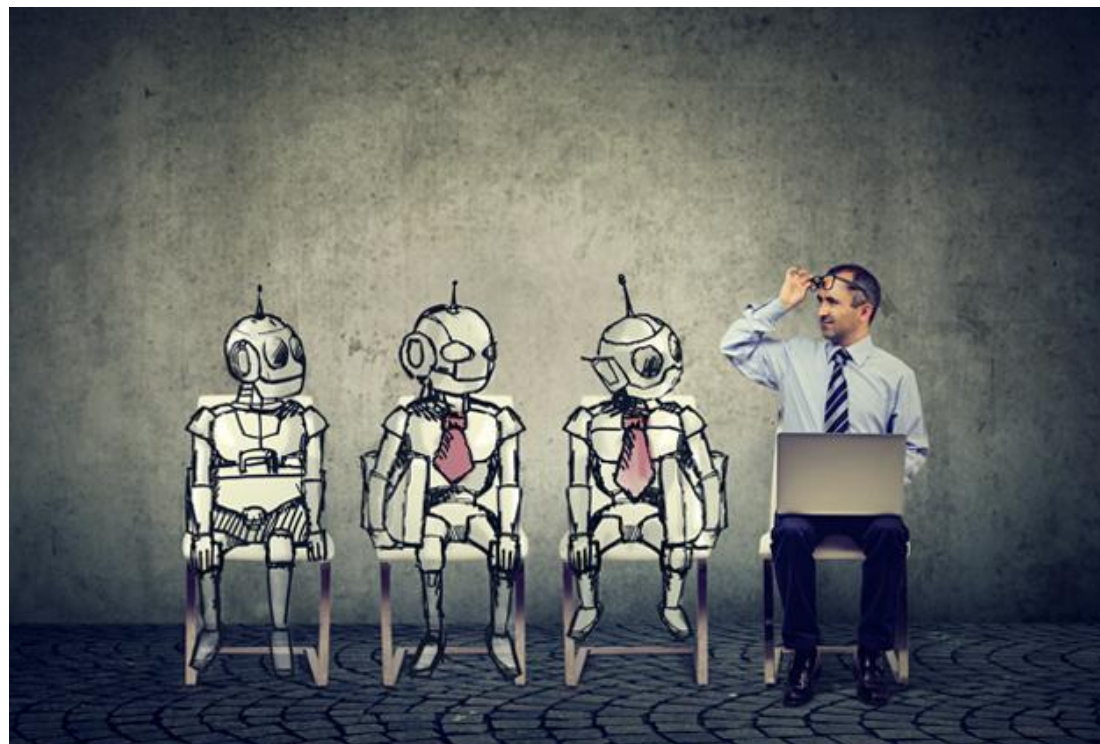


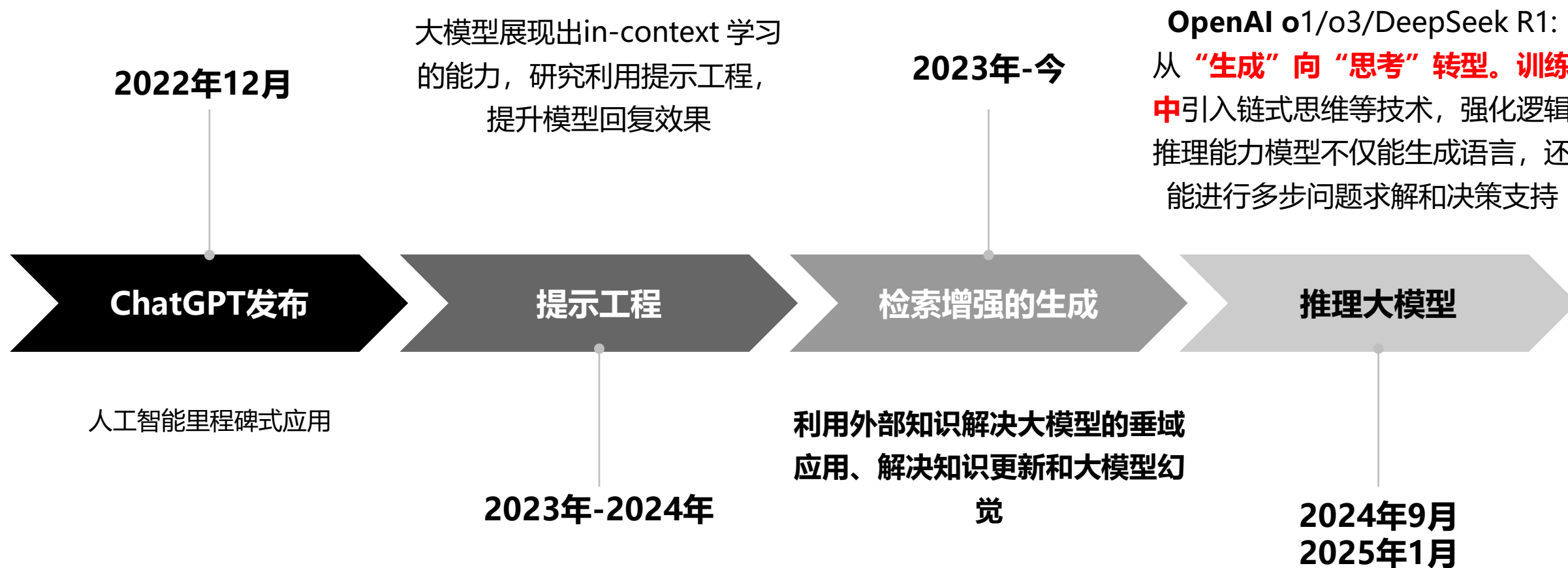


临界颠覆：人工智能重构企业生产力逻辑

张诚

26.4.24-4.26

这3年过得有点快：大模型技术的迭代与演化



- 大模型正在从**单纯的生成**转向**更智能的理解与推理**



OpenClaw

THE AI THAT ACTUALLY DOES THINGS

Clears your inbox, sends emails, manages your calendar, checks flight status, and more.
All from WhatsApp, Telegram, or any chat app you already use.

NEW Introducing OpenClaw →

What People Say

"Today and tomorrow, OpenClaw is the closest to a true AI assistant I've ever used. Truly a game-changer."



"@openclaw is absolute game changer for me. The amount of things I done from my phone just during my breakfast is absolutely breathtaking."

@SedRicKCZ



"Just told Ema, my @openclaw to turn off the PC (and he did it) Executed perfectly."

@bangkokbuild

"That OpenClaw has a 'Hackable' install is a huge plus. This should be a standard for OSS projects."



"I asked it to take picture of the sky whenever it's pretty. It designed a skill and took a pic!"

@sinelecinia



"Finally got it. It's truly amazing."

@anishk



我们需要转变思维，不仅仅将 AI 视为工具，而应将其视为一种新的经济主体/方式

低成本适应：通用/基础模型能够更低成本**适应更多场景**

基础模型（如 GPT - 4）经过大规模数据训练的，成为灵活通用的人工智能平台。核心特点是**适应性强**，只需少量额外训练就能应用于各种下游任务，使企业无需开发昂贵的专用系统，而可以依赖少数共享模型

AI vs. 外包/内容生成

AI 的可扩展性，即相当于引入了大规模同质化外包。而知识外包要么是小规模同质化，要么是大规模但异质化的，颠覆了之前的外包模式和效果

比如软件外包/AI编程

2025年AI编程行业全景研究：从智能辅助到自主开发的范式转变

懂点AI 2025-07-07 19:27 海南

AI 导读

全球AI编程工具市场正以27.1%的年增速爆发，中国增速达38.4%领跑全球。从代码补全到全流程开发，AI正通过自然语言编程、多模态交互重塑软件工程，GitHub Copilot等工具已让开发者效率提升30%以上。

内容由AI智能生成

有用 | 分享

一、行业概述：AI 编程工具市场爆发式增长

1.1 市场规模与增长趋势

AI 编程工具市场正经历前所未有的增长。根据 QYResearch 的统计及预测，2023 年全球 AI 编码工具和助手市场销售额达到 48.6 亿美元，预计到 2030 年将达到 262.2 亿美元，年复合增长率 (CAGR) 高达 27.1%(2024-2030)(74)。另据 Spherical Insights 预测，全球 AI 编码工具市场规模将在 2032 年突破 295 亿美元(70)。

中国市场增长尤为迅猛。2023 年中国 AI 代码生成市场规模达到 65 亿元人民币，预计到 2028 年将增长至 330 亿元人民币，年复合增长率 (CAGR) 高达 38.4%，远超全球 24.3% 的平均水平(26)。这种增长势头反映了中国企业对 AI 编程工具的强烈需求和快速接纳。

比如软件外包/AI编程



在这个世界级编程竞赛中，这可能是人类最后一次战胜AI了。

大数据文摘 2025年07月25日 15:02 山西

以下文章来源于数字生命卡兹克，作者数字生命卡兹克



数字生命卡兹克

努力分享一些很新、很酷的AI干货，愿我们永远对世界保持好奇。

大数据文摘受权转载自数字生命卡兹克

作者：卡兹克

昨晚，一场轰轰烈烈的人类 vs AI的比赛结束了。

而人类，这一次，险胜。

暂时守住了胜利。



Psycho

@FakePsyho

人类暂时胜利了！

我完全累垮了。我算了下，过去三天只睡了10个小时，几乎快撑不住了。

休息好后我会发布更多关于比赛的内容。

(说明一下，这些是临时结果，但我的领先应该足够大)

[翻译帖子](#)

Rank	User	Score	A
1	Psycho	45,245,838,577 / 200	45,245,838,577 / 200

AtCoder World Tour Finals (简称AWTF) 是由日本知名的在线编程竞赛平台AtCoder举办的全球编程赛事

有两个赛道，Algorithm (算法) 和Heuristic (启发式)。

Algorithm (算法)：赛场给出 4-6 道独立题，每题都有唯一正确输出；你必须在限定时间里一次写对，系统才给满分

Heuristic (启发式)：赛场只给一大题，最优解本身没人知道，官方只公布评分函数；选手可以反复提交、调参、用贪心或模拟退火等花式办法把分数往上抬，最后按得分高低排名

这一场比赛，就是AWTF2025的Heuristic (启发式) 赛道的比赛，7月16号开始的，早上9点到晚上7点，一共比赛10个小时

赢者通吃 Winer Takes All?

我们需要转变思维，不仅仅将 AI 视为工具，而应将其视为一种新的经济主体/方式

可扩展知识：AI 能够**规模化**应用**知识**（而不仅是信息）

一种**全新的自动化形式**：与早期技术不同，能够像人类一样通过学习获取**隐形知识**/可编码知识
核心优势在于其**可扩展性**：人类专家的知识受限于个人时间，而 AI 的知识一旦被编码到模型中，就可以通过计算资源在整个经济中**大规模部署和复制**，**边际成本低**

AI vs. 传统机器人

AI更依赖数据和计算（而不是运行环境/规则约束），改进可以立即部署到所有硬件上，不再存在版本的不同。计算资源会被分配给最先进的AI 模型，而传统机器人往往以具有不同能力的老旧版本形式共存

比如智驾付费...

离谱！奔驰远程互联要收费，付费解锁真成“无底洞”？

向快乐出发 2025-04-01 11:16

想象一下，有一天你收到一条消息：“亲爱的用户，您的刹车使用次数即将用尽，请及时充值。”这听起来像是个玩笑，但它却预示着汽车行业未来的一个重要趋势——功能付费解锁。

不只是新兴的电动汽车品牌如特斯拉、蔚来和威马等在玩转这个功能，就连宝马、奔驰这样的传统豪华车品牌也加入了“付费订阅车辆功能”的行列。不久前，奔驰就推出了一项服务，用户只需支付4998元/年，就能解锁EQS车型的后轮转向功能，此举在市场上引起了不小的轰动。而最近，奔驰又推出了新的付费服务——车辆远程互联功能。



城南旧事

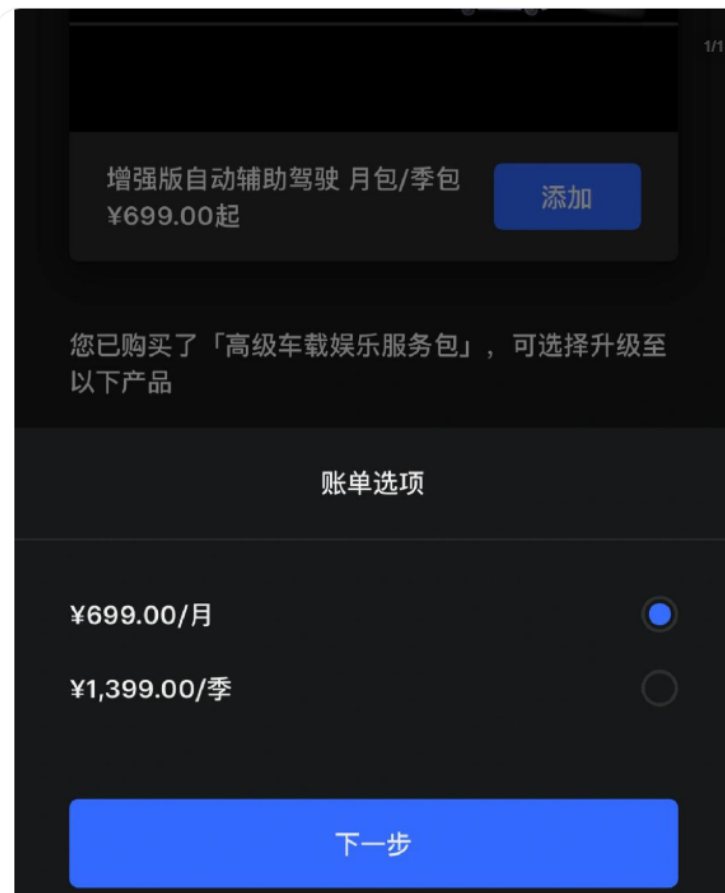
关注

699 元的EAP冲不冲？

今天特斯拉App更新了EAP（增强版自动驾驶）订阅功能，简单来说就是花点钱解锁一些高级驾驶辅助功能。具体包括：

1. ****高速根据导航自动驾驶****：这个功能说实话我体验一般，甚至免费送的时候我都直接关掉了，感觉用起来不太顺手。
2. ****AP打灯自动变道****：这个还不错，打转向灯就能自动变道，挺方便的，算是我觉得EAP里最实用的功能了。
3. ****智能召唤****：这个功能听起来很酷，但实际能用到的场景不多，感觉有点鸡肋。
4. ****自动泊车****：嗯.....这个功能我连免费版都懒得用，更别说花钱订阅了。总的来说，EAP的功能有点“食之无味，弃之可惜”的感觉，虽然69块钱不算贵，但要不要冲，我还真得好好想想。你呢？觉得值不值？

03-23



当年乐视没干成的事，智驾可否实现？

我们需要转变思维，不仅仅将 AI 视为工具，而应将其视为一种新的经济主体/方式

独立工作：AI 智能体/代理能够**独立**/自主工作

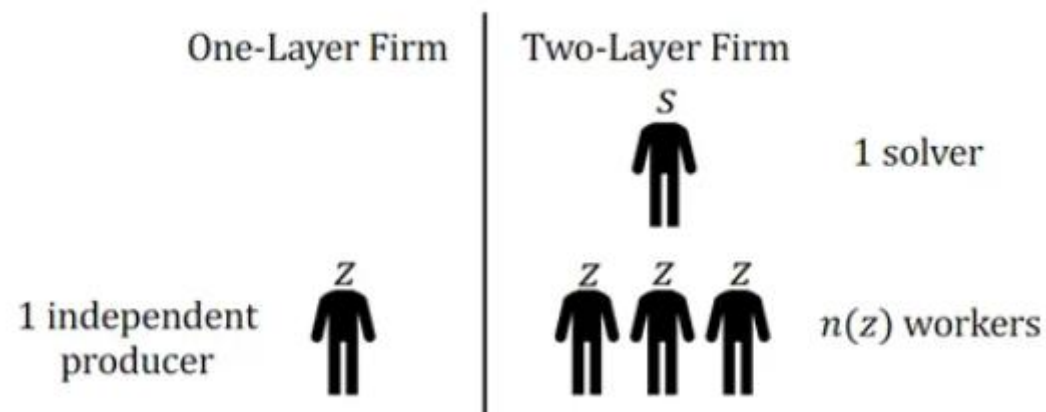
基于基础模型的能力，AI 正在发展为能够**自主执行**认知工作的 " 代理 "

AI vs. 信息技术 (ICT)

AI存在独立工作方式，而且不同独立 vs 辅助的部署方式，更需要注意不同方式下的成本效益
(并非所有自动化任务都具有好的ROI)

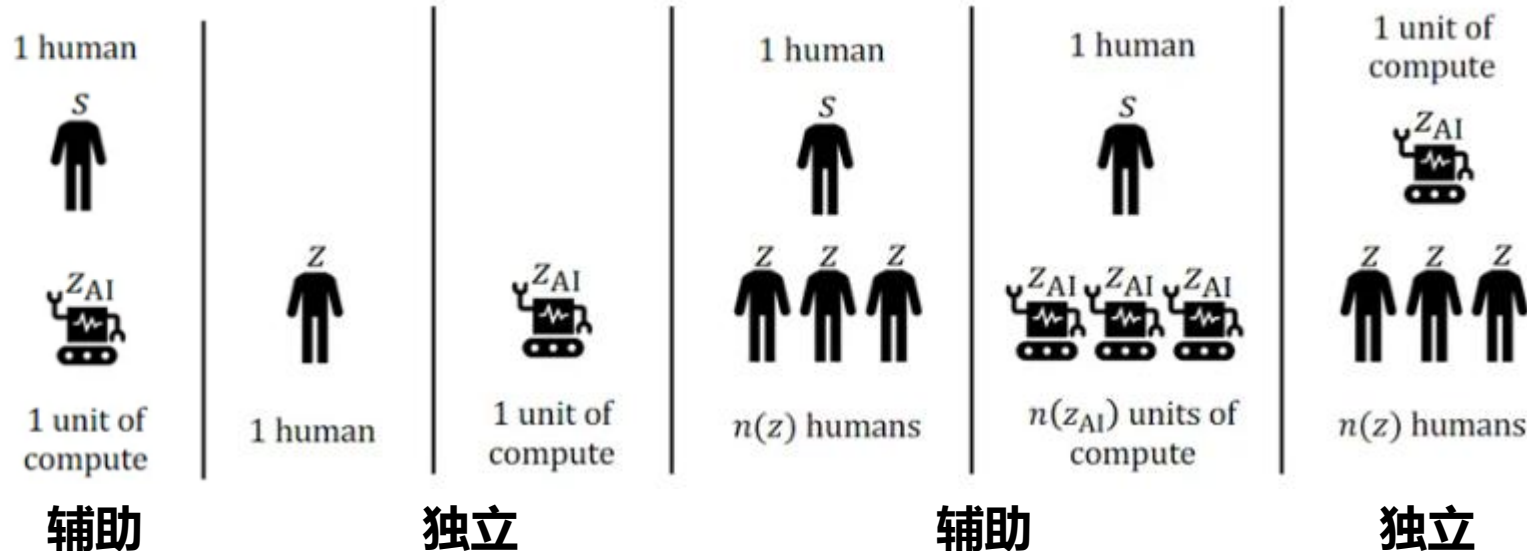
人机协同存在更多的挑战

传统管理



辅助和独立分别适合
哪种岗位/企业?

