和存储需求。FP8混合精度训练在保证训练准确性的基础上，显著降低了计算成本，使得大规模模型训练更加可行。

模型蒸馏技术

DeepSeek采用模型蒸馏技术，通过将知识从大型复杂模型（教师模型）迁移到小型高效模型（学生模型），实现性能和效率的双重优化。DeepSeek选择了多个开源模型作为蒸馏的目标模型，包括Qwen 系列和Llama 系列



- 推理效率提升：**蒸馏后的模型参数量大幅减少，例如 DeepSeek-R1-Distill-Qwen-7B 的参数量仅为7B，相比原始的 DeepSeek-R1（671B参数），计算复杂度显著降低。
- 性能优化：**在代码和数学基准测试中，蒸馏技术显著提升了模型性能。例如，在基准测试中，蒸馏后的 DeepSeek-V2.5 模型在 Pass@1 和 Length 指标上均显著优于基线模型。

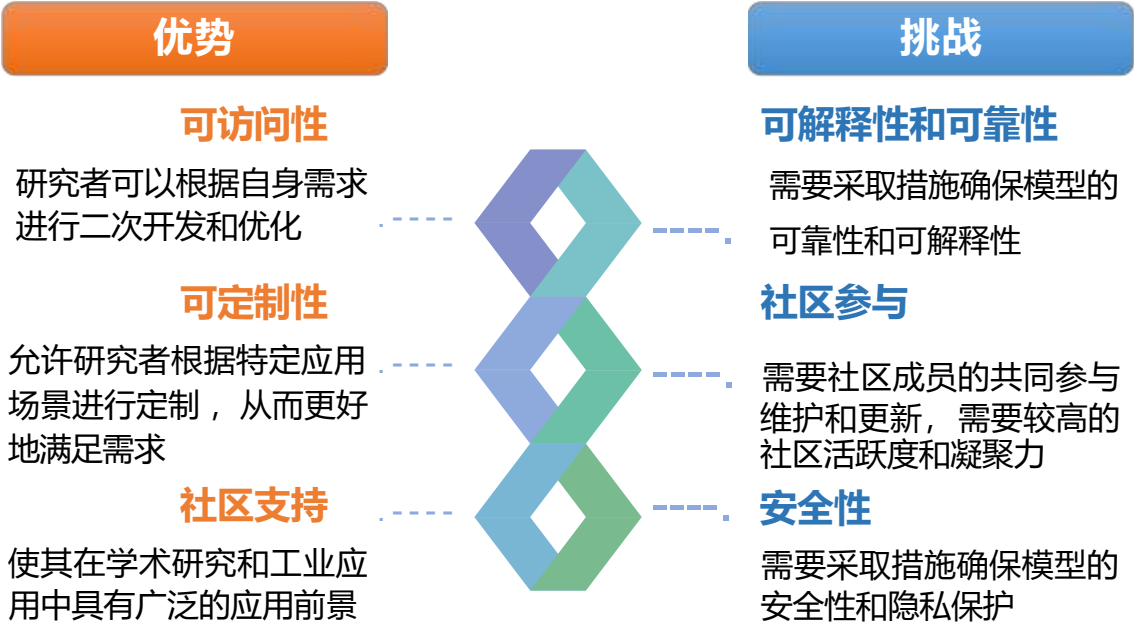
策略优化：开源特性，成本优势

DeepSeek采用开源策略，公开模型权重和技术报告，允许开发者自由使用、修改和分发其技术，促进了AI领域的创新和合作。

■ 开源策略

DeepSeek R1 采用 MIT 许可协议开源发布，允许全球的研究者和开发者免费使用和修改模型。这种开放策略促进了 AI 技术的普及和发展。

■ 开源模型的优势与挑战



DeepSeek 通过技术创新和优化策略，大幅降低了模型训练和推理成本，使其在性价比上远超 OpenAI 等竞争对手。

■ 成本优势

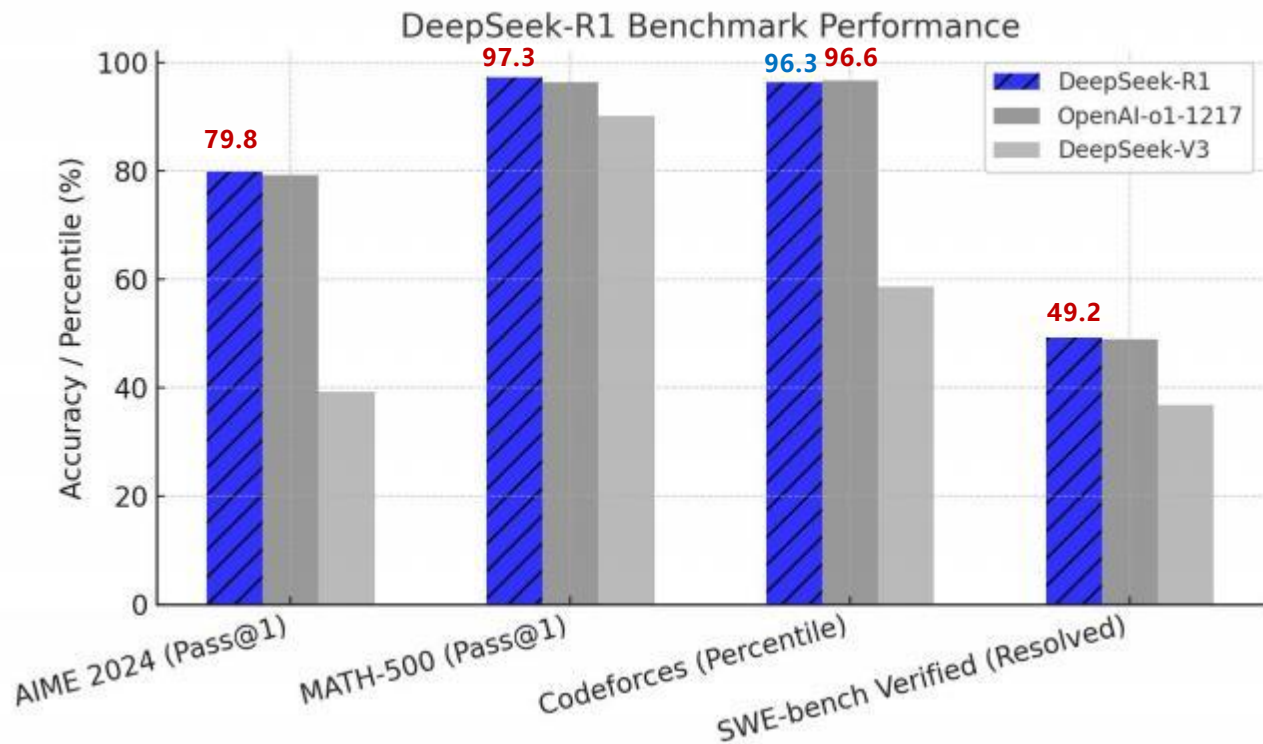
- **训练成本：**DeepSeek V3 的训练成本仅为 557.6 万美元，远低于其他国际大公司的训练成本。这种低成本策略使得更多企业和开发者能够负担得起高性能 AI 模型的训练和使用。
- **调用成本：**DeepSeek R1 的 API 服务定价为每百万输入 tokens 1 元（缓存命中）/4 元（缓存未命中），每百万输出 tokens 16 元，输出 API 价格仅为 OpenAI o1 的 3%。这种低廉的 API 价格进一步降低了使用门槛。

模型	训练成本	调用成本 (输入/百万 tokens)	调用成本 (输出/百万 tokens)
DeepSeek-V3	557.6万美元	0.14美元 (缓存未命中) / 0.014美元 (缓存命中)	0.28美元
DeepSeek-R1	未明确 (推测低于V3)	0.14美元 (缓存命中) / 0.55美元 (缓存未命中)	2.19美元
OpenAI GPT-4o	10亿美元	2.5美元 (缓存未命中) / 1.25美元 (缓存命中)	10美元
OpenAI o1	未明确 (推测高于GPT-4o)	15美元 (缓存未命中) / 7.5美元 (缓存命中)	60美元
Claude 3.5 Sonnet	5亿美元	3美元	15美元

测试评估：对标顶尖，能力出众

推理任务表现

- **数学推理能力对标顶尖模型：**DeepSeek R1 在 AIME 2024 基准测试中得分 79.8% (pass@1) ，略优于 OpenAI-o1-1217；在 MATH-500 测试中，取得 97.3%，表现与 OpenAI-o1-1217 相当，远超其他模型。
- **代码生成能力达专家级水平：**DeepSeek R1在编程任务中，Elo评分达 2029，超越 96.3% 的人类参赛者；在工程任务中DeepSeek-R1表现略优于 DeepSeek V3，这对开发人员在实际任务中有潜在帮助。



知识类任务表现

- **教育类知识问答能力突出：**在 MMLU 、MMLU-Pro等测试中，DeepSeek R1成绩超越 OpenAI-4o等其他闭源模型。

其他任务表现

- 在**创意写作、问答、编辑、摘要**等任务中，DeepSeek R1 表现优异。
- **非考试类智能处理能力强大：**在 AlpacaEval 2.0 和 ArenaHard 中，胜率分别为 87.6% 和 92.3%。

Benchmark (Metric)		Claude-3.5-Sonnet-1022	GPT-4o-0513	DeepSeek-V3	OpenAI-o1-mini	OpenAI-o1-1217	DeepSeek-R1
Architecture		-	-	MoE	-	-	MoE
# Activated Params		-	-	37B	-	-	37B
# Total Params		-	-	671B	-	-	671B
English	MMLU (Pass@1)	88.3	87.2	88.5	85.2	91.8	90.8
	MMLU-Redux (EM)	88.9	88.0	89.1	86.7	-	92.9
	MMLU-Pro (EM)	78.0	72.6	75.9	80.3	-	84.0
	DROP (3-shot F1)	88.3	83.7	91.6	83.9	90.2	92.2
	IF-Eval (Prompt Strict)	86.5	84.3	86.1	84.8	-	83.3
	GPQA Diamond (Pass@1)	65.0	49.9	59.1	60.0	75.7	71.5
	SimpleQA (Correct)	28.4	38.2	24.9	7.0	47.0	30.1
	FRAMES (Acc.)	72.5	80.5	73.3	76.9	-	82.5
	AlpacaEval2.0 (LC-winrate)	52.0	51.1	70.0	57.8	-	87.6
	ArenaHard (GPT-4-1106)	85.2	80.4	85.5	92.0	-	92.3
Code	LiveCodeBench (Pass@1-COT)	38.9	32.9	36.2	53.8	63.4	65.9
	Codeforces (Percentile)	20.3	23.6	58.7	93.4	96.6	96.3
	Codeforces (Rating)	717	759	1134	1820	2061	2029
	SWE Verified (Resolved)	50.8	38.8	42.0	41.6	48.9	49.2
	Aider-Polyglot (Acc.)	45.3	16.0	49.6	32.9	61.7	53.3
Math	AIME 2024 (Pass@1)	16.0	9.3	39.2	63.6	79.2	79.8
	MATH-500 (Pass@1)	78.3	74.6	90.2	90.0	96.4	97.3
	CNMO 2024 (Pass@1)	13.1	10.8	43.2	67.6	-	78.8
Chinese	CLUEWSC (EM)	85.4	87.9	90.9	89.9	-	92.8
	C-Eval (EM)	76.7	76.0	86.5	68.9	-	91.8
	C-SimpleQA (Correct)	55.4	58.7	68.0	40.3	-	63.7

本地部署：灵活高效，协同优化

DeepSeek的本地部署在性能上表现出色，能够满足不同应用场景的需求，尤其是在端侧和端云协同场景。通过合理的硬件配置和优化策略，DeepSeek可以在本地环境中高效运行，为用户提供强大的AI支持。

端侧部署能力

DeepSeek 在端侧部署中展现出较强的适应性和灵活性。

模型轻量化

DeepSeek通过蒸馏技术优化小模型（1.5B/7B/8B/14B/32B/70B参数规模），使其在本地部署中表现出色，适合存储和计算资源有限的端侧设备。

硬件兼容性

支持英特尔、英伟达等主流硬件平台，并可通过AnythingLLM和Ollama等工具实现PC本地部署，保护数据隐私的同时满足定制化需求。

离线能力

DeepSeek 支持完全离线部署，适合网络条件受限的场景（如工业物联网、偏远地区）。

实时性

在端侧设备上，DeepSeek能够满足实时性要求，例如在智能家居、自动驾驶等场景中，推理延迟低至毫秒级。

端云协同优化

DeepSeek的本地部署与云端计算相结合，实现高效的计算和传输。例如，其蒸馏模型在端侧SoC（系统级芯片）上的表现，显著降低了硬件门槛，同时提升了用户体验。

任务分配与负载均衡

数据传输与延迟优化

模型更新与协同训练

对比优势：高性价比，技术普惠

“与国内外顶尖同类产品比较，DeepSeek践行强化逻辑推理（R1）与长文本效率（V3）的差异化技术路线，其在性能和成本方面展现出出色的性价比，尤其在训练成本和开源透明度方面具有明显优势。

公司	模型	产品类型	核心功能	优点	缺点
DeepSeek	DeepSeek R1	开源推理模型	复杂推理、数学解题、代码生成	逻辑推理能力顶尖； 开源生态支持自定义；训练成本低	长文本生成能力弱于 V3 工程类任务上稍逊于OpenAI O1
DeepSeek	DeepSeek V3	开源大语言模型	多语言处理、长文本生成、代码生成	MoE 架构效率高；长文本处理强； 中英文混合场景优化	在推理能力上稍逊于R1 在特定任务上稍逊于OpenAI O1
OpenAI	OpenAI O1	闭源推理模型	复杂推理、文本生成	企业级 API 生态完善； 多模态交互流畅；开发者工具丰富	训练成本高；闭源且费用高昂； 中文支持弱于本土模型
OpenAI	GPT-4o	闭源大语言模型	多语言处理、文本生成、 创意内容创作	全模态能力行业领先； 实时交互响应快；商业化成熟度高	训练成本高；运营成本高 数据隐私争议大
Meta	Llama 3.2	开源大语言模型	多语言支持、内容生成、 信息检索	完全开源免费；社区支持广泛； 多语言基础能力均衡	多模态功能缺失； 长文本生成质量不稳定
Anthropic	Claude-3.5	闭源推理模型	对话系统、内容生成、 逻辑推理	对话逻辑连贯性强； 伦理安全性高；文档分析能力突出	中文支持较弱； 闭源且 API 访问受限
百度	文心一言	闭源大语言模型	多语言处理、复杂的语言理解和文本生成	中文场景优化最佳； 多模态搜索整合；本土行业适配性强	国际竞争力不足； 上下文窗口较小

革新技术标准： 低本高能， 开放共创



DeepSeek的成功促使AI行业重新审视技术应用与发展方向。其低成本、高性能的模型为AI技术的普及提供了实际范例，推动了AI技术在训练成本、模型效能和开源生态方面的新标准的形成。

创新技术路径

DeepSeek通过算法优化与架构创新（如MLA、MoE结构），将训练成本降至行业1/10，打破了传统AI巨头依赖“规模法则”的垄断局面。其FP8混合精度训练和开源原生FP8权重，显著降低了中小团队的技术门槛，推动AI技术民主化。