AI管理



我们需要转变思维，不仅仅将 AI 视为工具，而应将其视为一种新的经济主体/方式

学会如何不被忽悠？

如何识别你（投资）的团队是否在技术上真实可信，而非用模糊术语忽悠人，比如：

给我讲讲你如何改的注意力矩阵

你们为什么要改注意力，而不是改数据或训练目标？

强化学习如何用到训练里的

为什么选择用强化学习，主要解决了哪一个之前解决不了的问题

去掉最创新的那个模块，模型性能降低多少，在哪些任务降低

最不愿意让模型做的任务是什么？

如果上下文长度再翻一倍，架构可能最先崩的是哪里？

现在模型里，最不敢/能开源的是哪一部分，为什么？

我们还要操心企业里面如何用好AI：AI可以帮助更好的补货么？



Figure 3: A JRE Vending Machine Operation Worker



Kawaguchi, K. (2021). When will workers follow an algorithm? A field experiment with a retail business. *Management Science*, 67(3), 1670-1695.

实施结果

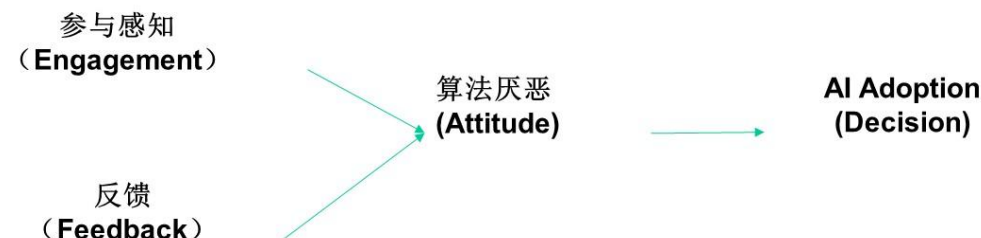
组1：将原有AI算法（未采纳员工的预测知识）结果告知员工；

组2A：将新AI 算法（采纳员工知识作为先验）结果告知，并告知算法融合了员工知识

组2B：将新AI 算法（采纳员工知识作为先验）结果告知，但并未告知算法融合了员工知识

组3：未告知员工算法结果

组1，2B，3结果无显著差异
组2A工作改善最高



做了几个 B 端 AI 项目的落地之后，感受如图所示。(算了，给钱就行...)



5小时前

...

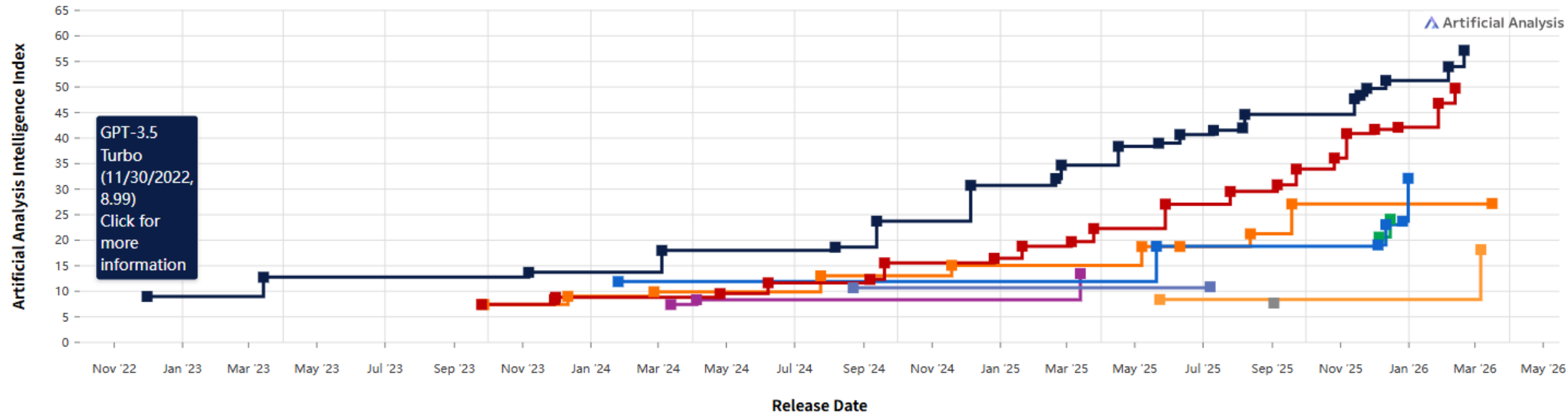
AI说是落地了，没起到什么用，空有一个AI的名头。比如现在很多家推RAG项目，有的效果做出来不好，还不如传统检索。

甚至，我们还要关心：国家之争

Frontier Language Model Intelligence By Country, Over Time

Artificial Analysis Intelligence Index v4.0 incorporates 10 evaluations: GDPval-AA, τ^2 -Bench Telecom, Terminal-Bench Hard, SciCode, AA-LCR, AA-Omniscience, IFBench, Humanity's Last Exam, GPQA Diamond, CritPt

■ United States ■ China ■ South Korea ■ United Arab Emirates ■ France ■ Israel ■ Canada ■ India ■ Switzerland

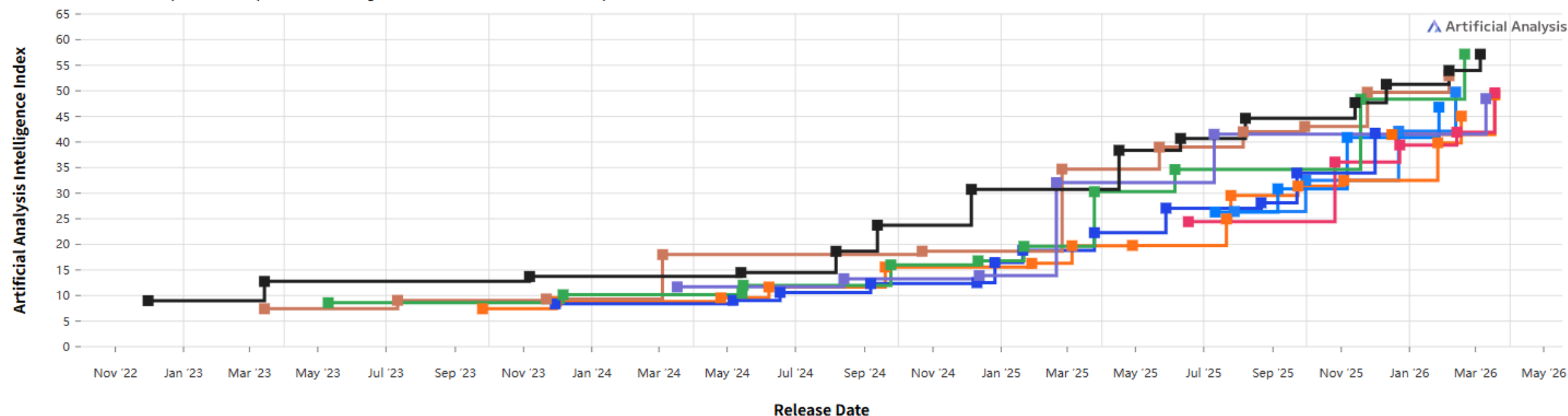


甚至，我们还要关心：模型之争

Frontier Language Model Intelligence, Over Time

Artificial Analysis Intelligence Index v4.0 incorporates 10 evaluations: GDPval-AA, τ^2 -Bench Telecom, Terminal-Bench Hard, SciCode, AA-LCR, AA-Omniscience, IFBench, Humanity's Last Exam, GPQA Diamond, CritPt

Alibaba Anthropic DeepSeek Google Kimi MiniMax OpenAI xAI Xiaomi Z AI



DeepSeek gives China's chipmakers leg up in race for cheaper AI

By Reuters

February 13, 2025 12:46 PM GMT+8 · Updated 3 days ago



[1/2] The Deepseek logo is seen in this illustration taken on January 29, 2025. REUTERS/Dado Ruvic/Illustration/File Photo [Purchase Licensing Rights](#)



Summary Companies

- DeepSeek models boost options for Chinese chipmakers like Huawei
- DeepSeek's open-source nature seen aiding AI adoption in China
- Nvidia still leads in AI chips despite Chinese advancements

DeepSeek-R1 已发布并开源，性能对标 OpenAI o1 正式版，在网页端、APP 和 API 全面上线，[点击查看详情](#)。

deepseek

探索未至之境

开始对话

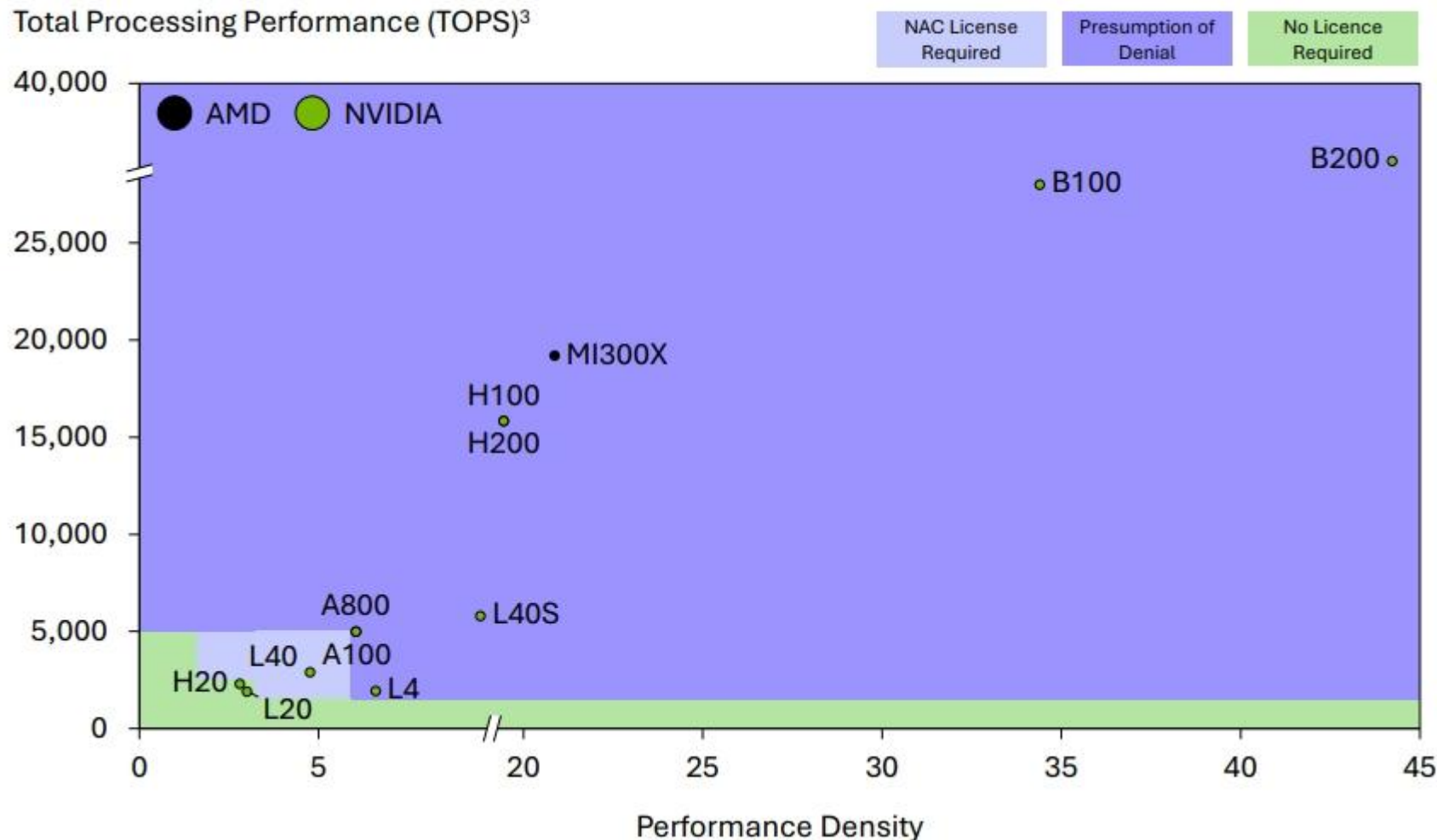
免费与 DeepSeek-V3 对话
使用全新旗舰模型

获取手机 App

DeepSeek 官方推出的免费 AI 助手
搜索写作阅读解题翻译工具

US Accelerators Prohibited for Export to China^{1,2}

Total Processing Performance (TOPS)³



Total Processing Performance (TPP): 每秒万亿次操作（Tera Operations per Second, TOPS）为单位

Performance Density: 性能密度评估单位物理尺寸上的性能表现，即TPP除以芯片的面积Die Size

即便如此

3月1日, DeepSeek官方通过社交媒体账号公布, 理论上成本利润率为545%

- 平均占用226.75*8的H800 GPU。如GPU租赁成本为2美金/小时, 总成本为\$87,072/天
- 输入+输出token总数为776B, 按R1定价, 一天收入为\$562,027, 成本利润率545%

如果我们把LLM看作产品 (而非技术)

- 性价比高的产品 vs. 高性能 (高成本) 产品

DeepSeek一体机

32B/70B/671B 满血版

开箱即用 本地部署 满足个人及部门需求

满血版劲爆价

8.98万/台



工作站一体机（单用户）

适配模型	型号	配置	性能参数	适用场景
32B模型	W559V2 单卡方案	CPU: AMD 7542*1 内存: 64G GPU: AMD 7900XTX 24G*1	单用户吞吐 25 tk/s	个人基础智能化办公场景
	W359V3 单卡方案	CPU: AMD 9124*1 内存: 64G GPU: 600W 32G*1	单用户吞吐 50 tk/s	个人高性能智能办公场景
70B模型	W559V2 双卡方案	CPU: AMD 7543*2 内存: 256G GPU: AMD 7900XTX 24G*2	单用户吞吐 12 tk/s	个人高性能智能办公场景
	W359V3 双卡方案	CPU: AMD 9124*1 内存: 256G GPU: 600W 32G*2	并发数 ≤ 15 单用户吞吐 28 tk/s	部门50人左右并发且参数需求较高的智能办公场景
671B满血版	W359V3 CPU+内存	D5-R1-671B-Q4_K_M CPU: AMD 9375F*1 内存: 96G 6400 DDR5*12	并发数 ≤ 5 单用户吞吐 12tk/s	科研建模、医疗影像分析等纯内存个人满血版DeepSeek一体机。可成为科研人员、新晋工作者及高安全需求用户的“AI工作伴”

服务器一体机（多用户）

适配模型	型号	配置	性能参数	适用场景
32B模型	U629V2	CPU: AMD 7543*2 内存: 256G GPU: NV 4090D 48G*2	并发数 130 吞吐量 2000 token	企业级智能客服、教育科研、开发运维等场景需求，实现高性能与低成本的平台
	U628V2	CPU: Intel 5318Y*2 内存: 256G GPU: NV 4090D 48G*2	并发数 90 吞吐量 1400 token	中小企业内部智能客服、内部知识库管理、自动化文档处理、数据分析报告生成等
70B模型	G558V3	CPU: Intel 5418Y*2 内存: 512G GPU: NV 4090D 48G*4	并发数 90 吞吐量 1400 token	中小企业内部大模型计算、多模态AI、自动编程、大规模金融和医疗分析等场景
	G659V2	CPU: AMD 7543*2 内存: 512G GPU: NV 4090D 48G*8	并发数 90 吞吐量 1300 token	中小企业内部大模型计算、多模态AI、自动编程、大规模金融和医疗分析等场景
	G659V3	CPU: AMD 9334*2 内存: 512G GPU: NV 4090D 48G*8	并发数 120 吞吐量 1800 token	中小企业内部大模型计算、多模态AI、自动编程、大规模金融和医疗分析等场景
	G658V3	CPU: Intel 5330*2 内存: 512G GPU: NV 4090D 48G*8	并发数 120 吞吐量 1800 token	中小企业内部大模型计算、多模态AI、自动编程、大规模金融和医疗分析等场景
	G678V3	CPU: Intel 5330*2 内存: 1T GPU: 600W 32G*8	并发数 180 吞吐量 2100 token	中大型企业内部的商业智能客户、自动化内容生产以及内部知识管理、同时支持数据分析和建模、帮助企业将海量数据转化为业务决策等全方位场景
	G679V3	CPU: AMD 9334*2 内存: 1T GPU: 600W 32G*8	并发数 180 吞吐量 2100 token	中大型企业内部的商业智能客户、自动化内容生产以及内部知识管理、同时支持数据分析和建模、帮助企业将海量数据转化为业务决策等全方位场景

DS推理模型的硬件需求

模型参数	CPU要求	内存要求	显存要求 (GPU)	硬盘空间	适用场景
1.5B	6核 (现代多核)	16GB	4GB (如:GTX 1650)	5GB+	实时聊天机器人、物联网设备
7B	8核 (现代多核)	32GB	8GB (如:RTX 3070)	10GB+	文本摘要、多轮对话系统
8B	10核 (多线程)	32GB	10GB	12GB+	高精度轻量级任务
14B	12核	64GB	16GB (如:RTX 4090)	20GB+	合同分析、论文辅助写作
32B	16核 (如i9/Ryzen 9)	128GB	24GB (如:RTX 4090)	30GB+	法律/医疗咨询、多模态预处理
70B	32核 (服务器级)	256GB	40GB (如:双A100)	100GB+	金融预测、大规模数据分析
671B	64核 (服务器集群)	512GB	160GB (8x A100)	500GB+	国家级AI研究、气候建模