01

重塑定价逻辑

DeepSeek V3模型以557.6万美元的训练成本，实现了与GPT-4o相当的性能，生成速度提升至60 TPS。这种“低成本高性能”模式不仅挑战了OpenAI、Google等巨头的市场地位，还迫使行业整体降价（如字节豆包降价85%），重塑了AI服务的定价逻辑。



C2

推动研发转型

DeepSeek的全栈开源策略（模型权重、训练代码均采用MIT协议），吸引了全球开发者参与，形成了强大的社区生态。这种开放模式加速了技术迭代，削弱了闭源巨头的技术壁垒，推动全球AI研发从“封闭垄断”向“开放协作”转型。



03

重塑产业格局：打破桎梏，竞争活跃

DeepSeek R1 的全球影响力正在重塑 AI 产业格局，特别是在中美之间的技术竞合中。同时，也为全球 AI 产业的发展提供了新的机遇和挑战。

■ 中美技术竞合

DeepSeek的创新不仅打破了美国AI产业的技术壁垒，也为中国AI产业在全球科技竞争中提供了新的突破口。

DeepSeek的成功推进中国AI产业的发展，同时也促进了中美两国在AI领域的竞争与合作，推动全球AI技术的多元化发展。

DeepSeek的横空出世

给美国科技市场带去巨大冲击

- 受其影响，美国芯片巨头英伟达的股价暴跌17%，博通下跌17%，AMD下跌6%，微软也下跌了2%。
- DeepSeek的应用程序在苹果应用商店的下载量一举超越了ChatGPT，荣登免费应用程序排行榜榜首。

■ 活跃市场竞争

DeepSeek的崛起改变了AI市场的竞争格局，促使国际科技巨头加快技术创新的步伐，加大研发投入，推出新的模型和应用，以应对竞争。

Open AI

上线新一代推理模型o3系列的mini版本，并首次免费向用户开放其基础功能。o3-mini专注于数学、科学和工程等领域的复杂推理任务，其性能和成本效益均优于之前的o1系列。

谷歌

发布新一代Gemini 2.0系列模型，包括Gemini 2.0 Pro、Gemini 2.0 Flash、Gemini 2.0 Flash-Lite和Gemini 2.0 Flash Thinking，旨在提升AI能力并提高性价比。

■ 全球AI产业链升级

DeepSeek的崛起带动了全球AI产业链上下游的发展。其低成本高性能的模型降低了大模型的投资、开发、运营成本，推动了国产AI芯片、云平台、操作系统等产业的发展。

技术深化：突破局限，能力提升



DeepSeek R1展示了强化学习技术和算法创新在 AI 领域的巨大潜力，但其仍然处于发展阶段，存在一定局限性和优化空间。未来，随着技术的不断进步和创新，DeepSeek R1 可能会在以下几个方面实现进一步的突破：

通用能力提升

目前，DeepSeek R1在函数调用、多轮对话、复杂角色扮演和 JSON 输出等任务中的能力不及 DeepSeek-V3。未来，DeepSeek计划探索如何利用长推理链来增强在这些任务的表现。

解决语言混杂问题

DeepSeek R1当前只针对中文和英文进行了优化，这可能在处理其他语言的查询时导致语言混杂问题。DeepSeek计划在未来的更新中解决这一局限。

优化提示工程

目前模型对提示较为敏感，少样本提示会持续降低其性能。因此，建议用户使用零样本设置，直接描述问题并指定输出格式，以获得最佳效果。

软件工程任务

DeepSeek-R1 在软件工程基准测试中的表现未能显著超越 DeepSeek-V3。未来版本将通过在软件工程数据上实施拒绝采样或在强化学习过程中引入异步评估来提高效率。



场景拓展：创新推动，垂直深耕

DeepSeek R1将通过强化学习和多模态融合等技术手段，进一步提升推理能力、优化语言理解和生成效果，并拓展在复杂任务中的应用边界；同时，将深耕垂直领域，如教育、金融、医疗等，为不同领域提供更精准、高效的解决方案。

技术创新推动

多模态融合

DeepSeek未来可能会在多模态融合方面进一步探索，将自然语言处理、计算机视觉等技术更深度地结合。

具身智能探索

与机器人等硬件深度融合，实现物理世界的智能交互。这将拓展其在工业制造、物流配送等领域的应用。

自进化系统构建

通过自动合成训练数据，持续迭代模型能力。这将使其能够更好地适应不同垂直领域不断变化的需求，提升在各领域的应用效果。

垂直领域深耕



医疗领域

DeepSeek已经在医疗辅助诊断方面有所应用，未来有望进一步深化，如通过流程优化，提高诊断的准确性和效率。通过与医疗设备的结合，实现更精准的医学影像分析和疾病预测。



金融领域

未来，DeepSeek可能会进一步拓展到金融风险防控、智能投顾、金融产品创新等领域，通过深度分析金融市场数据和用户行为数据，为金融机构提供更全面、精准的决策支持。



教育领域

目前DeepSeek在教育辅助方面已经展现出独特优势，未来，其可能会与在线教育平台、教育机构等合作，开发更多个性化的学习方案和智能辅导工具，满足不同学生的学习需求。



法律领域

DeepSeek在法律文书处理方面已经具备一定的能力。未来，其有望进一步拓展到法律咨询、案件预测、法律知识图谱构建等领域，为法律专业人士和普通用户提供更便捷、高效的法律服务。



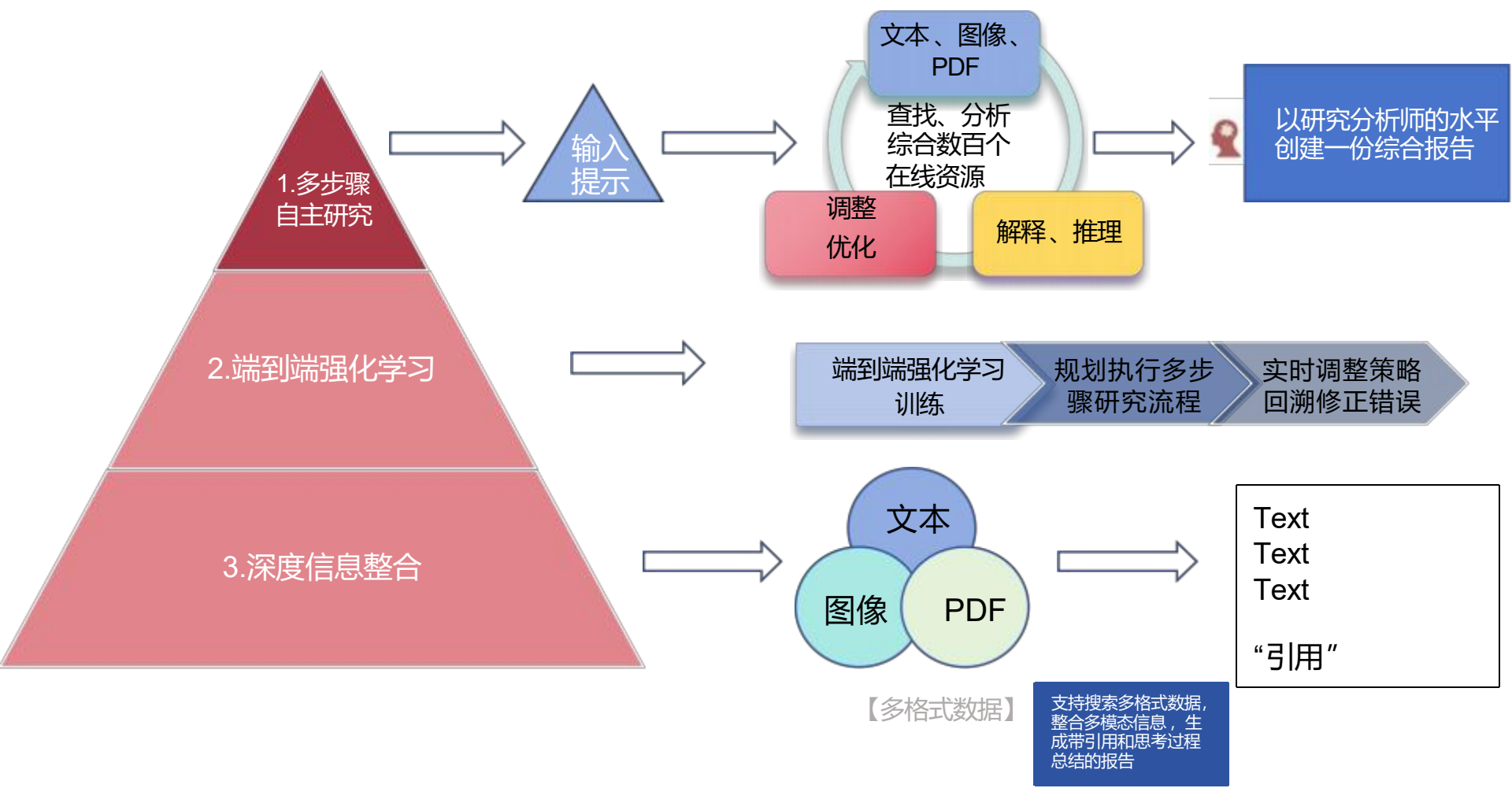
工业领域

DeepSeek在工业质检智能化方面已经取得显著成效。未来，其可能会进一步拓展到工业生产流程优化、设备故障预测与维护、供应链管理等领域，提供更高效的工业生产和运营的解决方案。

DeepResearch: 智能协作， 自主研究

「核心功能」

多步骤自主研究、端到端强化学习、深度信息整合



实际使用 图源@宝玉

在 ChatGPT 中, 选择「message composer」中的 deep research 并输入查询

可以附加文件或电子表格, 为问题添加上下文。一旦开始运行, 侧边栏将显示所采取的步骤和使用的来源摘要。

基准测试：精度提升，行业领先

表现： 人类终极考试， 准确率突破 26.6%

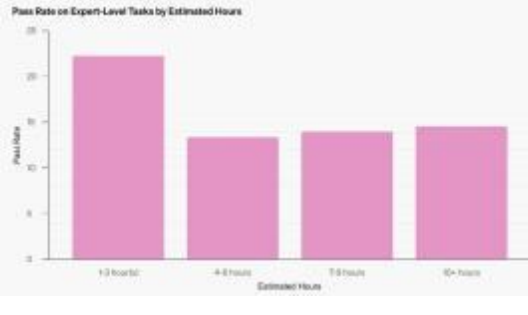
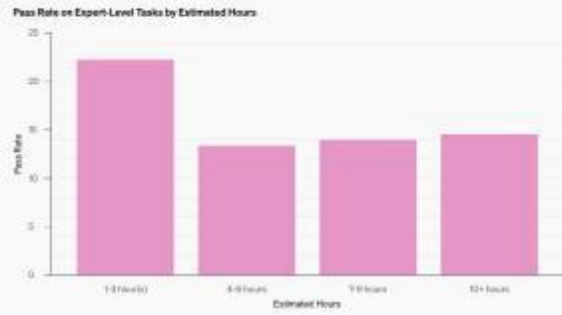
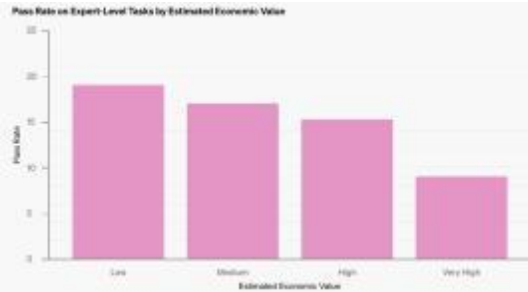
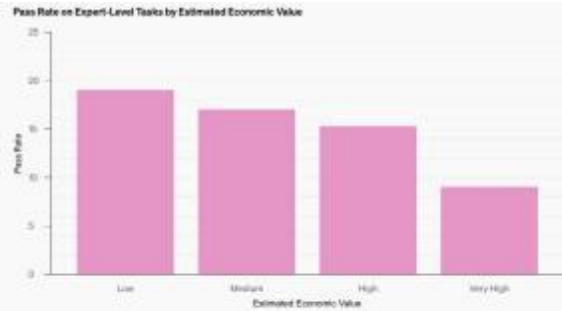
GAIA测试

这项测试包括3000多个多项选择题和简答题，
涵盖了从语言学到火箭科学、古典文学到生态学的100多个学科。

Model	Accuracy (%)
GPT-4o	3.3
Grok-2	3.8
Claude 3.5 Sonnet	4.3
Gemini Thinking	6.2
OpenAI o1	9.1
DeepSeek-R1*	9.4
OpenAI o3-mini (medium)*	10.5
OpenAI o3-mini (high)*	13.0
OpenAI deep research**	26.6

* Model is not multi-modal, evaluated on text-only subset.
**with browsing + python tools

GAIA	Level 1	Level 2	Level 3	Avg.
Previous SOTA	67.92	67.44	42.31	63.64
Deep Research (pass@1)	74.29	69.06	47.6	67.36
Deep Research (cons@64)	78.66	73.21	58.03	72.57



准确率是此前 OpenAI o1 模型的近三倍

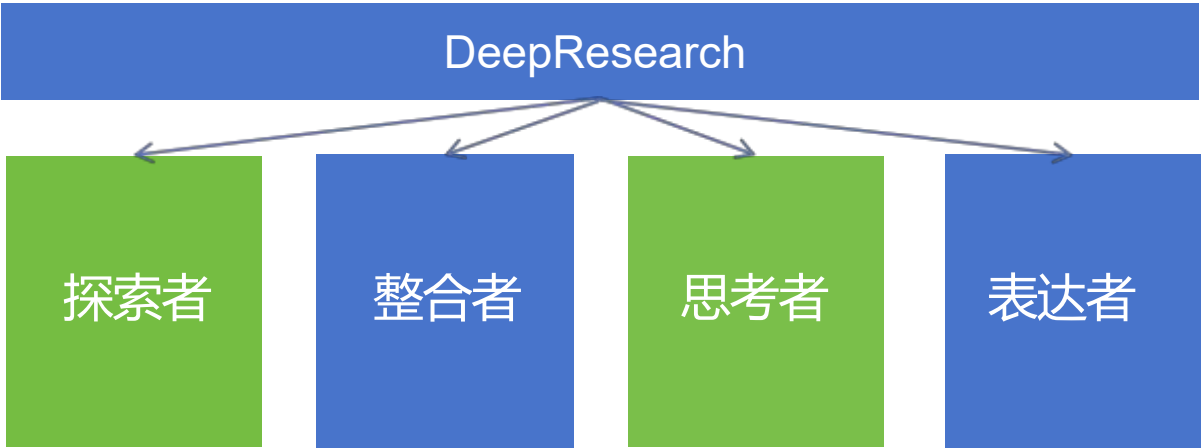
来源：<https://openai.com/index/introducing-deep-research>

技术协同：多步推理，快速输出

与GPT-4o对比

类别	DeepResearch	GPT-4o
功能目标	自动化多步骤研究任务，收集、综合、分析、输出报告	语言生成，支持多种自然语言任务
任务执行方式	多模块协同，逐步执行复杂任务	单输入文本生成输出，处理单一任务
研究能力	处理复杂学术、法律、市场研究，支持多轮分析	生成创意内容，提供建议，适度推理分析
输入输出格式	支持图像、PDF等多种格式输入输出	主要依赖文本输入输出
模块协作	多个模块协同工作（探索者、整合者、推理者等）	单一模型，无模块化协作