FOMO (Fear of Missing Out, 怕错过)

“比肩” ChatGPT之父奥特曼，中国 “AI教父” 李一舟靠卖课年入5000万

中美两大AI巨头



据飞瓜数据显示，2023年李一舟售卖的199元AI课《每个人的人工智能课》，一年内卖出约25万套，销售额约5000万。而一张网络流传的截图显示，李一舟本人通过AI课程，在3年内收入超亿元。

据说在创业者中，90%都有抑郁症。这种说法虽然有些夸张，但在创业的高压之下，心理和情绪出问题的并不少见。

四个月前刚刚宣布要进军AI的王慧文⁺，不幸成为其中之一。

据美团在港交所的公告显示，因个人健康原因，王慧文不得不提出辞去公司非执行董事、董事会提名委员会成员和公司的授权代表证券上市规则职务的请求，而这一决定将从次日起正式生效。

在此之前，王慧文已经离开了他创办的光年之外⁺。随后在6月29日，美团发布公告，**宣布将以现金约2.33亿美元、债务承担约3.67亿元及现金1元，完成人工智能公司光年之外境内外主题100%股权的收购。**

四个月的时间，两个“老王”在AI赛道上完成了一次“二人转”。



王慧文

我的人工智能宣言：
5000万美元，
带资入组，
不在意岗位、薪资和 title，
求组队。

境内收购事项的代价包括，境内买方应向王慧文支付1元，光年之外债务3.67亿元。公告期内，境内光年之外（不包括一流科技）的未经审核账面值为3.48亿元。2023年4月，境内光年之外完成收购一流科技，截至2022年12月31日，一流科技的经审核账面值为1558.5千万元，资产净值为57万元。

目前几乎没有经营活动的光年之外账上仍有2.84亿美元，也就是说此次收购价几乎是按照净现金的平价收购。

备注：老王的投入算在债务里。此前王慧文的5000万美元带资进组实际上是以可换股债券形式，将钱借给了光年之外。而王慧文的退出方式，是得到366,924,000元的债权支付，换算成美元是5071万。



不怕AI超过人,只怕
会用AI的人超过我

数字/智能鸿沟

Deepseek 满血版 R1/V3

无需部署 打开即用
立刻响应

Deepseek 满血版

网页版 支持手机/电脑及APP

DeepSeek满血版R1不限次数

¥9.9 1000+人付款

满2件8.9折 24小时内发 包邮

店>

DeepSeek

实付 ¥18.9

查看明细>

1 2件8.9折, 一次买多款>

颜色分类 (6)

88 大图

手机端app版【不限次数永久使用】安装直接使用不需要云端部署 ¥9.9

电脑端app版【不限次数永久使用】安装直接使用不需要云端部署 ¥9.9

手机电脑端app版【不限次数永久使用】安装直接使用不需要云端部署 ¥18.9

手机电脑端网页版【不限次数永久使用】无需安装直接使用 ¥18.9

电脑端本地部署【满血版最新更新8B 14B 32B模型】有详细安装视频+使用指南+指... ¥19.9

软件版+电脑版+本地部署电脑端 ¥29.9

发了你四个链接

哪个好就用哪个

另外3个我之前就有了。。

已读

Kimi 你也有了?



CZ  BNB 
@cz_binance

X.com

Translated by Grok

号称安装龙虾后就啥也不用做了。

之后所有的时间都在调整那个啥也做不了的龙虾。

9:55 AM · 3/9/26 · **494K** Views



1.1K



329



2.3K



149



因此，我们还要走进技术：近年LLM发展之路

网站使用LLM服务

笔记本使用LLM服务

编程调用LLM

自动编程调用LLM

微信指挥自动编程调用LLM

Qwen3.5-千问

复旦的同学为什么这么优秀？

复旦大学的同学之所以表现出色，并非单一因素所致，而是源于一套贯穿招生、培养、评价和文化塑造全过程的系统性育人体系。这套体系旨在激发学生潜能，培养他们成为能够思考大问题、建构大视野、砥砺大情怀的未来领军人才。

🎯 选拔与培养：从源头注入卓越基因

复旦的优秀人才培养，从招生环节就已开始布局，并通过多元化的培养模式为不同类型的学生提供成长路径。

- 强基计划：通过“强基计划”招收对基础学科有浓厚兴趣和天赋的学生。这些学生遵循荣誉项目路径，修读更高难度的课程，并很早就能进入实验室参与科研。例如，计算与智能创新学院的大二学生谢志得，在大一暑假就已加入自然语言处理实验室，研究大模型效率问题。
- 荣誉项目：这是一个高挑战度的培养模式，将前沿科学问题融入教学，吸引拔尖学生参与。学生通过“闯关式”学习，不断突破自我，积累研究经验。化学系本科生张和杨就通过荣誉项目接触前沿研究，为他日后攻读博士打下了坚实基础。
- 卓博计划：针对有志于学术研究的顶尖本科生，提供本博贯通的个性化培养。通过选拔最优秀的学生，配备最优秀的导师，让他们提前接触研究生阶段的科研项目，培养国际一流的“高精尖缺”创新人才。

🌱 教育理念：打破边界，激发创新潜能



复旦的教育理念核心是打破传统学科壁垒，赋予学生极大的自主权，鼓励他

笔记本上使用

```
命令提示符 - ollama run deep x + v
C:\Users\zc_ch>ollama run deepseek-r1:1.5b
>>> 复旦的同学为什么这么优秀
Thinking...
嗯，用户问的是复旦的同学为什么那么优秀。首先，我得考虑复旦有哪些独特的优势和教育资源，比如地理位置、师资力量、课程设置这些。然后，可以提到复旦的历史深厚，拥有丰富的学府资源，这对学生的影响很大。

接下来，学术环境也不能忽视，复旦的图书馆资源丰富，学生可以根据自己的兴趣选择不同的方向。科研项目也是复旦特别之处之一，他们有很多国家级和省级课题，学生有机会参与研究，提升综合能力。

还有，复旦大学在社会服务方面也很突出，比如扶贫、教育公平，这些都为学生们提供了广泛的发展机会。此外，复旦的环境和气候条件优越，学生生活舒适，更容易投入学习中。

最后，结合当前的社会发展趋势，比如数字化转型和个性化发展，复旦的学生可能会更加适应未来，从而获得更好的发展机会。
...done thinking.

复旦作为我国顶尖的高等学府之一，拥有悠久的历史、丰富的人才资源和卓越的教学设施。以下几方面可以解释为什么复旦的同学会显得特别优秀：

### 1. **历史积淀与教育资源**
- 复旦是我国重要的高等学府，拥有悠久的历史传统和深厚的文化底蕴。学校在教育体系中占据重要地位，为学生提供优质的教育资源。
- 教师队伍 consists mainly来自国内顶尖高校，他们在学术研究、教学和人才培养方面具有较高的声誉。

### 2. **丰富的校园资源**
- 复旦有多条国家级科研项目和国家级课题，学生可以参与科研活动，提升综合能力。此外，学校还为学生提供了多种社团和课外活动平台，有助于学生全面发展。
- 学生还可以通过参加各类学术讲座、竞赛等学习机会，拓展视野并增强竞争力。

### 3. **国际化氛围**
- 复旦注重培养学生的国际视野，鼓励学生参与国际交流项目。同时，学校还积极为留学生提供奖学金和生活补贴，帮助他们克
```

编程调用LLM

```
import requests
import json
import re

def llmchat(question):
    host="http://localhost"
    port="11434"
    model='deepseek-r1'
    #model="deepseek-r1:1.5b"

    url = f"{host}:{port}/api/chat"

    headers = {"Content-Type": "application/json"}

    data = {
        "model": model,
        "options": {
            "temperature": 0.6
        },
        "stream": False,
        "messages": [{
            "role": "user",
            "content": prompt
        }]
    }

    response=requests.post(url,json=data,headers=headers,timeout=6000)
    answer=response.json().get("message").get('content')
    print(answer)
    print('\n')
    # print(json.dumps(response.json(), indent=2, ensure_ascii=False))

prompt = "复旦为什么这么优秀" #在这里输入问题
print('\n')
llmchat(prompt)
```

<think>

嗯，用户问“复旦为什么这么优秀”，这个问题看似简单但内涵很深啊。首先得判断用户的意图——可能是学生在择校阶段想了解高校优势，也可能是校友想确认母校的地位，甚至可能是普通网友对学术话题感兴趣。

用户用“这么优秀”这种口语化表达，说明ta可能不是专业研究者，而是希望得到通俗易懂的解释。但也不能太浅显，毕竟复旦确实值得深入分析。需要兼顾科普性和权威性，既要讲清楚客观事实，又要体现人文关怀。

从问题本身看，“为什么”背后藏着两个层面：一是具体优秀指标（学术、历史等），二是价值判断依据（国际认可度）。用户可能潜意识里想比较中国高校特色，或者寻找个人报考/工作的理由支撑。回复时得注意结构化呈现优势领域，但避免变成枯燥的罗列。

复旦的优势确实多到难以概括，不过核心可以归纳为三点：历史底蕴、师资力量和国际化程度。比如华中科技大学也有很强实力，但用户特意提到复旦，说明对人文社科或长三角地区更关注？回复里应该适当对比同类院校突出特色。

要注意平衡学术性和可读性。用具体数据支撑观点（如校友贡献占比），也要有生动案例（周强教授）。最后那段“软实力”总结很重要——把抽象优势转化为生活化的体验，比如提到上海的咖啡馆文化，能让用户产生共鸣。

突然想到用户可能隐含焦虑：为什么自己没考上复旦？所以结尾要强调优秀与个人努力有关，避免让回答变成招生广告。用校友网络和地域资源作为补充说明比较稳妥。

</think>

复旦大学之所以能在中国高校中脱颖而出，并享有国际声誉，其成功源于多方面因素的综合作用：

一、历史底蕴深厚

- **中国现代高等教育起源**：1905年创建于上海，前身“复旦公学”是中国最早按现代模式建立的大学之一。作为中国最早的私立高等学府之一（与南洋公学同根同源），其百年积淀形成了独特的学术传统。
- **抗战时期贡献**：在西南联合大学时期（1937-1946年）参与重要科研工作，如“人工降雨”研究，展现了知识分子的家国情怀。

自动编程调用LLM

写一个程序：通过本地安装的ollama调用 deepseek-r1:1.5b，回答：复旦的同学为什么这么优秀

SOLO Coder

> 思考过程

我来帮你写一个调用 Ollama 和 deepseek-r1:1.5b 模型来回答问题的程序。让我先了解一下当前项目结构。

camera-mcp

Thought

camera-mcp.py

requirements.txt

Thought

我来为你创建一个调用 Ollama 的 deepseek-r1:1.5b 模型来回答问题的程序。

ollama_deepseek_qa.py ollama_deepse... +59 -0 查看变更

Thought

程序已创建，现在让我运行它来测试：

camera-mcp 在沙箱中 在终端查看

```
$ cd d:\camera-mcp; python ollama_deepseek_qa.py
```

Thought

程序已成功创建并运行！

回答结果

程序通过 Ollama 调用 **deepseek-r1:1.5b** 模型，回答了“复旦的同学为什么这么优秀”这个问题。模型给出的回答涵盖了多个方

ollama_deepseek_qa.py > ...

变更已完成，请确认是否采纳。

```
1+ """
2+ Ollama DeepSeek-R1:1.5b 问答程序
3+ 使用本地安装的 Ollama 调用 deepseek-r1:1.5b 模型回答问题
4+ """
5+
6+ from openai import OpenAI
7+ import os
8+
9+ def ask_ollama(question: str, model: str = "deepseek-r1:1.5b") -> str:
10+     """
11+     通过 Ollama 调用大模型回答问题
12+
13+     Args:
14+         question: 要回答的问题
15+         model: 使用的模型名称，默认为 deepseek-r1:1.5b
16+
17+     Returns:
18+         模型的回答
19+     """
20+     client = OpenAI(
21+         base_url="http://localhost:11434/v1",
22+         api_key="ollama"
23+     )
24+
25+     response = client.chat.completions.create(
26+         model=model,
27+         messages=[
28+             {
29+                 "role": "user",
30+                 "content": question
31+             }
32+         ]
33+     )
```

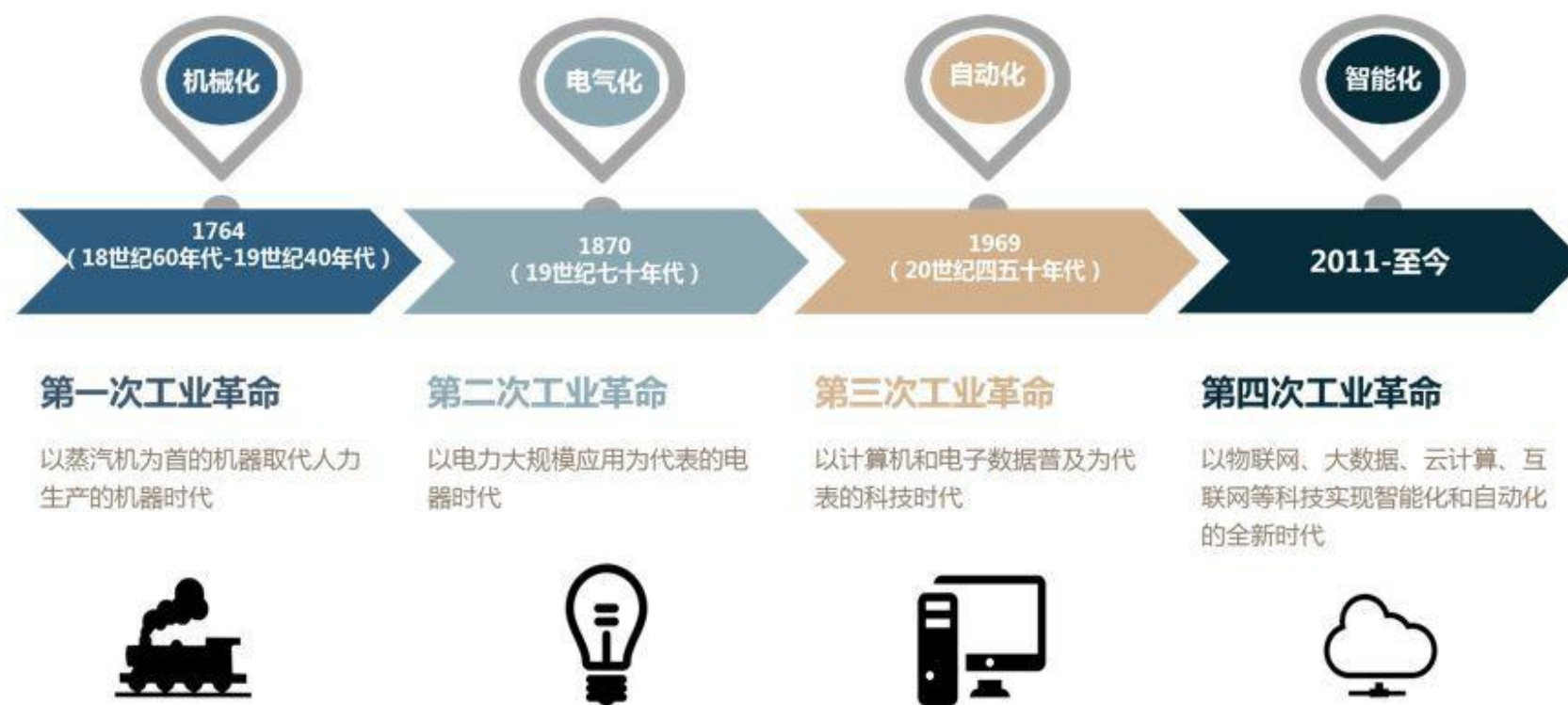
微信指挥自动编程调用LLM



让我们正式开始

对技术的管理和投资首先需要了解技术

通用人工智能



历次技术突破只是人类智能的产物，而唯独通用人工智能是“**智能”本身的革命**”

人工智能

- 技术：规则和逻辑推理的系统
- 应用：基于规则的医疗诊断系统

机器学习

- 技术：决策树、反向传播神经网络、支持向量机、协同过滤
- 应用：信用评分、推荐系统

深度学习

- 技术：AlexNet (2012)、卷积神经网络，循环神经网络
- 应用：人脸识别等

生成式人工智能

- GAN、Diffusion、Transformer...

1950s

1980s

2010s

2015s

2022.12

2024.9

2025.1

人工智能的早期探索

感知智能

认知智能

推理智能

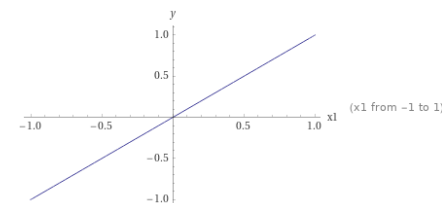
- 1. 如何理解人工智能及算法
- 2. AI/大模型技术的应用要点
- 3. 以及，如何不被忽悠...

理解算法： $Y = F(x)$

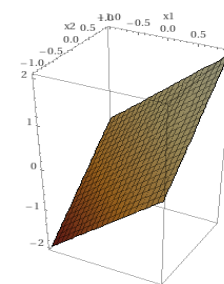
构建Y和X之间的F(.)