相比传统GPT-4o模型， Deep Research在多步推理、数据验证、处理速度和信息追溯性方面表现出明显优势。这些提升有助于模型在复杂任务中的表现更好，特别是在需要高可靠性和高效执行场景中。



应用场景1：学术研究，助力科研

文献综述加速

DeepResearch能迅速梳理海量文献，提炼关键信息，显著提升文献综述效率。

技术报告生成

基于深度学习模型，自动生成高质量技术报告，确保研究成果的准确传达。

自动实验设计

基于已有实验数据自动生成最优实验设计，预测可能的实验结果，并提出资源最小化、效能最大化的实验方案。

"未来知识"生成器（预测性科研）

分析过去几十年各领域的论文发展轨迹，利用深度时间序列预测技术，自动生成某一领域在未来5-10年的潜在研究主题、理论突破、以及可能的新技术趋势。

学术研究案例：明确需求，报告生成

■ 团队自测案例

通过百度网盘分享的文件：deep Research功能深入研究.docx

链接：<https://pan.baidu.com/s/1pyaygXqFXvRe-ln7gn5gOA?pwd=fn7s> 提取码：fn7s



01

科研场景实测：

生物学研究生输入“CRISPR技术在
肿瘤免疫治疗中的最新进展”



02

获得：

1. **近三年124篇**核心论文摘要
2. 关键临床**试验数据**
3. 汇总**技术路线对比图谱**
4. 待突破**方向预测**
5. **符合APA格式**的参考文献库

应用场景2： 金融分析， 市场预测

通过自动化数据收集、整合、推理与报告输出， 提供全面的市场趋势预测和投资决策支持。

场景应用

股票市场分析

风险管理与投资组合优化

宏观经济预测

数据分析

自动化处理财务报表， 挖掘隐藏的投资机会，
评估潜在风险， 优化资产配置策略。

智能预测

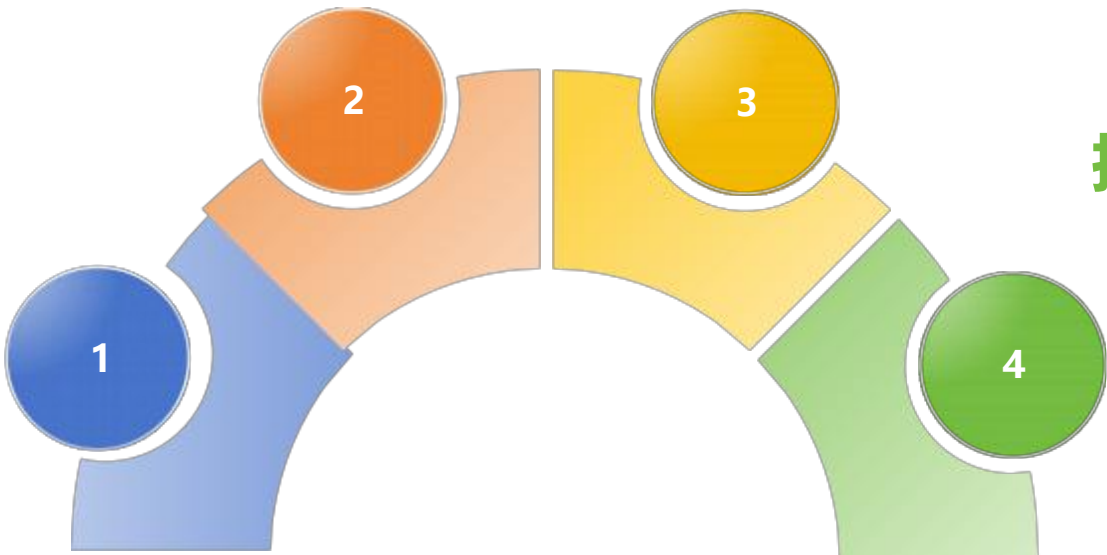
运用先进算法预测市场走势， 辅助金融机构和个人投资者做出更明智的选择。

市场洞察

DeepResearch整合全球金融市场动态， 实时追踪行业趋势， 为投资者提供深度分析。

报告生成

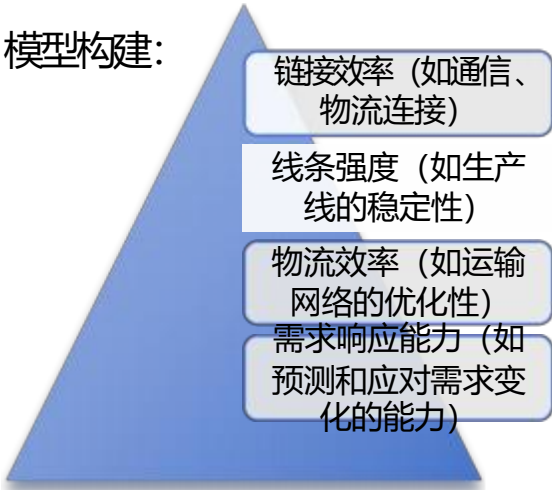
一键生成专业级投资风险评估报告，
支持定制化需求， 提升决策效率。



金融分析案例：数据整合，供应链优化

- 数据来源：全球12个交易所的财报数据
- 提取来自全球主要交易所（如纽约证券交易所、道琼斯指数等）的半导体相关财报和数据

➤ 模型构建：



➤ 情景模拟：

- ❑ 建立基于5种不同情景（如需求波动、突发事件、技术革新）的供应链模拟模型。
- ❑ 使用Deep Research提供的可视化工具生成可解释性的分析报告，展示各情景对供应链压力及影响的具体路径。

1. 数据获取

数据解析过程

- ❑ 来自行业研报机构的178份半导体供应链风险分析报告。
- ❑ 解读各研报的核心观点、关键指标及预测方法。
- ❑ 建立行业报告的质量评估体系，识别高价值研报并进行分类。

2. 模型构建与供应链脆弱性评估

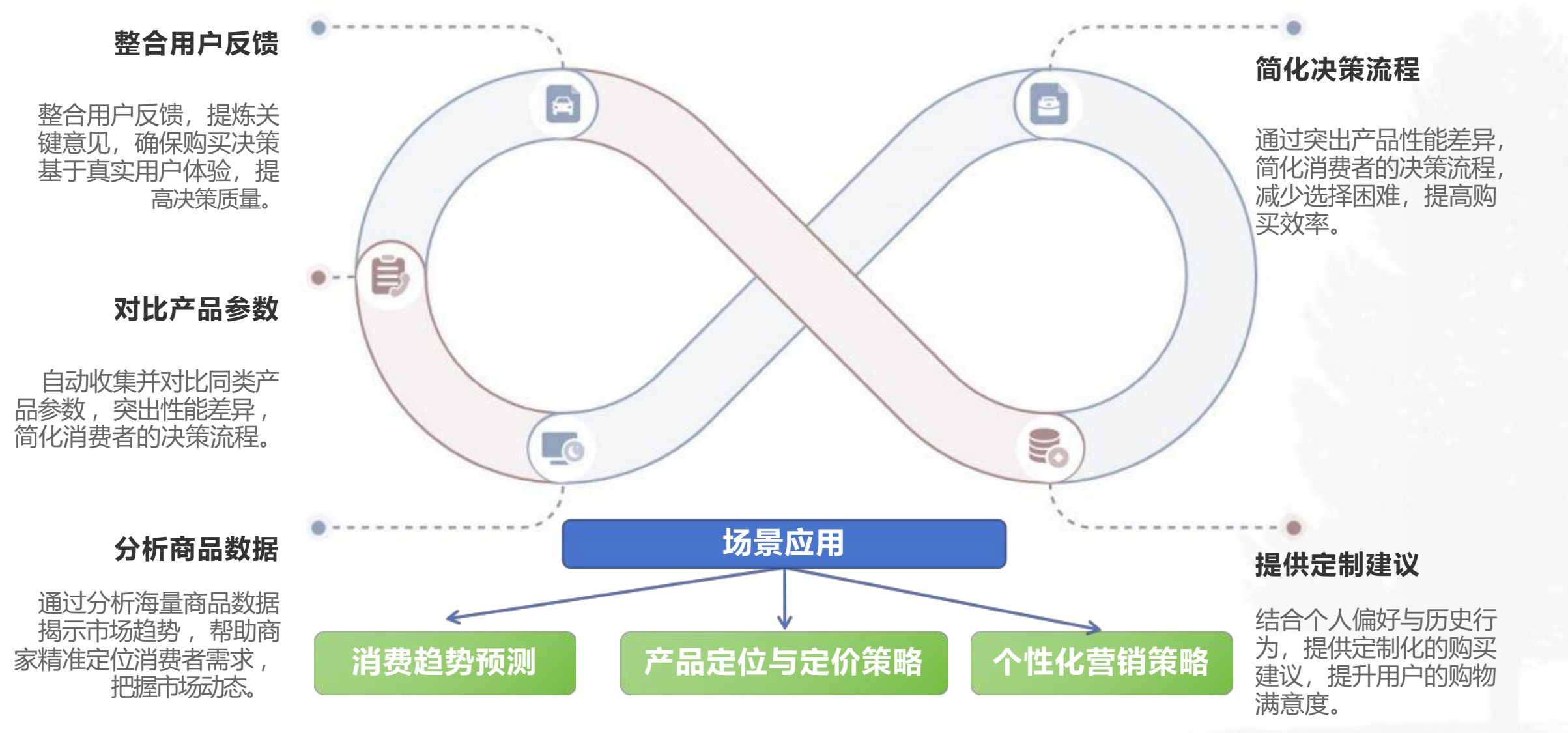
供应链脆弱性评

- ❑ 使用层次分析法对各关键因素进行权重评估，最终得出半导体供应链的脆弱性等级。
- ❑ 分析各研报中对供应链脆弱性的描述，并结合数据来源和模型构建结果，识别高风险区域。

3. 情景模拟与建议

- 在供应链风险最高的环节加强协同协作，并提供透明的沟通机制。
- 加强内部风险管理框架的设计，建立应急响应和恢复计划。
- 定期更新模型和数据来源，确保预测准确性和前瞻性。

应用场景3：消费决策，个性推荐



消费决策场景案例：需求识别，产品匹配

用户诉求：拟购买滑雪板，对滑板的使用场景、款式、颜色、价位、适用场景、预计购买国家提出要求，指定DeepResearch给出建议



问题识别

面对海量市场数据，DeepResearch识别用户的关键信息并进行网络信息匹配。



解决方案

利用 DeepResearch的 智能分析能力，自动筛选出符合用户每一部分诉求的关键词，精准度和匹配度大幅提升



成果亮点

通过自动化报告生成，显著提升了搜索结果与用户预期的匹配度，提高用户的决策效率，辅助用户消费决策。

Deepresearch直接解决了用户的每一部分需求，**从板材、规格到颜色，亮点以及可操作的滑行技巧**。同时还通过滑雪板的信息介绍，分析了板材性能与用户可能使用场景和需求进行了匹配分析

GPT-4o				Deep research	
Jones Mountain Twin	The Jones Mountain Twin is an aggressive directional twin board that blends the stability of a freeride board with the playfulness of a freestyle board. Its hybrid camber-rocker profile provides excellent edge hold for	Mid-range to Premium	Available through international retailers that ship to Japan.	CAPITA Defenders of Awesome (D.O.A.)	Medium Flex (5/10) twin that's camber-dominant with slight rocker
				All-Mountain Freestyle Twin	at tips ("Resort V1" profile). This hybrid camber gives it pop and edge hold for carving, yet added uplift at the tips improves float in powder. The true twin, blended radial sidecut
					grips well in turns, and the mid flex makes it stable at speed but agile in trees. It's a proven go-anywhere deck - CAPITA's #1 seller - so you get a

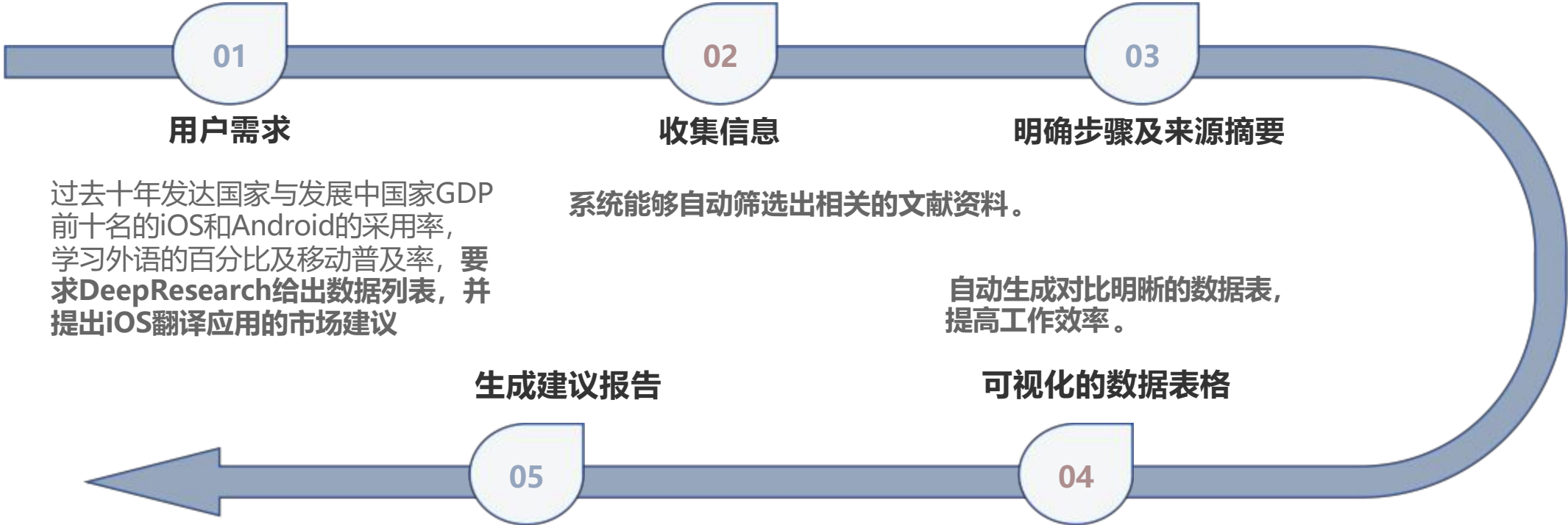
GPT-4o		Deep research					
powder snow. Hokkaido is famous for its light, dry powder. Boards with a wider nose and rocker profiles can enhance floatation in deep snow. - Temperature: The region experiences cold temperatures, which keep the snow light and fluffy. Ensure your gear, including outerwear and gloves, is suitable for low temperatures. - Local Brands: Japanese brands like Gentemstick and Moss Snowstick design boards specifically for local conditions, offering performance advantages in Hokkaido's unique snow. When selecting a snowboard, consider demoing different models if possible to find the best fit for your riding style and the specific conditions you'll encounter.		doesn't blend in - the black edition has big white GNU letters, and some variants include neon graphics. Table 1: All five recommended snowboards are medium-flex, all-mountain oriented boards that balance carving performance and powder capability. Each model is available through major retailers in Japan. Comparison of Key Specifications: <table><thead><tr><th>Snowboard Model</th><th>Flex (Rating)</th><th>Shape (Stance)</th><th>Camber Profile</th></tr></thead></table>		Snowboard Model	Flex (Rating)	Shape (Stance)	Camber Profile
Snowboard Model	Flex (Rating)	Shape (Stance)	Camber Profile				