理解算法： $Y=F(x)$

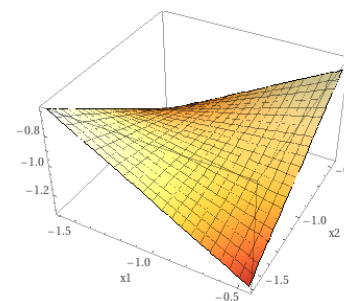
$$Y=ax_1$$



$$Y=ax_1+bx_2$$



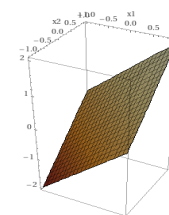
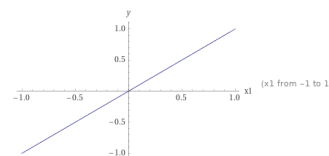
$$Y= ax_1+bx_2+cx_1x_2$$



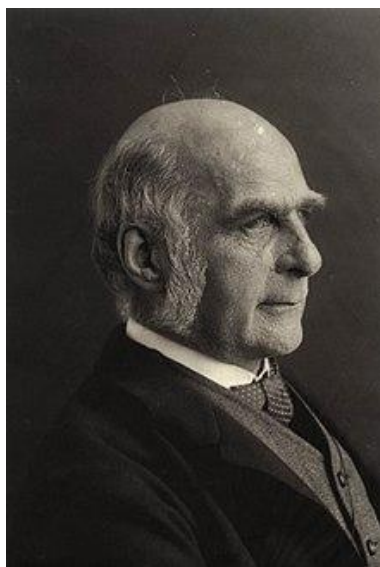
回归

- $y = ax + b$
 - 在可以观测到一组y和x的时候，我们可否得到a，b系数？
- 如果可以，这样的方法可否用在更一般的函数上，比如：

$$y = \sum_{i=1}^k \alpha_i x_i + b + \varepsilon = \alpha_1 x_1 + \alpha_2 x_2 + \cdots + b + \varepsilon$$

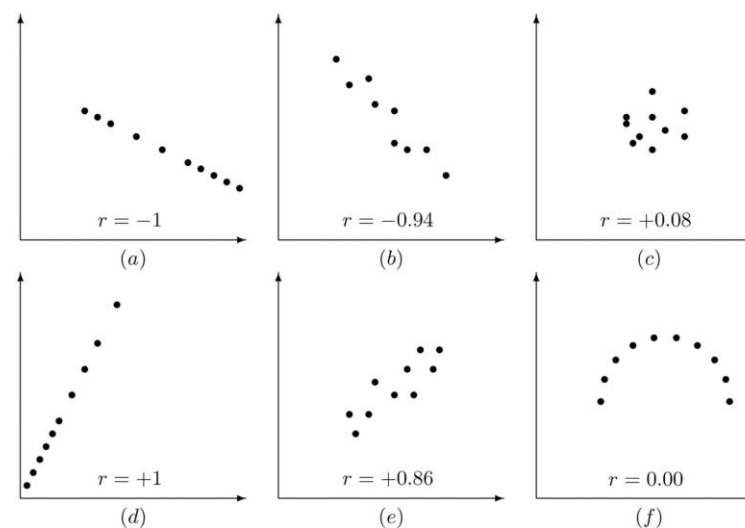


统计学



Sir Francis Galton
1822-1911

高尔顿爵士分别在1885年和1888年提出了回归和相关的概念，研究遗传和变异的问题，如父母的身高和孩子的身高的关系，身高和体重的关系等。



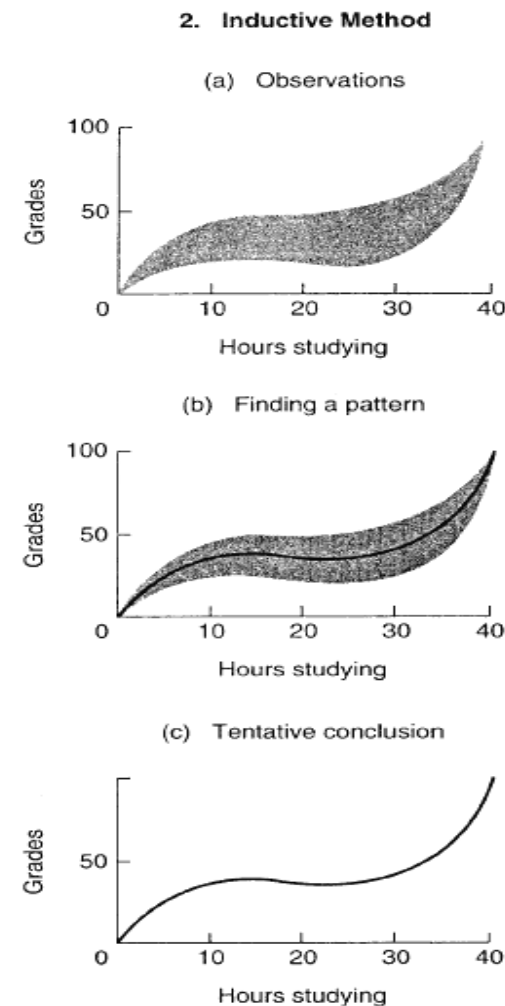
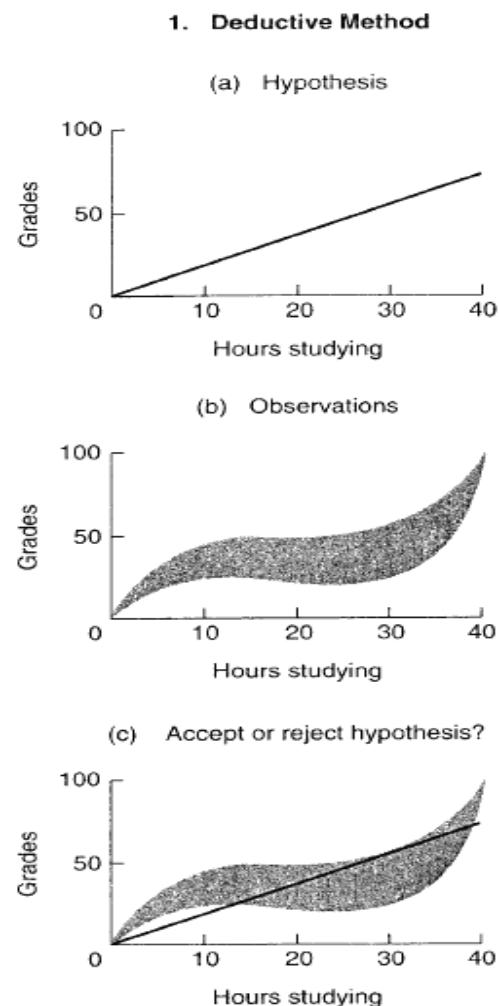
面对复杂的社会,我们如何建立已有知识可能还不足以指引的模型?

$$Y_{ijkt} = F(X_{ij}, X_{ik}, X_{it}, X_{ijt}, \dots)$$

i人j单位k地t时间

Figure 2-2 Deductive and Inductive Methods

演绎：
从一般规律推到具体现象



归纳：
从具体现象推到一般规律

AI本质: 根据输入 (X) 和输出 (Y) 找出两者的关系 (函数 $F()$)

$$Y = F(X)$$

统计回归

OLS/Logistic回归

.....

机器学习

分类 (决策树, 贝叶斯)

聚类 (K-mean)

关联分析 (MBA)

链接挖掘 (Pagerank)

神经网络 (NN)

相比统计先假设函数 (的样子) 再用数据去验证 (称之为假设检验), 计算机是生成/穷举最符合数据之间关系的函数

AI工作逻辑

1. 定义问题及对应的数据 (Y)

通过企业表现预测市场涨跌 (分类或者连续)

通过人/机器/企业运营行为识别风险

.....

2. 定义数据体现的特征, 和数据之间的关系 (X和F(X), F()的建立就是机器学习的过程)

用户生理心理特征, 社会变量

企业行动, 绩效指标, 产业信息, 投资者行为, 市场信息.....

常用AI算法

链接挖掘 (谷歌搜索, 网红估价)

Pagerank (流量密码)

分类 (决策预测, 风控信贷, 运营优化)

决策树C4.5/CART/Adaboost/RF/GBDT

神经网络NN/CNN/RNN/GPT: 深度学习/图片视频声音
强化学习

朴素贝叶斯Naive Bayes (贝叶斯网络)

矢量支持机SVM

聚类 (用户画像, 市场细分)

K-mean (分群分类)

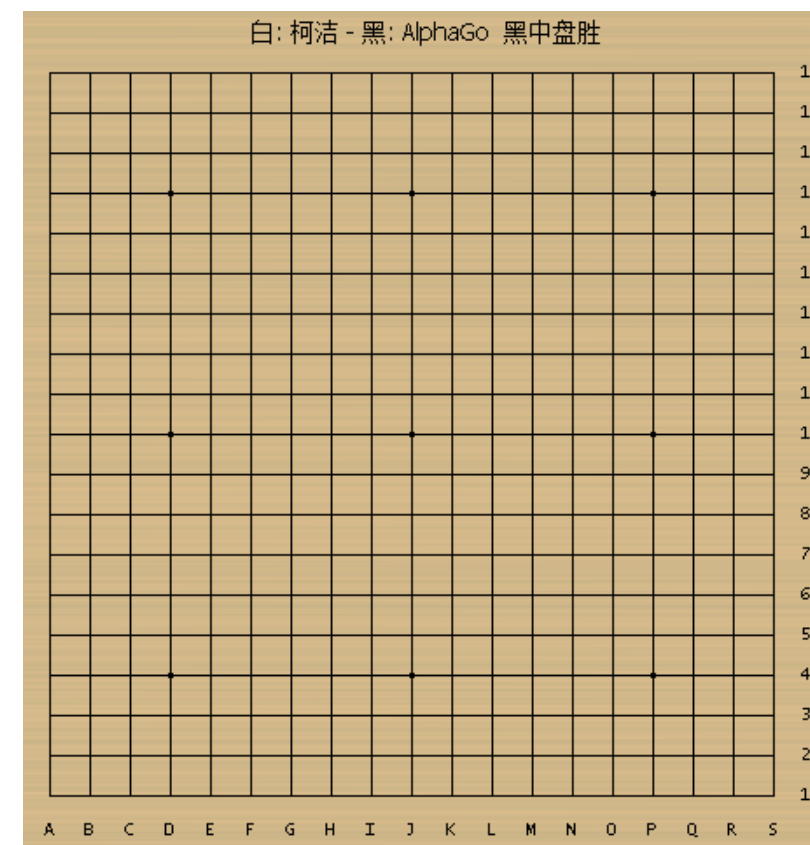
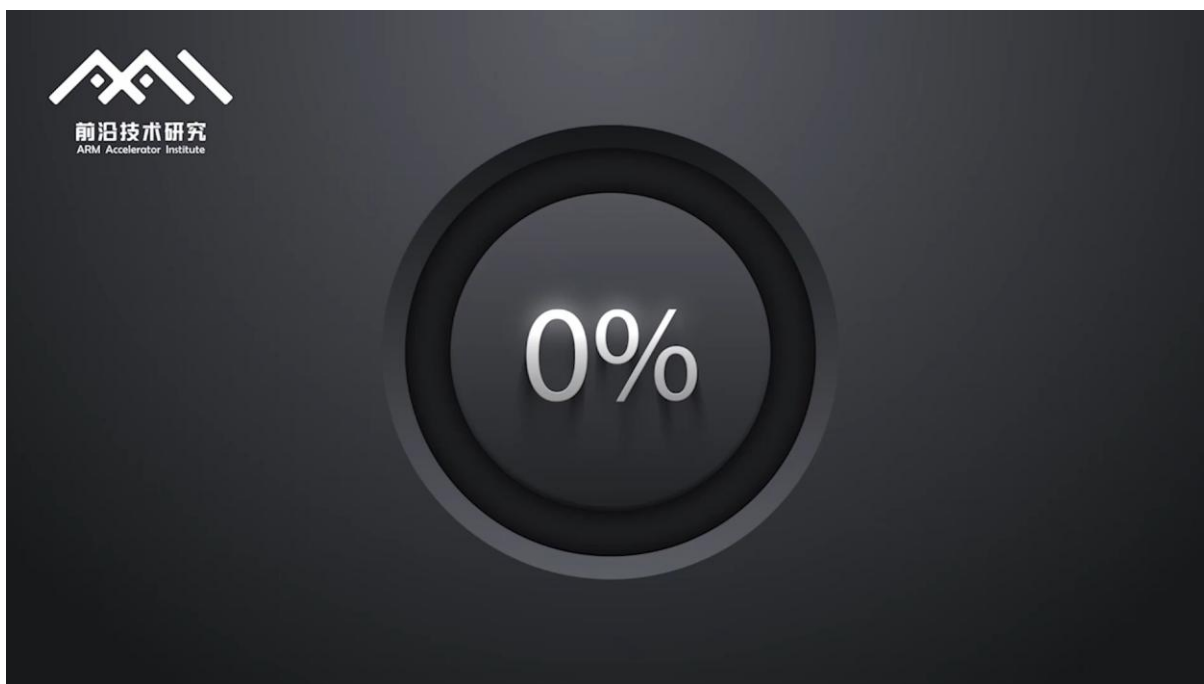
K-近邻(KNN, 聚类思路和分类的结合)

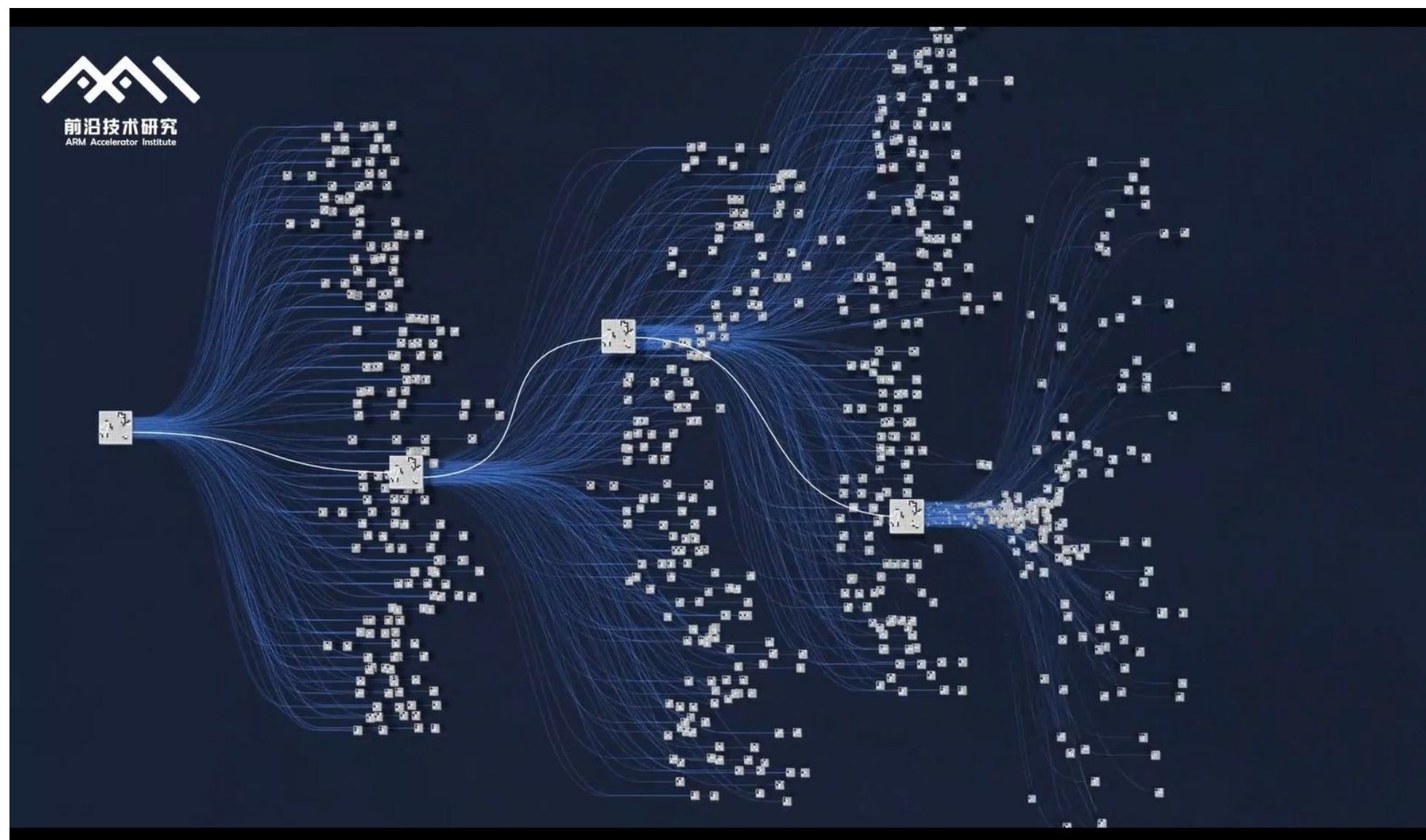
关联分析 (产品推荐)

关联规则Apriori (购物篮分析)

第一步：从现象中定义问题

难不难？



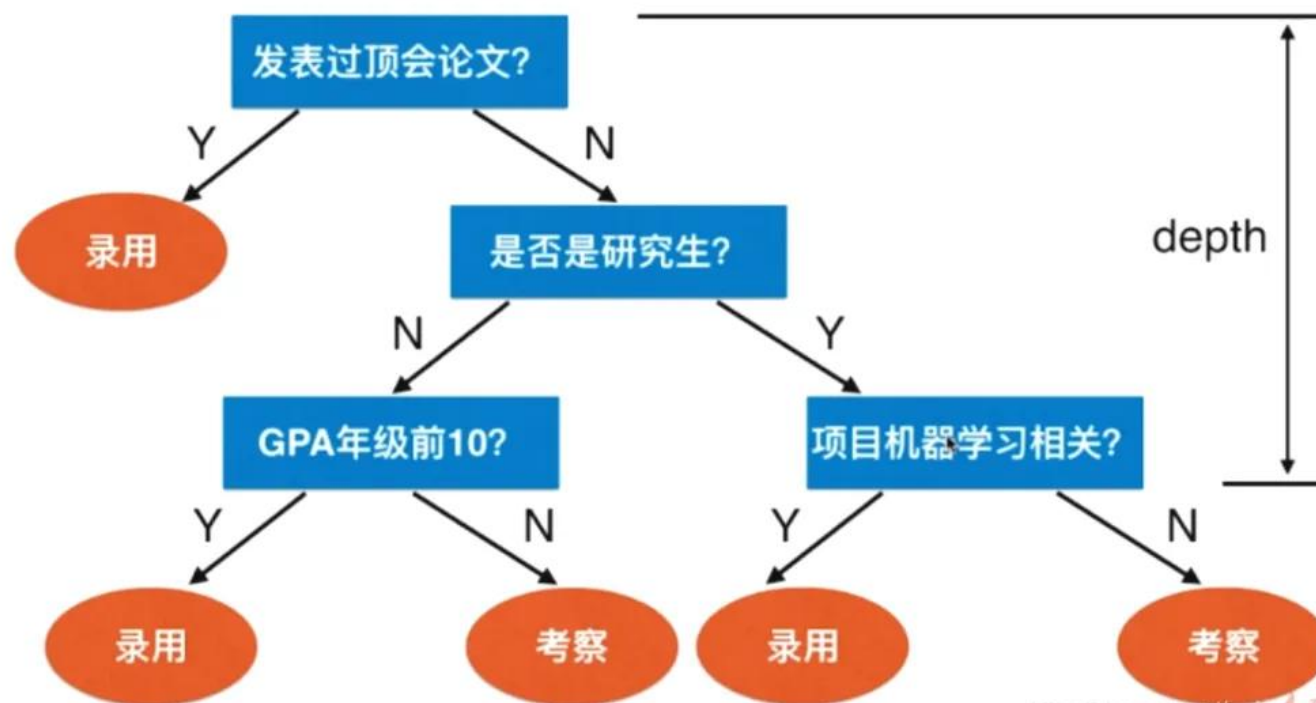


要不要他=发表+发表*研究生*(GPA+项目经验)

vx:15086893362

什么是决策树

招聘机器学习
算法工程师

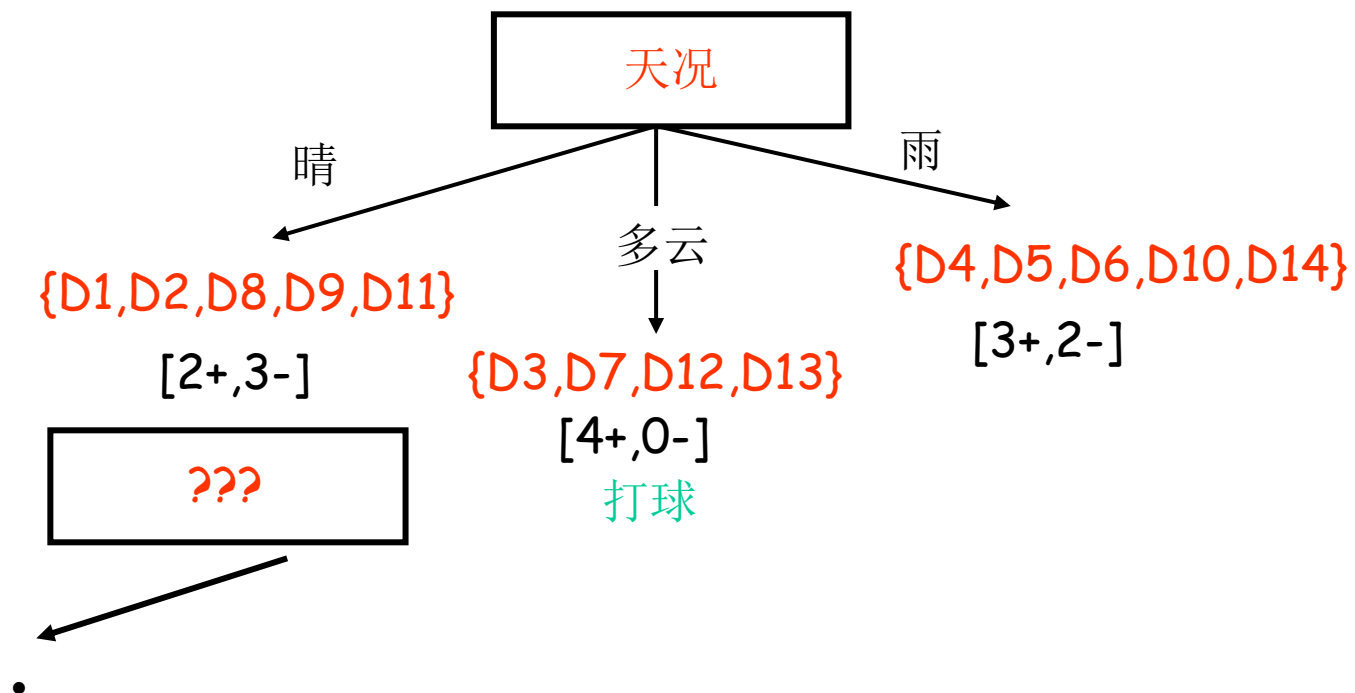


https://blog.csdn.net/weixin_41864004 课程网

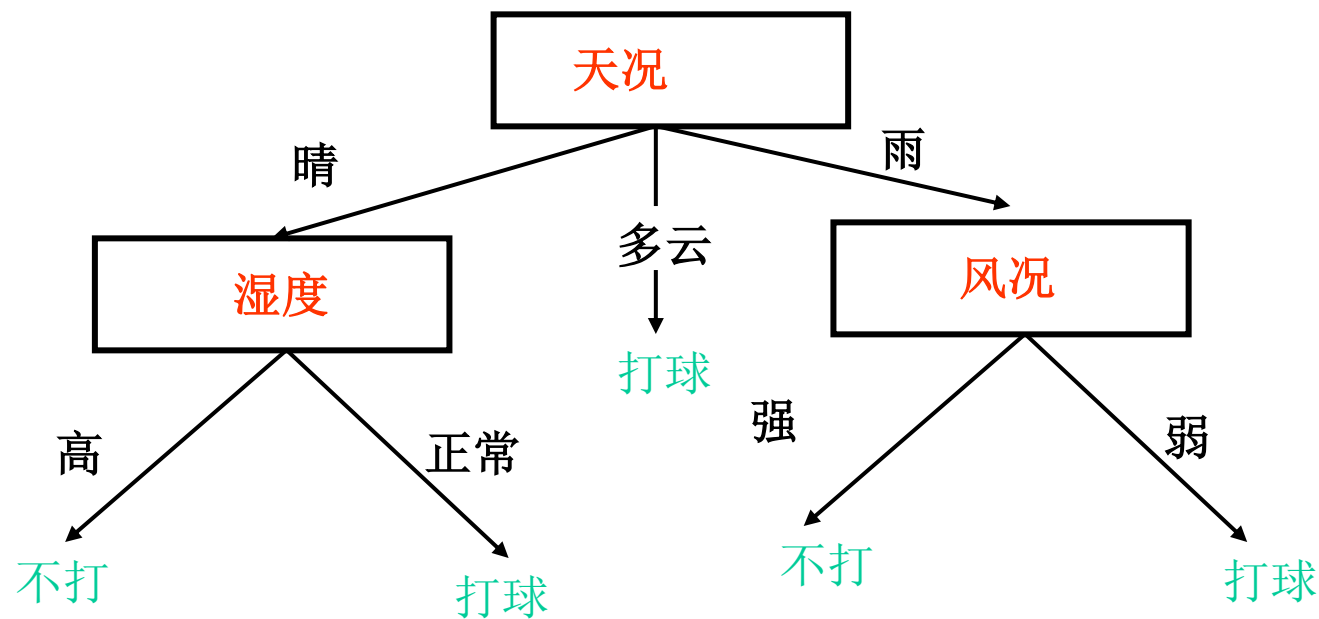
也许定义规则并不难，那么如何从数据中寻找规则呢？

做好客户服务

天	属性				打球
	天况	温度	湿度	风况	
1	晴	热	大	无	否
2	晴	热	大	有	否
3	多云	热	大	无	是
4	雨	中	大	无	是
5	雨	冷	正常	无	是
6	雨	冷	正常	有	否
7	多云	冷	正常	有	是
8	晴	中	大	无	否
9	晴	冷	正常	无	是
10	雨	中	正常	无	是
11	晴	中	正常	有	是
12	多云	中	大	有	是
13	多云	热	正常	无	是
14	雨	中	大	有	否

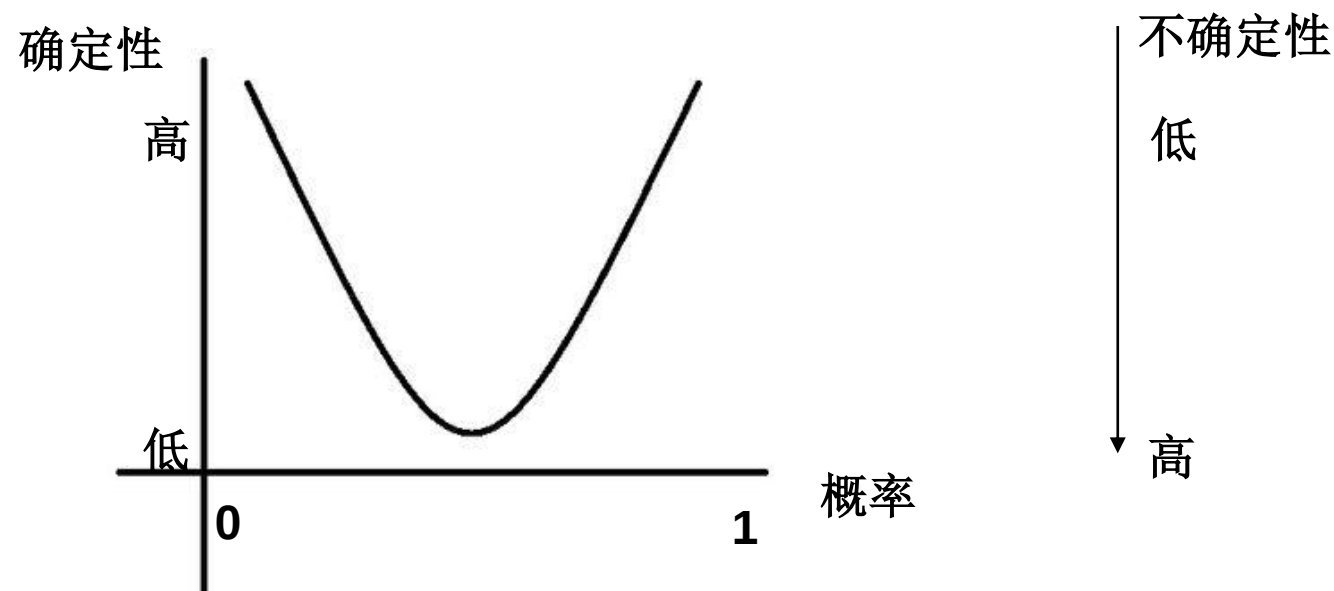


决策树



$(\text{天况} = \text{晴} \wedge \text{湿度} = \text{正常})$
 $\vee (\text{天况} = \text{多云})$
 $\vee (\text{天况} = \text{雨} \wedge \text{风况} = \text{弱})$
 打网球 = 是

问题的本质：（不）确定性
即事物的发展多大概率（不）接近于0/1

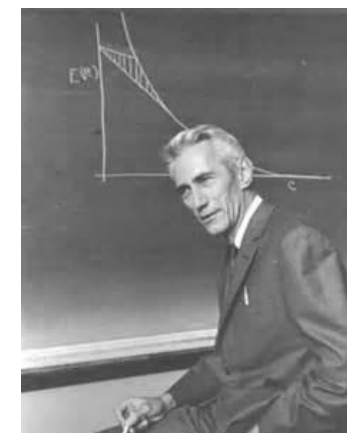


事物行为的不确定性

信息的基本作用就是消除决策的不确定性

信息熵（entropy）：不确定性

因此，获得更多信息的价值在于多大程度降低不确定性（即信息熵）



1916-2001

信息熵 $p_1 \log(p_1) + p_2 \log(p_2)$

$$- \sum_{k=1}^K p_k \log p_k$$

比如：观察A对同一件事物的6次选择，他6次做出了赞同的决定，3次做出了反对的决定

•那么不确定性可以用四则运算方式简单计算为：

• - $\left[(3/6) \log_2 (3/6) + (3/6) \log_2 (3/6) \right] = 1$

•可以看出，如果6次都是同1结果，熵为0（没有不确定）；如果两个选择一半一半，熵为1（完全不确定）（非负性）

•这个计算结果是连续变化的，不会出现重合（单调性）

•P可以不只两个，加法可以随意调换顺序而不改变结果（加法交换律）

$$\lim_{x \rightarrow 0} x * \log_2 x = 0$$

$$\log_2 0 = -\infty$$

$$0 * \log_2 0 = 0$$

举例计算

假设一个条件下分别有正负样本若干，如何计算该条件下的不确定性呢？

C1	0
C2	6

$$P(C1) = 0/6 = 0 \quad P(C2) = 6/6 = 1$$

$$\text{Entropy} = -0 \log 0 - 1 \log 1 = -0 - 0 = 0$$

C1	1
C2	5

$$P(C1) = 1/6 \quad P(C2) = 5/6$$

$$\text{Entropy} = - (1/6) \log_2 (1/6) - (5/6) \log_2 (5/6) = 0.65$$

C1	2
C2	4

$$P(C1) = 2/6 \quad P(C2) = 4/6$$

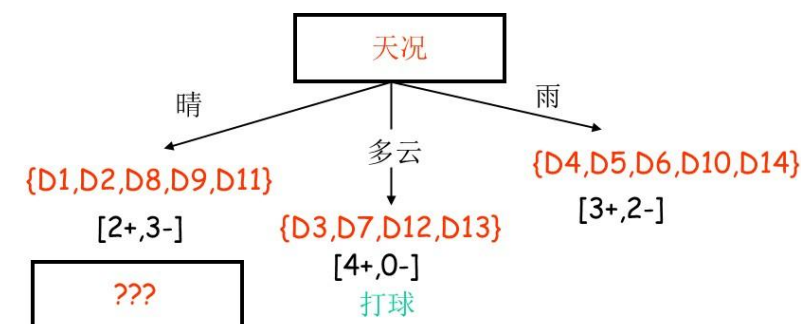
$$\text{Entropy} = - (2/6) \log_2 (2/6) - (4/6) \log_2 (4/6) = 0.92$$

$$\lim_{x \rightarrow 0} x * \log_2 x = 0$$

$$\log_2 0 = -\infty$$

$$0 * \log_2 0 = 0$$

新的规则能消减多大的不确定性？称之为增益



方法：用决策树某点与其分枝的不确定结果相减（初始为1，表示完全不确定）

值越大就说明消减信息熵越多，信息价值越大，称为信息增益 Gain， S_v 表示S的分枝

$$\text{增益}(S,A) \equiv \text{熵}(S) - \sum_{v \in \text{Values}(A)} |S_v|/|S| \text{ 熵}(S)$$

比如Outlook=Overcast时：entropy=1*log(1)+0*log(0)=0, Gain=1-0=1

$S_{\text{sunny}} = \{D1, D2, D8, D9, D11\} = \{2+, 3-\}$

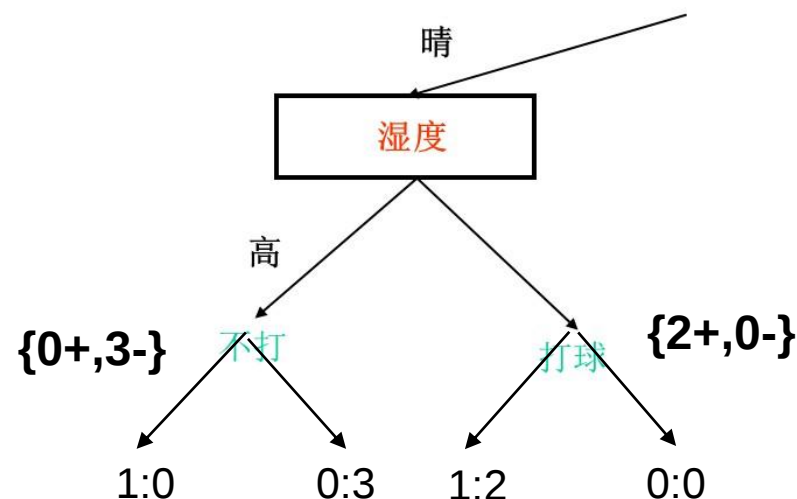
此时的熵为 0.97

即 $\text{Entropy}(S) = - \left[\left(\frac{2}{5} \right) \log_2 \left(\frac{2}{5} \right) + \left(\frac{3}{5} \right) \log_2 \left(\frac{3}{5} \right) \right] = 0.970$

假设我们用湿度去尝试

天	属性				打球
	天况	温度	湿度	风况	
1	晴	热	大	无	否
2	晴	热	大	有	否
3	多云	热	大	无	是
4	雨	中	大	无	是
5	雨	冷	正常	无	是
6	雨	冷	正常	有	否
7	多云	冷	正常	有	是
8	晴	中	大	无	否
9	晴	冷	正常	无	是
10	雨	中	正常	无	是
11	晴	中	正常	有	是
12	多云	中	大	有	是
13	多云	热	正常	无	是
14	雨	中	大	有	否

$$\text{Gain}(S, H) = \text{Entropy}(S) - \{3/5 * \text{Entropy}(S, H=\text{High}) + 2/5 * \text{Entropy}(S, H=\text{Normal})\}$$



$$-\sum_{k=1}^K p_k \log p_k$$

$$\begin{aligned} \text{Entropy}(S, H=\text{High}) &= -0/3 * \log(0) - 3/3 * \log(1) = 0 \\ \text{Entropy}(S, H=\text{Norm}) &= -2/2 * \log(1) - 0/2 * \log(0) = 0 \end{aligned}$$

Gain(Sunny, Humidity)

$$= \text{Entropy}(S) - \{3/5 * \text{Entropy}(S, \text{Humidity}=\text{High}) + 2/5 * \text{Entropy}(S, \text{Humidity}=\text{Normal})\}$$

$$= 0.970 - (3/5)0.0 - (2/5)0.0$$

$$= 0.970$$

假设我们用气温和风向为例去尝试

$$\begin{aligned} \text{Gain}(S, \text{Temperature}) &= 0.970 - (2/5)0.0 - (2/5)1.0 - (1/5)0.0 \\ &= 0.970 - (2/5)\{-1*\log(1) - 0*\log(0)\} - (2/5)\{-0.5*\log(0.5) - 0.5*\log(0.5)\} - (1/5)\{-1*\log(1) - 0*\log(0)\} \\ &= 0.570 \end{aligned}$$