应用场景4：商业研究，决策支持

提供深入的**市场趋势、消费者行为洞察及竞争格局**，为企业制定精准的商业策略提供支持。



商业研究场景案例：市场洞察，优化策略



GPT-4o	Deep research			
new iOS translation app powered by ChatGPT.				
Top 10 Developed Countries by GDP:				
1. United States				
2. Japan				
3. Germany				
4. United Kingdom				
5. France				
6. Italy				
7. Canada				
8. South Korea				
9. Australia				
10. Spain				

GPT-4o	Deep research			
9. Argentina				
10. South Africa				
Insights:				
• iOS and Android Adoption Rates: Globally, Android holds a dominant market share of approximately 70.93%, while iOS accounts for about 28.32% as of 2024. In the United States, iOS leads with a 58.81% market share. Other developed countries like Japan, Canada, Australia, and the United Kingdom also show a higher preference for iOS devices. In contrast, developing countries such as India and Brazil have a higher prevalence of Android devices.				
Interest in Language Learning: While specific				

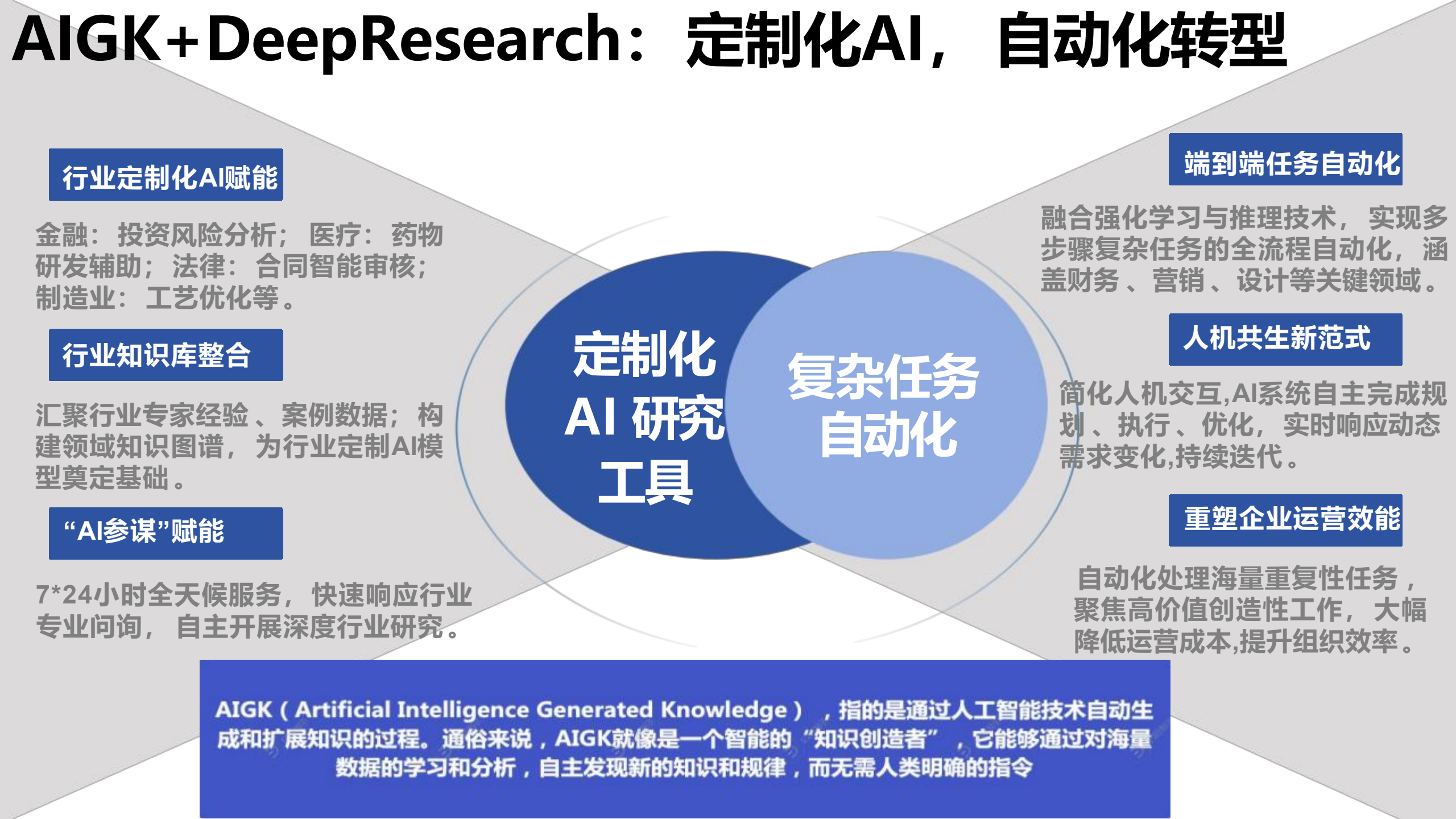
竞品分析：技术差异，性能对比

	 腾讯混元	 OpenAI Deep Research	 Google
技术架构	依赖传统RL框框架	基于OpenAI o3模型 支持多步骤策略实时调整与错误回溯	基于BERT等预训练模型 缺乏动态修正能力
性能表现	未提及类似基准测试结果 推测响应时间更长	在“人类最后考试”中准确率26.6%（O1模型的三倍），5-30分钟完成传统数小时任务。	仍依赖传统搜索+人工分析
应用广度	倾向于聚焦单一领域 (如金融)	覆盖学术、金融、消费、政策四大领域，案例已验证跨领域适配性	侧重通用搜索，缺乏垂直场景优化
用户体验	需分步提交指令	零门槛操作，用户仅需输入提示，系统自主规划研究流程	需手动配置参数或分阶段操作

通过百度网盘分享的文件：不同版本DeepResearch对比.zip

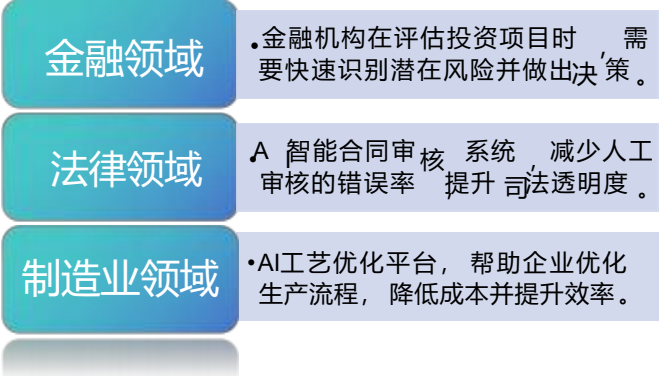
链接: https://pan.baidu.com/s/1vYEXJ_gpJMMYAXMXyYX3oQ?pwd=jjv9 提取码: jjv9