- 问题出在Temperature=Mild时的D8, D11

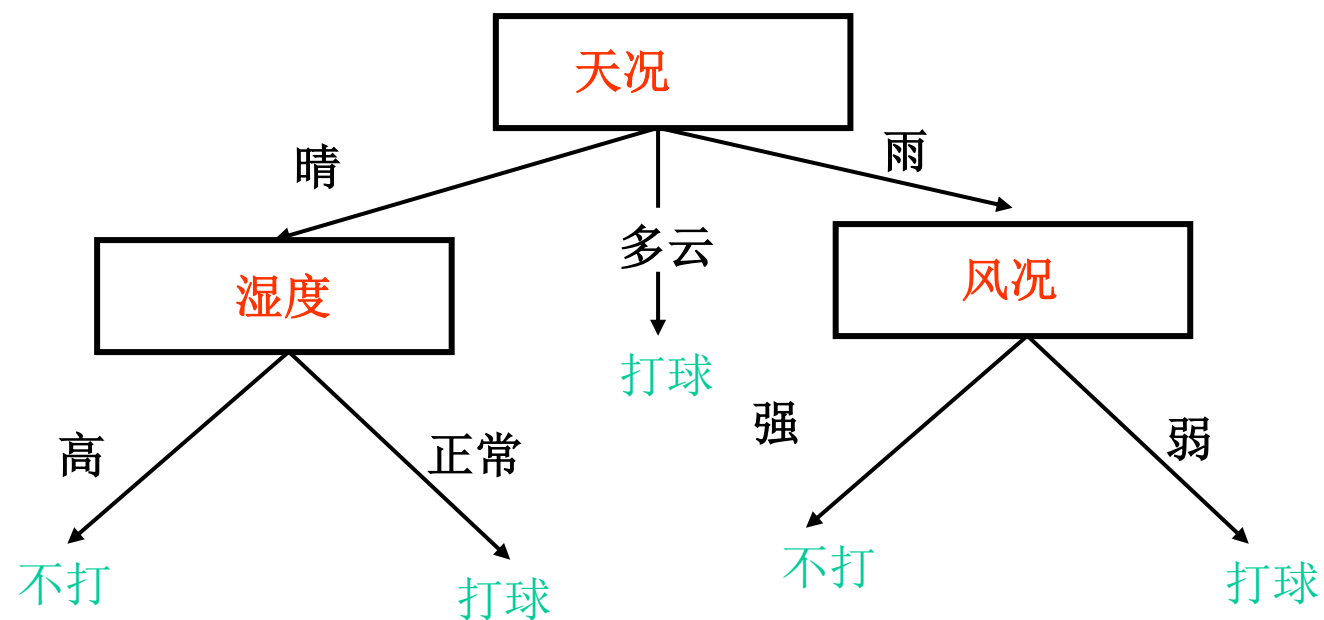
$$\begin{aligned} \text{Gain}(S, \text{Wind}) &= 0.970 - (2/5)1.0 - (3/5)0.915 = 0.019 \\ &= 0.970 - (2/5)\{-0.5*\log(0.5) - 0.5*\log(0.5)\} - (3/5)\{-2/3*\log(2/3) - 1/3*\log(1/3)\} \end{aligned}$$

$$\text{Gain}(S, \text{Humidity}) = 0.970 - (3/5)0.0 - (2/5)0.0 = 0.970$$

价值最大的是湿度，计算机用湿度作为晴天的下一个规则

天	属性				打球
	天况	温度	湿度	风况	
1	晴	热	大	无	否
2	晴	热	大	有	否
3	多云	热	大	无	是
4	雨	中	大	无	是
5	雨	冷	正常	无	是
6	雨	冷	正常	有	否
7	多云	冷	正常	有	是
8	晴	中	大	无	否
9	晴	冷	正常	无	是
10	雨	中	正常	无	是
11	晴	中	正常	有	是
12	多云	中	大	有	是
13	多云	热	正常	无	是
14	雨	中	大	有	否

决策树



$(\text{天况} = \text{晴} \wedge \text{湿度} = \text{正常})$
 $\vee (\text{天况} = \text{多云})$
 $\vee (\text{天况} = \text{雨} \wedge \text{风况} = \text{弱})$
 打网球 = 是

通信的数学理论，1948

如何测量信息

仅用0/1就可以表示无限的可能，使得二进制成为当代计算机的基础

这项发现也给人类带来了一个新的单词——比特 (bit)

B->KB->MB->GB->TB->PB->EB->ZB->YB->BB->NB->DB (2^{110} B)

Byte, Kilobyte, Megabyte, Gigabyte, Terabyte, Petabyte, Exabyte, Zettabyte, Yottabyte, Brontobyte, NonaByte, DoggaByte

A mathematical theory of communication

CE Shannon - Bell system technical journal, 1948 - Wiley Online Library

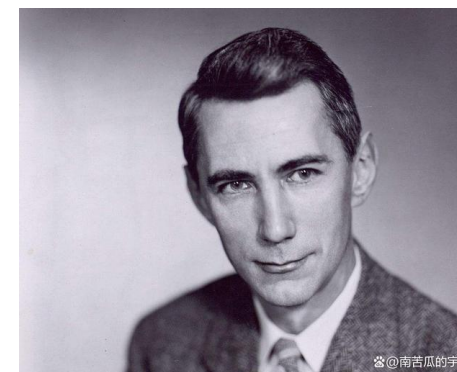
HE recent development of various methods of modulation such as PCM and PPM which exchange bandwidth for signal-to-noise ratio has intensified the interest in a general theory of communication. A basis for such a theory is contained in the important papers of Nyquist1 ...

☆ 99 Cited by 69155 Related articles All 143 versions »

[CITATION] The mathematical theory of communication

CE Shannon, W Weaver - 1998 - University of Illinois press

☆ 99 Cited by 45663 Related articles All 372 versions »



1916-2001

和财务决策树有什么区别？

分支有明确的先后顺序 vs 不知道顺序

背后是对商业逻辑的了解

相当于知道函数和不知道函数

数据量的差异

具体案例 vs 大数据

目标的清晰

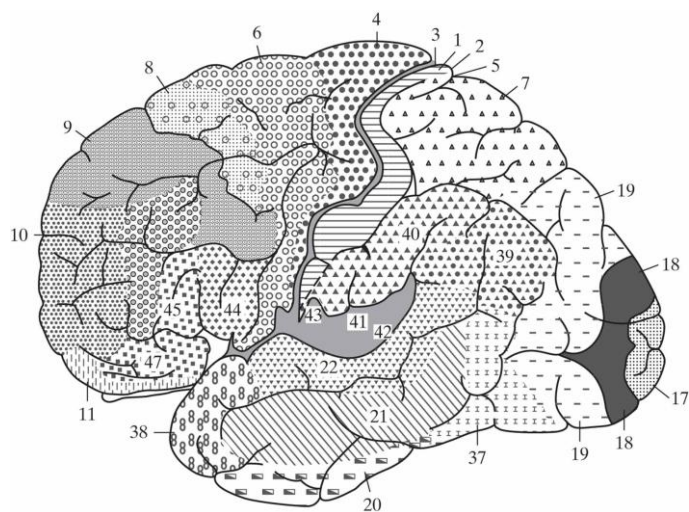
明确的ROI vs 明确的技术指标（准确率）

有没有这样的商业决策场景

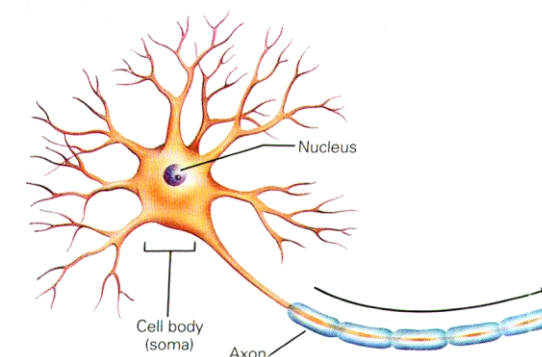
可以用ML替代人工规则（决策树）？

决策树、贝叶斯学习、向量支持机（小数据）、**神经网络...**

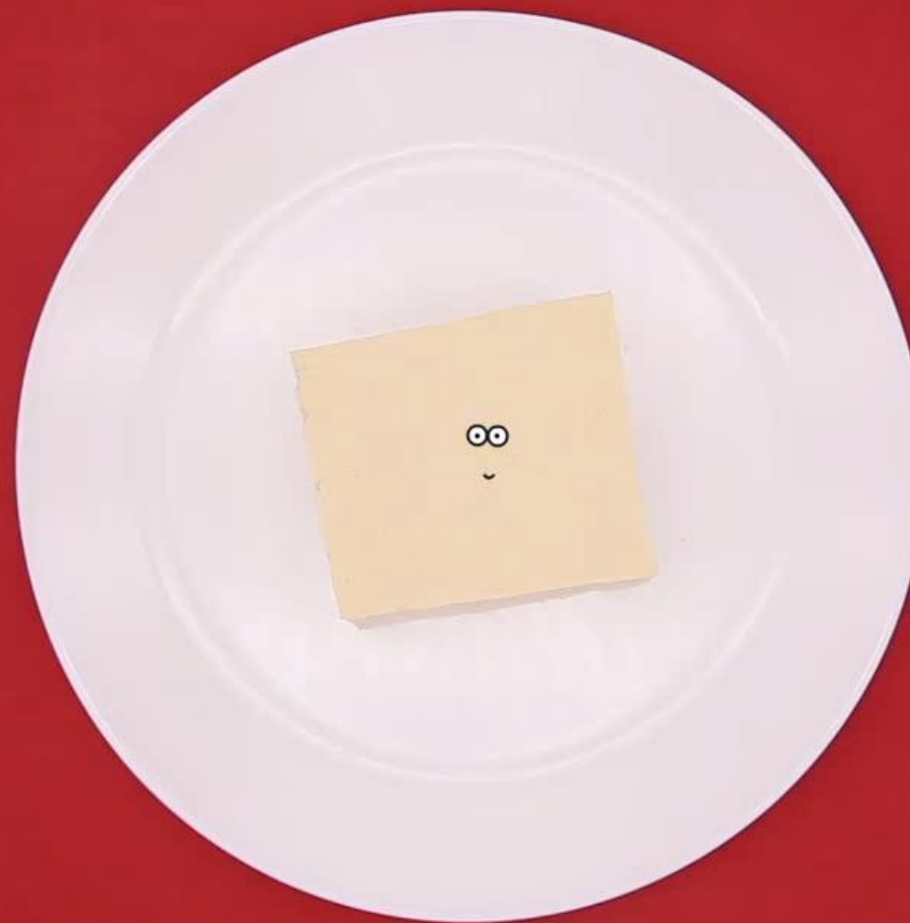
最依赖数据量和计算量的方法：神经网络/深度学习



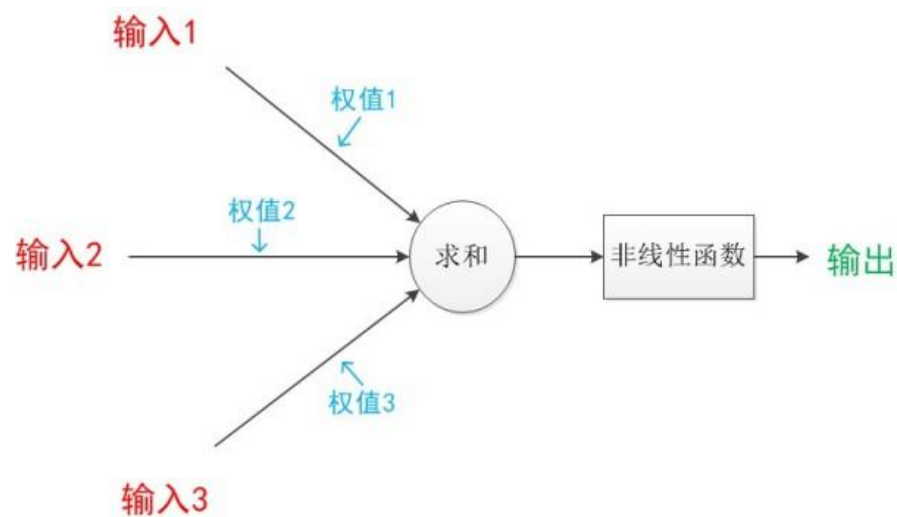
This image by Fotis Bobolias is licensed under CC-BY 2.0



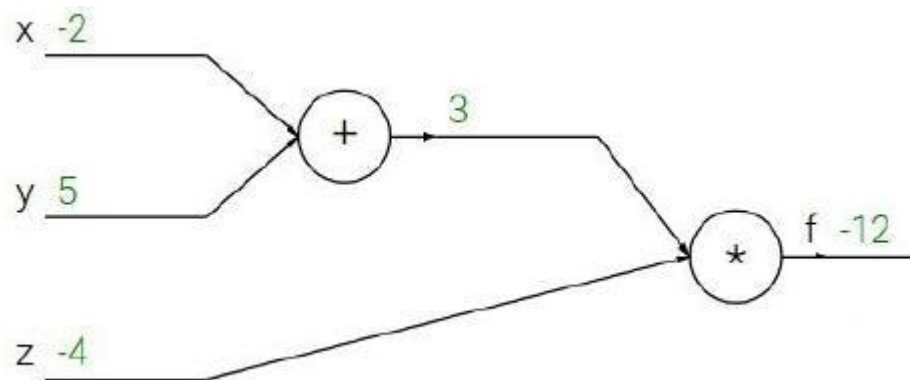
初中生物知识



一块上等的内酯豆腐

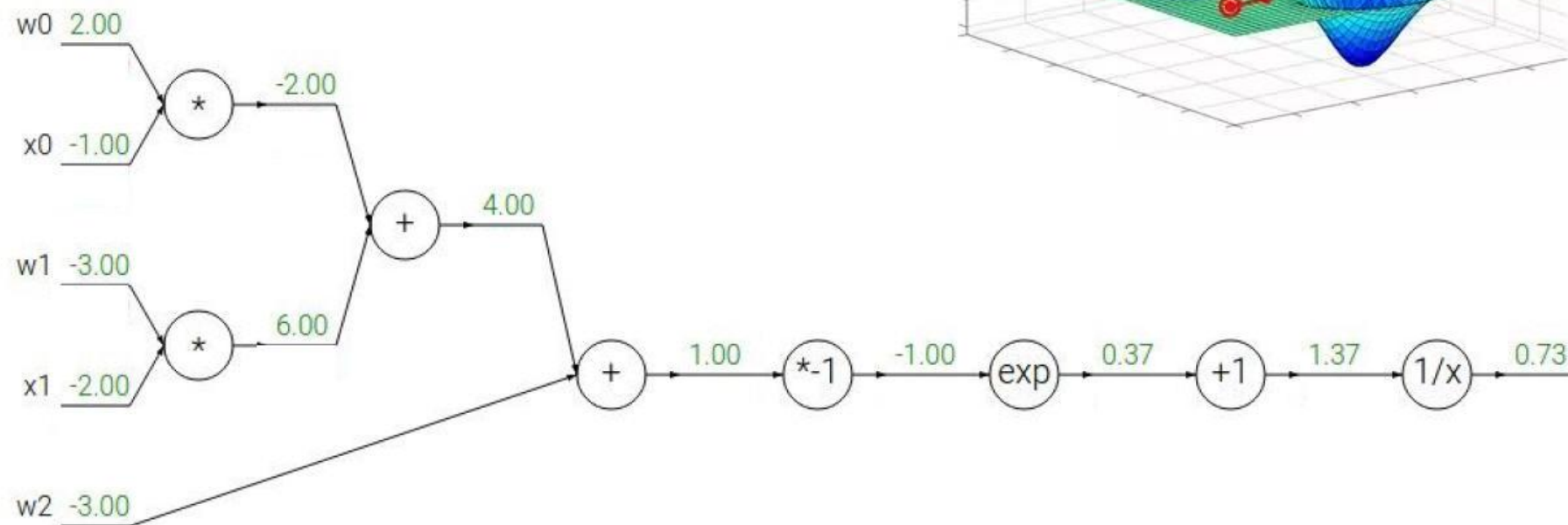
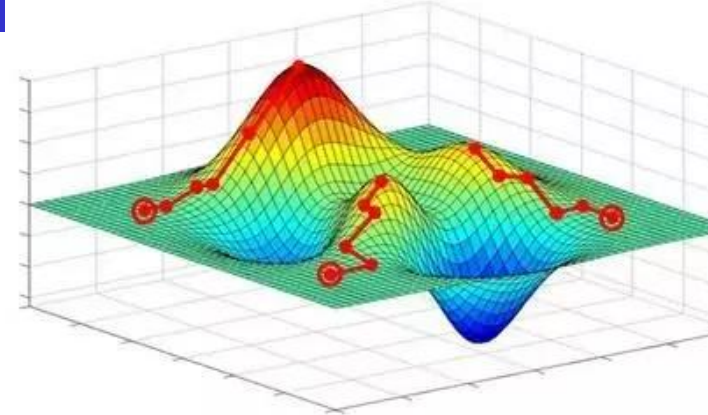


多个简单函数可以组合成复杂函数（如多元多次）



$$f(x, y, z) = (x + y)z$$

e.g. $x = -2, y = 5, z = -4$



$$f(w, x) = \frac{1}{1 + e^{-(w_0 x_0 + w_1 x_1 + w_2)}}$$

神经网络/深度学习对现实问题的解决思路

我们有什么？

输出结果 (Y: 流失、信用....)

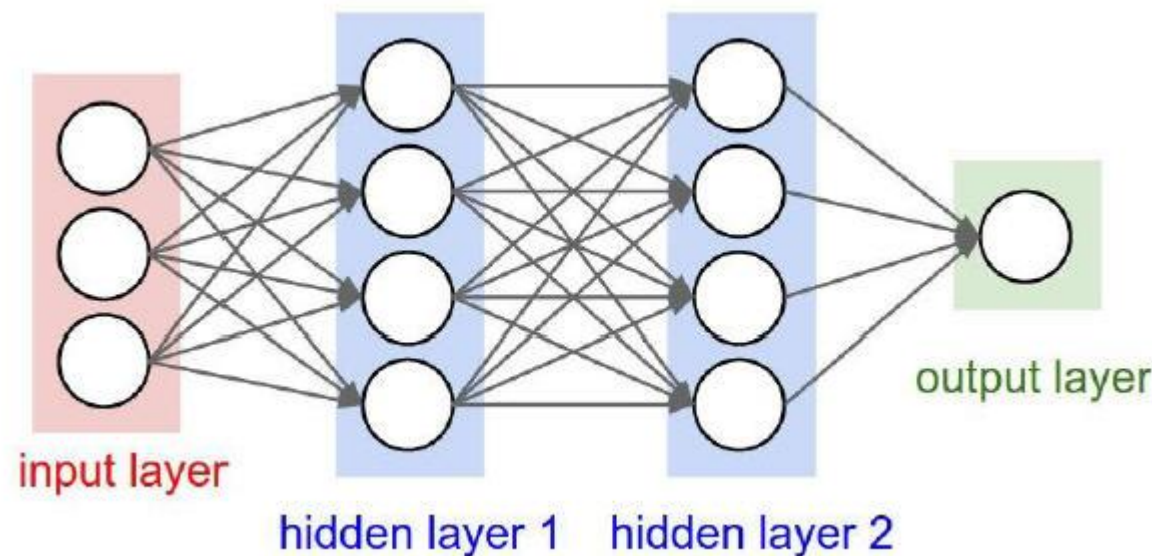
输入变量 (X: 年龄、性别、收入...)

我们想知道：

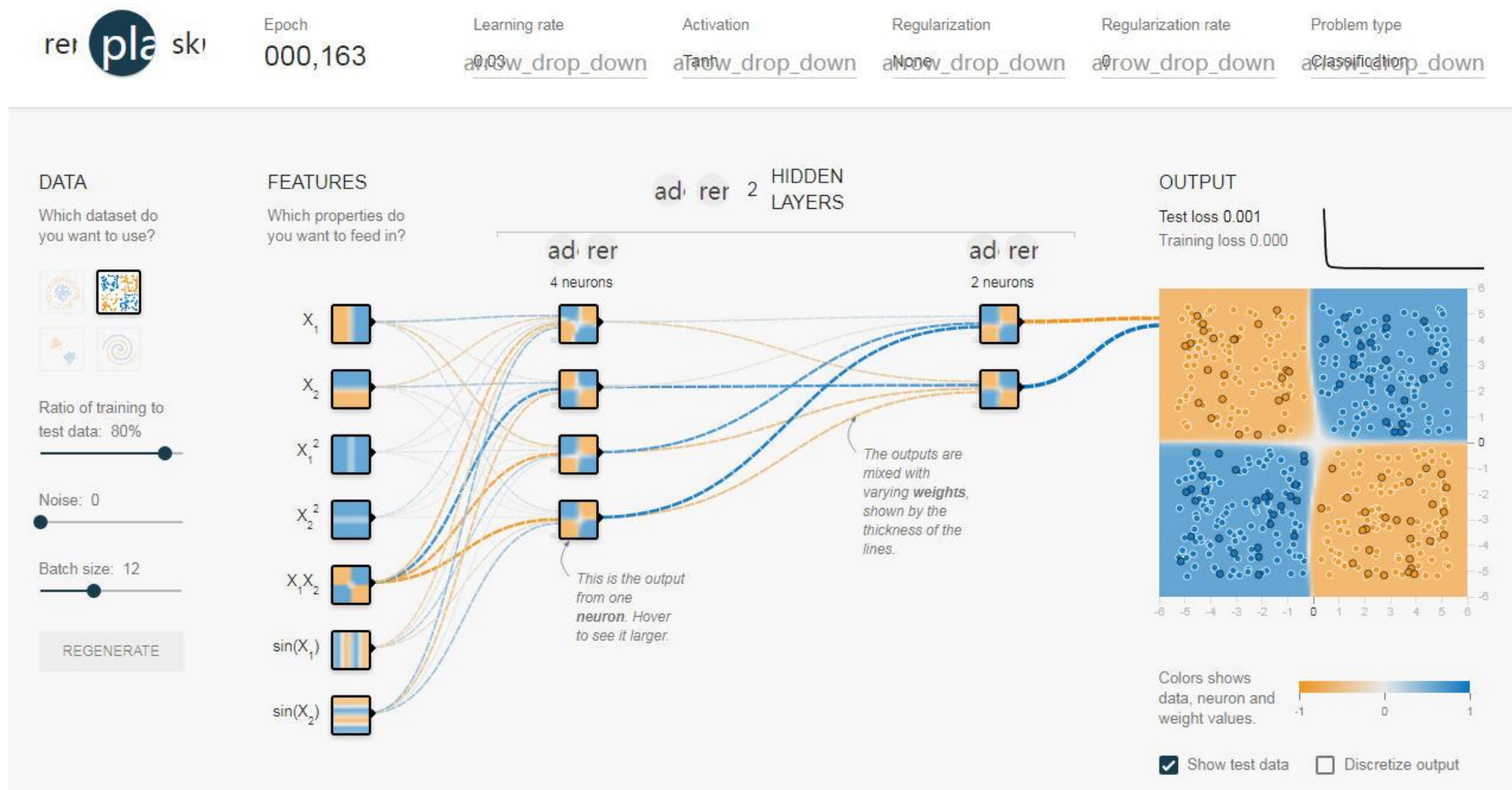
一个完美（事先不知道的）函数 $f()$ ，能够很好地描述输入变量和输出结果的关系，使得该函数的输出 Y' ，随着大量 X 的输入，越来越趋近 Y 。

所以

人类预设定层数和每层的神经元(参数)，由计算机通过连接各层中的神经元，去尝试各种组合（函数模块，运算规则），最终找到最合适表达 $Y=F(X)$ 的计算关系



<http://playground.tensorflow.org/>



数据：在二维平面内，若干点被标记成:黄色(-1)/蓝色(+1)，表示想要区分的两类。
线：表示权重. 它们控制了“如何组合”。神经网络的学习也就是从数据中学习那些权重。
输出：黄色背景颜色都被归为黄点类，蓝色背景颜色都被归为蓝点类。深浅表示可能性的强弱。

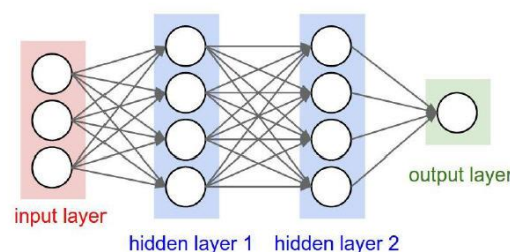
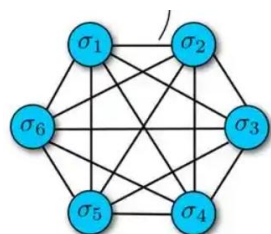
2024年诺贝尔物理学奖



John J. Hopfield



Geoffrey Hinton



Hopfield和Hinton设计并完善了人工神经网络及其计算

Hopfield受到对人脑认识的启发，神经元通过突触与其他神经元通信，设计了全联通神经网络。

Hinton进一步发展了神经网络，包括玻尔兹曼机（Boltzmann Machine，神经元之间通过权重相互连接，从而生成不同神经网络结构）和反向传播算法（更高效计算）

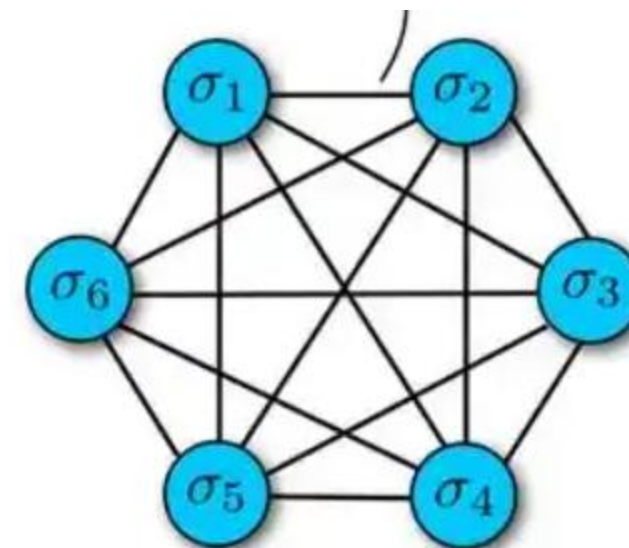
不同之处

Hopfield网络

每个神经元既是输入也是输出，并且可以更新自己的状态，信息可以在网络里循环

通过循环有效进行信息存储、更新和检索，用于记忆和联想

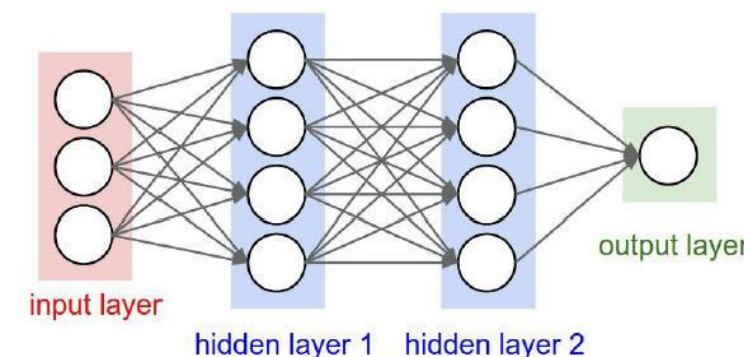
有不有点像我们现在关注的大模型记忆？



(反向传播算法对应的) 多层感知机 (MLP)

是一种有方向（称为Feedforward，前馈）神经网络，信息总是向前传递，不会返回到之前的层

通过学习实现分类或回归任务 $Y=F(x)$





理想

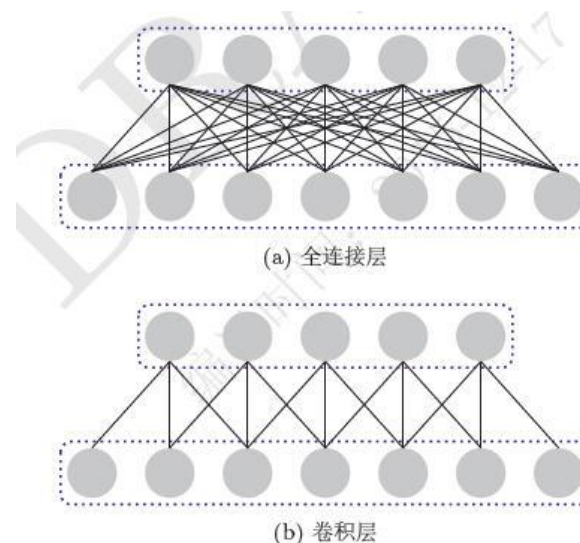


现实

深度学习：卷积神经网络 (Convolutional Neural Networks)

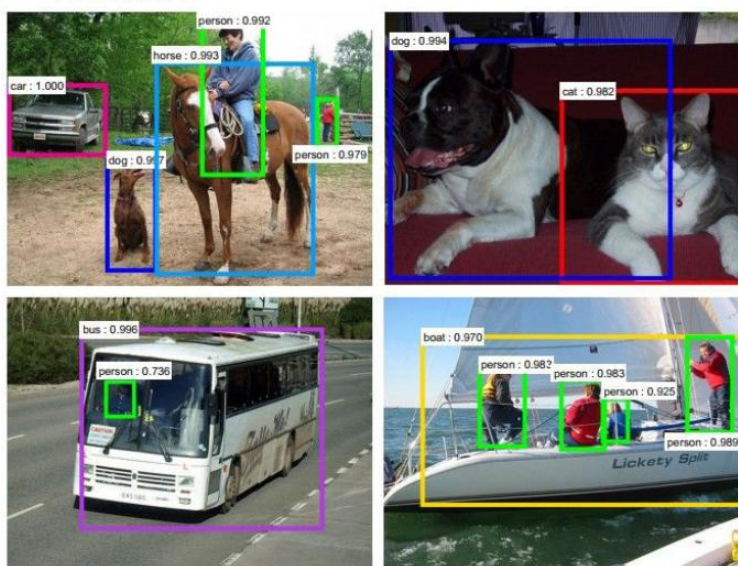
在一些场景下，可否让神经网络稍微聪明点地乱推，使得网络更有效率？

比如：一幅图片/棋谱格式固定，且总是有个明确且最终的表达含义



举例

Detection



Figures copyright Shaoqing Ren, Kaiming He, Ross Girshick, Jian Sun, 2015. Reproduced with permission.

[Faster R-CNN: Ren, He, Girshick, Sun 2015]

Segmentation

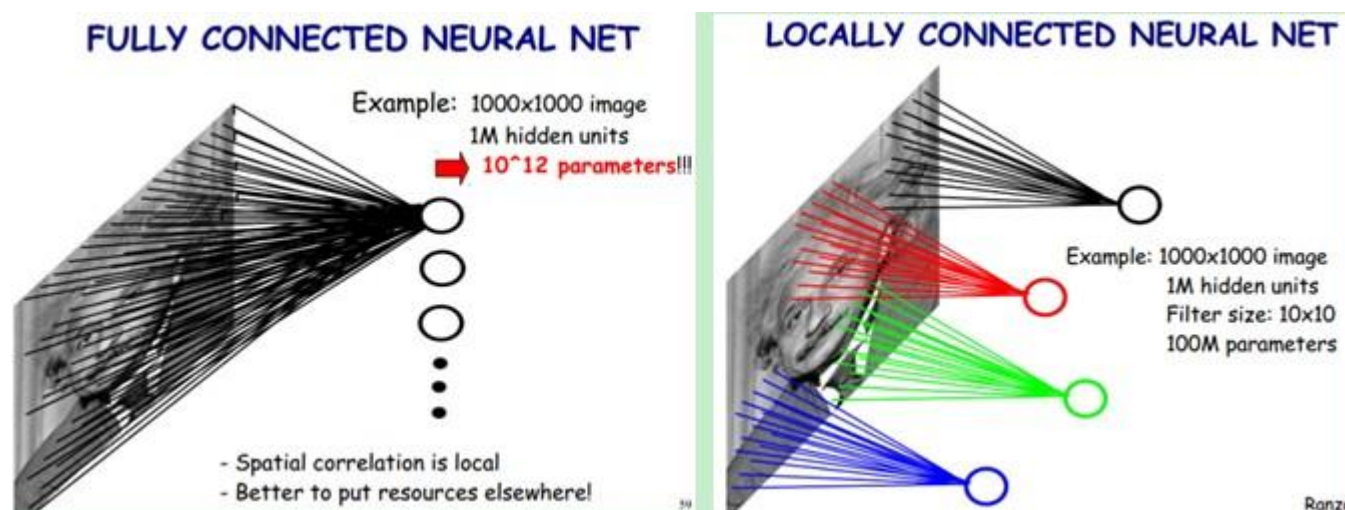


Figures copyright Clement Farabet, 2012. Reproduced with permission.

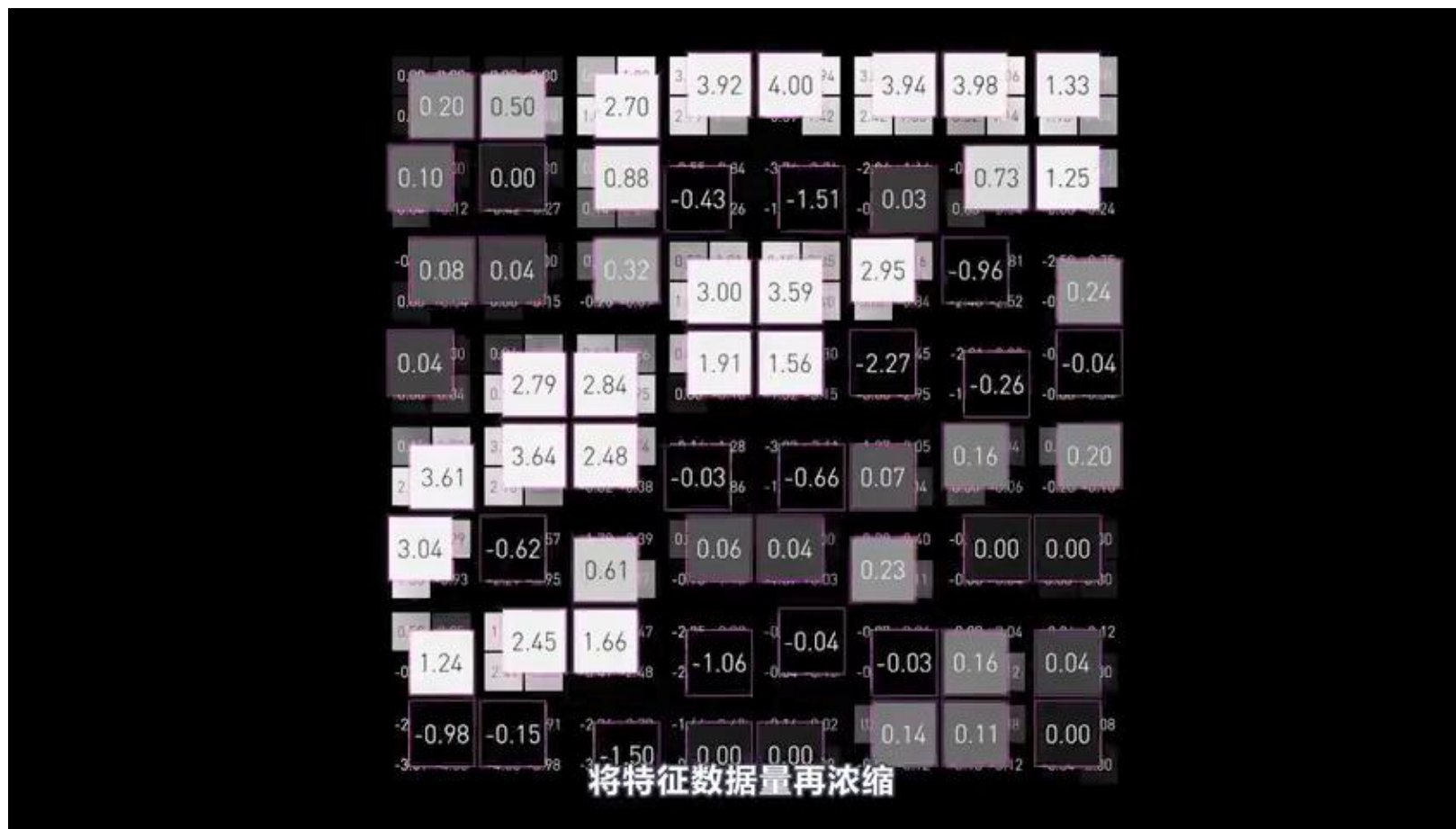
[Farabet et al., 2012]