--来自团队自测数据



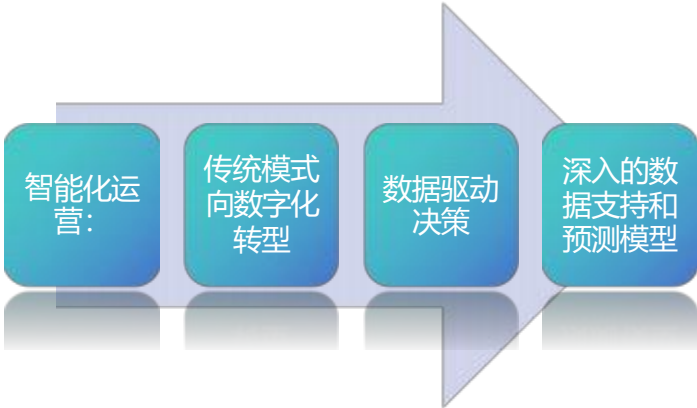
行业应用：AI定制，自动化决策

实施场景



数据来源：

- ❑ 收集来自不同行业的数据源，包括但不限于企业财务报表、行业政策文件、专家报告等。
- ❑ 将这些数据从多个平台获取，并进行清洗和预处理。



模型分析：金融行业投资风险智能预测

- 金融公司：基于实时数据，智能预测投资项目的风险等级（高、中、低）。
- 投资者：通过AI分析后，可获得投资决策的个性化建议。

- ✓ 金融机构决策支持：通过AI模型预测投资项目的风险等级，帮助银行做出更精准的投资决策。
- ✓ 投资者个性化决策：基于AI分析后，可获得特定投资标的的风险评分，进一步优化投资策略。

按照行业或主题对数据进行分类存储，例如：



- ❑ 快速响应能力：在各种行业需求瞬息万变的情况下，“AI参谋”能够提供即时的数据分析和决策支持，帮助客户迅速定位问题并制定解决方案。
- ❑ 自动化处理：系统通过算法自动识别异常数据、预测市场趋势，并生成快速反应的建议。

技术创新：流程自动，突破效能

● 端到端任务自动化

任务分析与状态跟踪

- 自动识别任务的基本要求和限制条件。
- 使用强化学习算法跟踪任务的状态变化（如预算使用、产品数量等）。

动态预测与优化

- 基于推理技术，实时预测未来的市场需求或用户行为。
- 根据预测结果优化资源分配和决策流程，确保高效性。

反馈与迭代

- 通过数据反馈（如实际的使用情况）更新模型参数。
- 进行持续优化，提升系统的适应能力和效率。

● 人机共生新范式

1. 自动化处理与智能化决策

- 人机协同：优化人机交互界面，减少人为干预，提升效率。
- 人机协作：支持多场景下的实时响应，确保快速精准决策。

2. 深度行业研究与数据驱动

- 基于深度学习的行业趋势预测模型，支持企业动态适应市场变化。
- 数据驱动的人工智能模型，实现专业预测和战略规划。

3. 复杂任务的全流程自动化

- 融合强化学习算法，自动识别高风险场景并提供相应建议。
- 深入分析数据，优化决策流程，确保全面覆盖核心业务环节。

● 重塑企业运营效能

数字化转型与高效管理

- 通过技术创新、优化结构和提升效率，实现企业的可持续发展。

自动化处理与智能化决策的突破

- 自动化任务：包括财务、营销和设计等领域的重复性操作。
- 通过算法和机器学习技术，自动识别高风险场景并提供相应的建议或优化方案。
- 智能化决策：融合强化学习与推理技术，实时预测和优化决策过程。
- 基于数据，动态跟踪任务的状态变化（如预算使用、产品数量等）。

聚焦高价值创造性工作

- 融合先进管理理念和技术手段，提升企业核心竞争力。高效操作减少人为错误，提高响应速度。
- 成本下降：自动化减少人工干预中的固定成本。

认知协作：异构智能，集群协作

三阶认知生成体系

异构智能体集群

- 数据勘探者（5个垂直领域AI）
- 逻辑架构师（3个推理引擎）
- 批判审查团（2个逆向思维AI+人类专家接口）

行动建议梯度

- stata 复制
- 立即布局（胜率>75%）
 - 二线电池厂商技术并购标的筛选（附潜在目标清单）
 - 持续跟踪（变盘阈值监测）
 - 4680电池量产良率 vs 特斯拉股价弹性系数
 - 风险对冲（黑天鹅防护）
 - 建议建立稀土期货空头头寸：每亿元持仓对应0.83风险抵消因子

技术亮点标注

- 创新性辩论机制：智能体间设置「认知温差」强制产生差异化视角
- 动态知识注入：每小时融合50+权威信源与2000+传感器数据流
- 可解释增强：点击任何结论可追溯完整逻辑链（示例见附录A）

AIResearch生成报告样本：《新能源汽车产业链投资机遇分析》

封面

- 生成标识：■ 本报告由XXX智能体集群经17轮辩论达成共识
- 时间戳：知识截止至2025-02-6 14:32:00

核心结论 (Consensus Core)

颠覆性机会

- 固态电池电解质材料：
[AI分歧度12%] 2026年成本下降曲线斜率较预期提升23%（卫星图像分析显示中试产线良率超预期）

风险预警

- 稀土永磁供应链：
[AI置信度91%] 缅甸政治动荡将导致2024Q4钕铁硼价格波动率超历史极值68%

多维论证矩阵

维度	数据侦探AI发现	逻辑架构AI推演	反事实审查AI挑战
技术演进	拆解23万份专利发现：磷酸锰铁锂研发集中度超预期	材料突破将重构电池厂商议价能力模型	若宁德时代专利诉讼败诉将逆转趋势
消费行为	充电桩评价数据揭示：北方用户冬季续航焦虑被低估58%	热管理系统升级需求催生200亿增量市场	假设油价跌至60美元将延缓电动化进程

动态知识图谱

python 复制

```
# 关键因果链可视化
Graph().node("政策").link("双积分制收紧").to("混动技术投入+32%")
      .node("锂价").link("期货曲线").to("钠离子电池商业化提速")
      .conflict("欧盟碳关税","国内补贴退坡").mark_red()
```

引入优化agent：复杂任务，实现自动化

传统方法的局限

- ❑ 目前AI主要是“助手”角色，需要用户提供明确指令，无法自主完成复杂任务。
- ❑ 现有AI工具难以跨多个子任务自动执行，仍需人工介入。

创新点

● AI 自主任务规划与执行(AI Agent)

AI 能够自主分解任务、规划步骤，并利用外部工具(如API、数据库、自动化流程)执行任务。

● 多 AI 代理协作

不同 AI 代理(市场分析 Agent、法律审核 Agent、财务预测 Agent) 可协同完成复杂任务，形成智能工作流。

● 任务反馈 & 自主学习

AI 在执行任务后自动优化策略，使任务执行效果不断增强。

》应用示例

- **智能法律顾问 A1**:自动读取合同，分析潜在法律风险，生成修改建议，并与企业法务系统对接完成合规审查。
- **企业 AI CEO**:结合市场数据、财务数据，自动生成年度战略规划，并动态调整业务目标。
- **智能招聘 A1**:筛选简历、面试候选人(语音/视频 AI 面试)、自动发送 offer，并完成 HR 系统录入。

增强知识图谱：多维解释，溯源路径

传统方法的局限

- ❑ 幻觉率过高，高价值信息过少，致使企业用户难以信任 AI 生成的行业研究和决策结果。

创新点

知识图谱增强 LLM(LLM+KG)

结合 AI 生成知识(AIGK)与行业知识图谱，使 AI 具备强逻辑推理能力。

可追溯的 AI 研究报告

所有 AI 生成的内容提供可溯源数据，确保数据可信度。

可解释的 AI 运行决策

AI 在做出决策时，会提供基于知识图谱的逻辑推理路径，增强可解释性。

应用示例

➤ 金融风险评估与决策支持：

通过结合金融知识图谱和AIGK技术，AI能够提供透明的决策过程和可解释的投资建议，增强金融决策的信任度。

➤ 医疗诊断与个性化治疗：

利用医学知识图谱和AIGK推理，AI不仅给出治疗方案，还能解释每个建议背后的医学依据，提升医生和患者的信任。