受生物学上的发现影响：在人类听觉系统、本体感觉系统和视觉系统中的神经元，只有特定区域能感受特定刺激。比如视网膜上特定区域的刺激才能够激活相应神经元

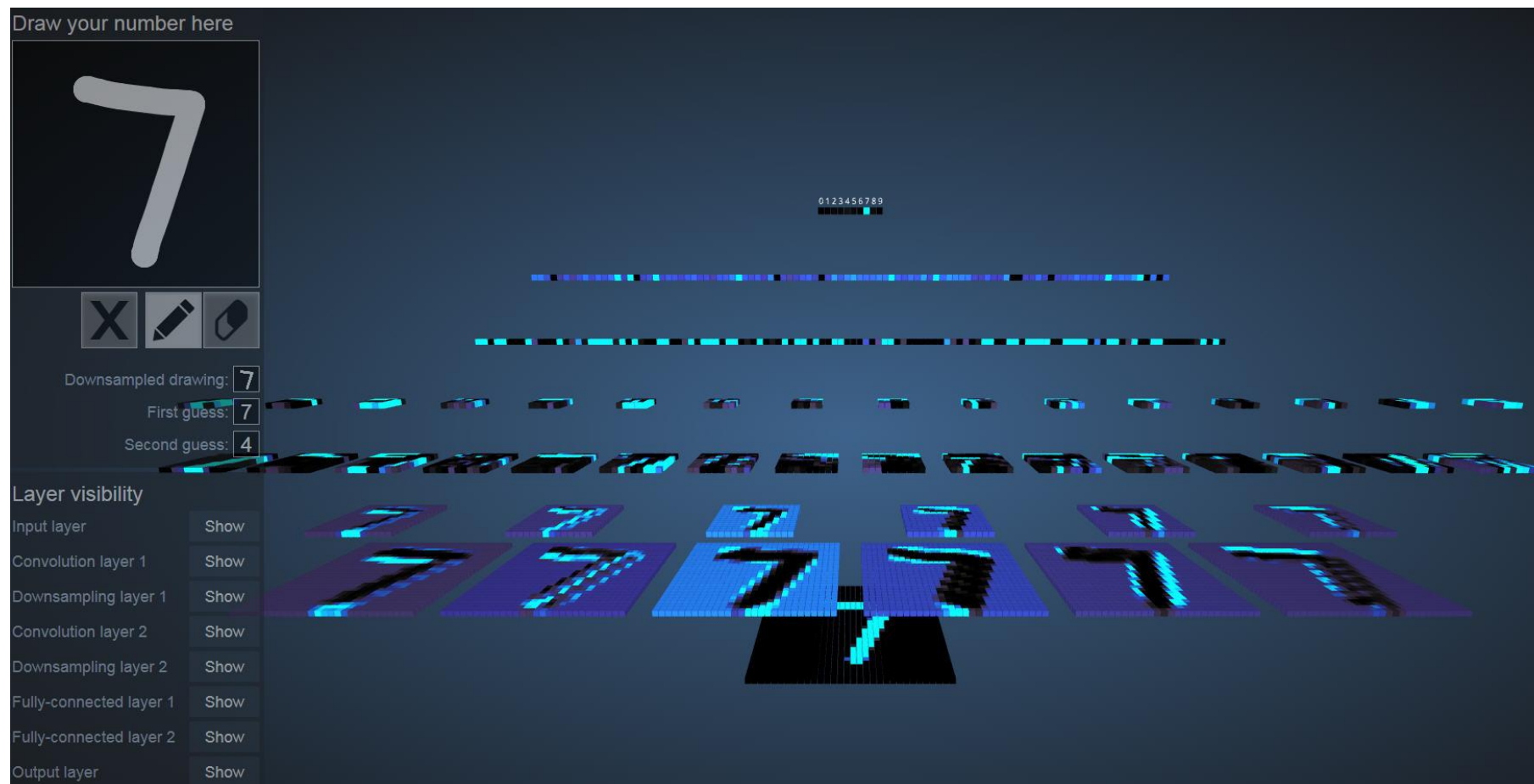
局部感知：形象地说，就是模仿你的眼睛，人在看东西时，目光是聚焦在一个相对很小的局部。



看段视频吧

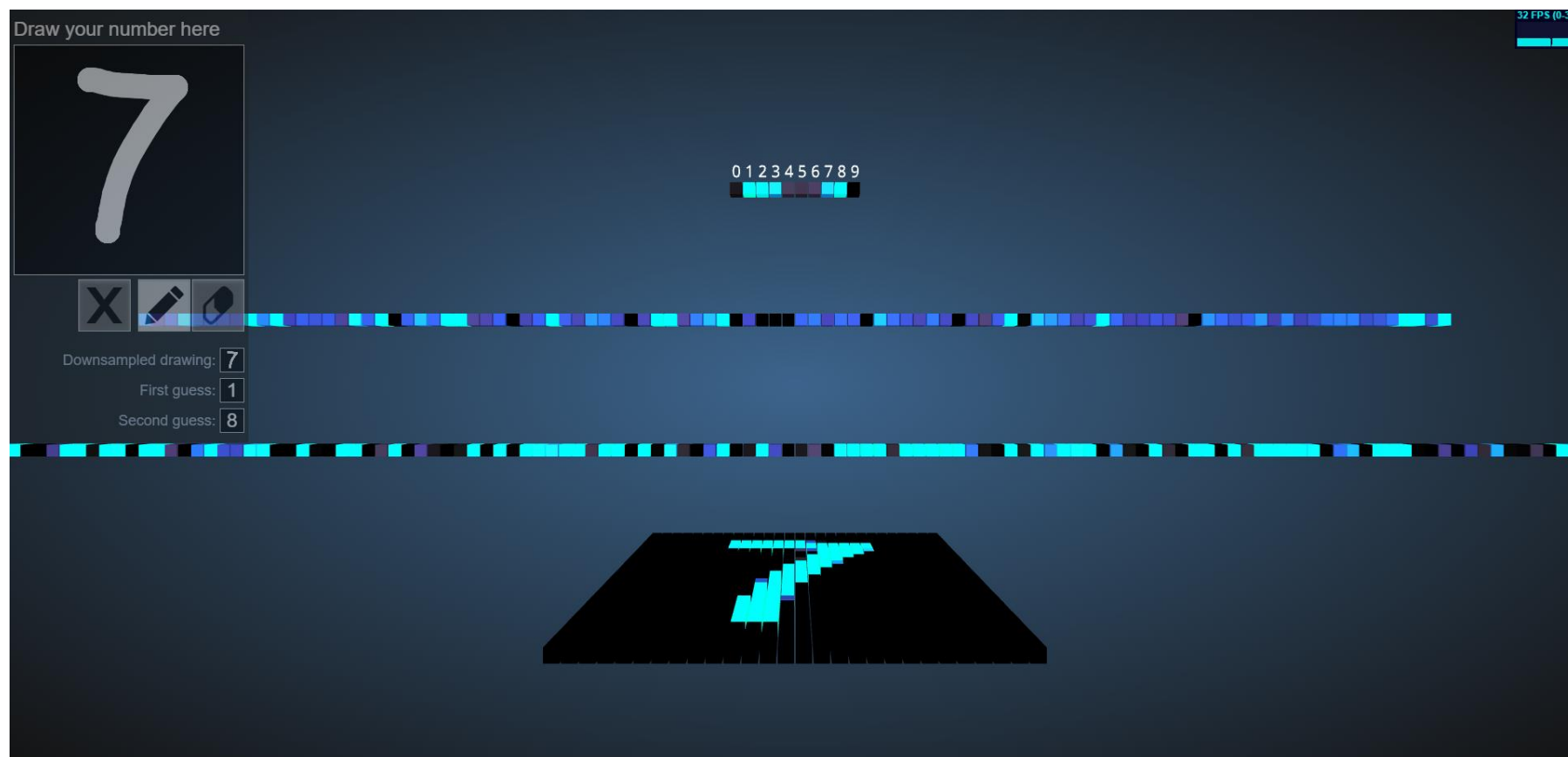


https://adamharley.com/nn_vis/cnn/3d.html

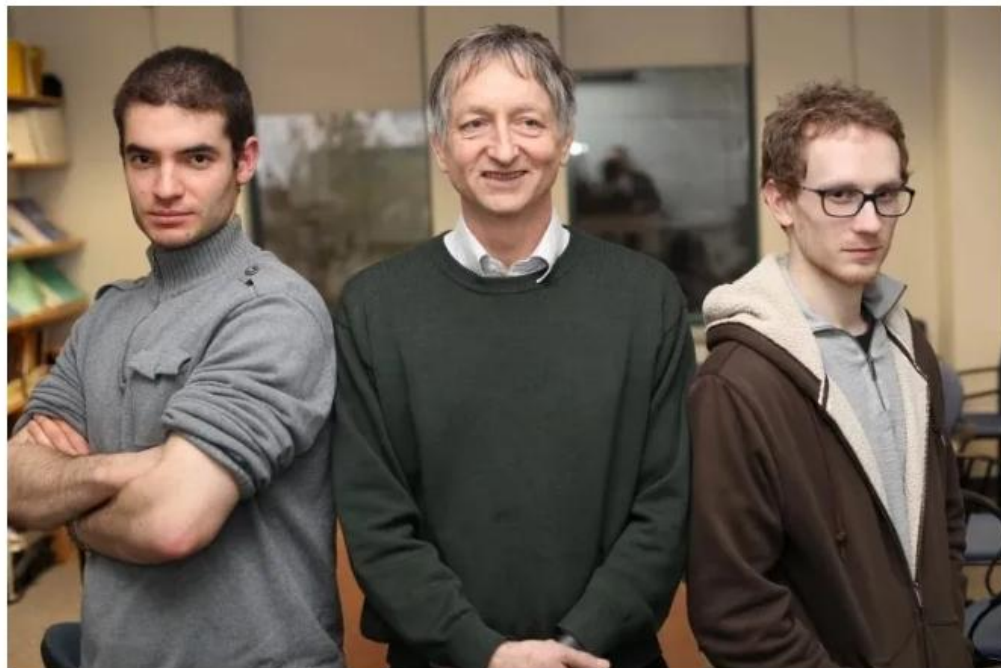


对比NN

https://adamharley.com/nn_vis/mlp/3d.html

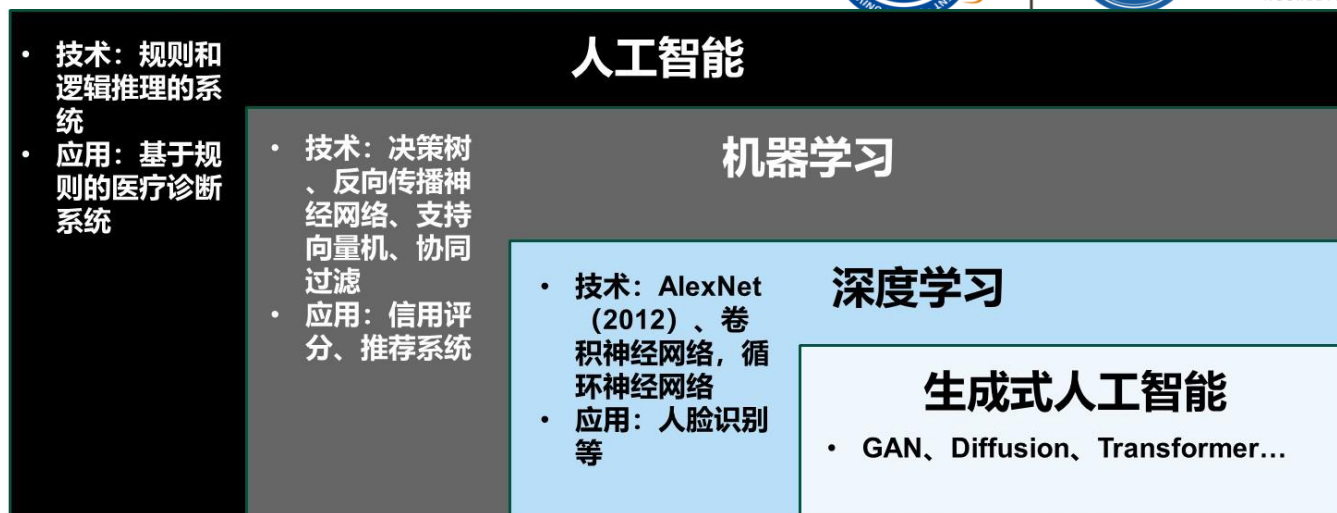


在 2012 年，来自多伦多大学的 Alex Krizhevsky、Ilya Sutskever、Geoffrey Hinton 等人提出了一个名为「AlexNet」的神经网络，赢得了 2012 年大规模视觉识别挑战赛 ImageNet 的冠军。



三位都是 AI 领域里响当当的人物。Geoffrey Hinton 被誉为「深度学习之父」，后来获得了 2018 年的图灵奖、2024 年的诺贝尔物理学奖；Ilya Sutskever 是 OpenAI 的联合创始人及前首席科学家，也是 AlphaGo 论文的众多作者之一。冠名该模型的 Alex Krizhevsky 也是 CIFAR-10 和 CIFAR-100 数据集的创建者，不过他却逐渐对研究失去了兴趣，于 2017 年 9 月离开了谷歌。

在描述当年的 AlexNet 项目时，Geoffrey Hinton 总结道：「Ilya 认为我们应该做这件事，Alex 让它成功了，而我获得了诺贝尔奖。」



Imagenet classification with deep convolutional neural networks

[A Krizhevsky, I Sutskever](#) - Advances in neural ..., 2012 - proceedings.neurips.cc

We trained a large, deep convolutional neural network to classify the 1.3 million high-resolution images in the LSVRC-2010 ImageNet training set into the 1000 different classes. On the test data, we achieved top-1 and top-5 error rates of 39.7% and 18.9% which is considerably better than the previous state-of-the-art results. The neural network, which has 60 million parameters and 500,000 neurons, consists of five convolutional layers, some of which are followed by max-pooling layers, and two globally connected layers with a final ...

☆ Save 📄 Cite Cited by 141389 Related articles All 89 versions 🔗

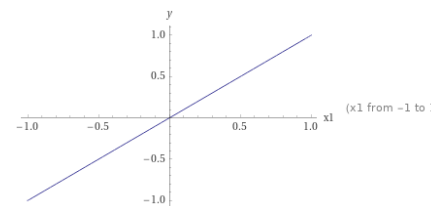
ImageNet classification with deep convolutional neural networks

[A Krizhevsky, I Sutskever, GE Hinton](#) - Communications of the ACM, 2017 - dl.acm.org

We trained a large, deep convolutional neural network to classify the 1.2 million high-resolution images in the ImageNet LSVRC-2010 contest into the 1000 different classes. On the test data, we achieved top-1 and top-5 error rates of 37.5% and 17.0%, respectively, which is considerably better than the previous state-of-the-art. The neural network, which has 60 million parameters and 650,000 neurons, consists of five convolutional layers, some of which are followed by max-pooling layers, and three fully connected layers with a final ...

☆ Save 📄 Cite Cited by 35871 Related articles All 11 versions 🔗

$$Y=F(X)$$



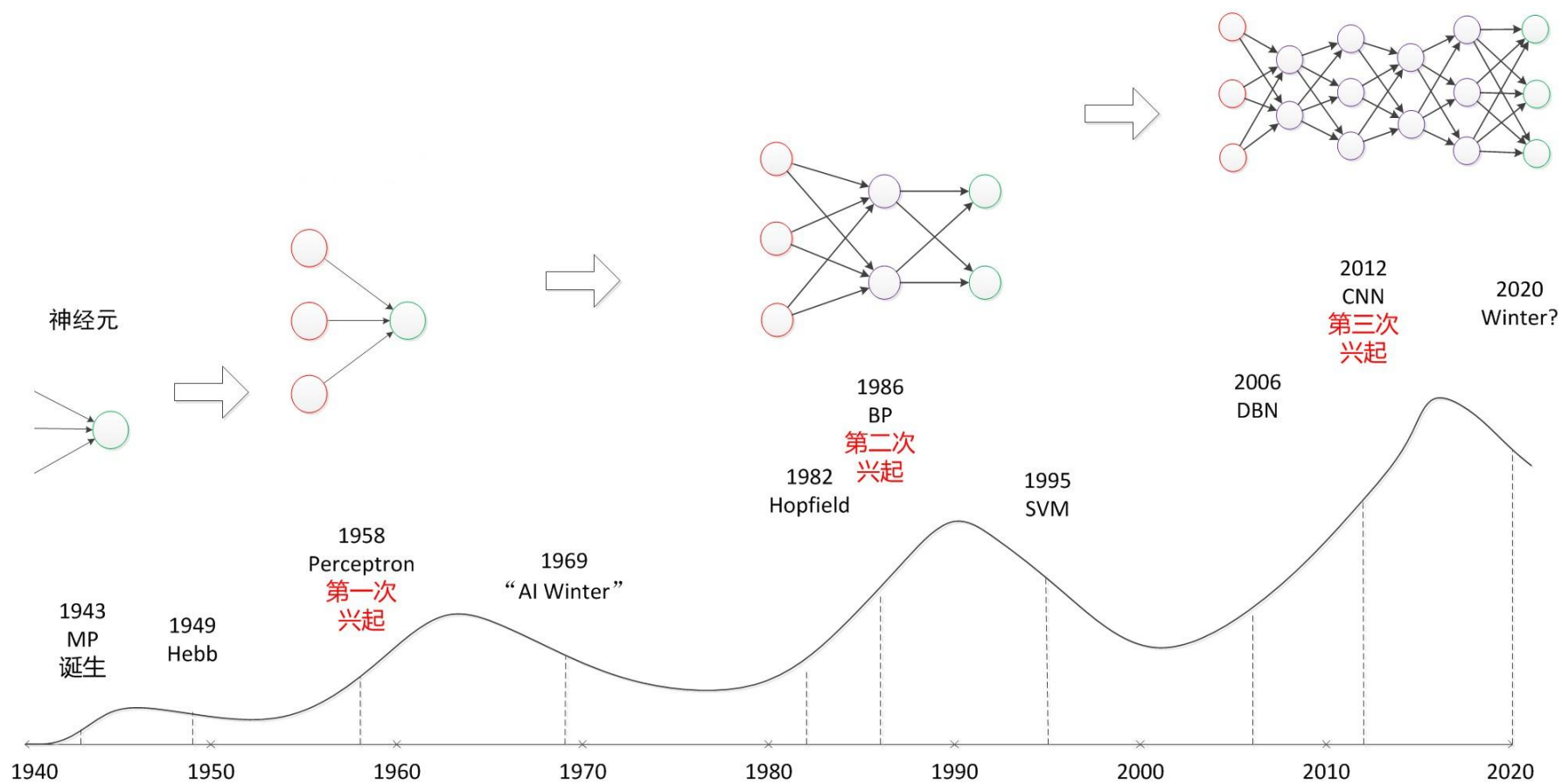
如今最先进的人工智能技术是上一代机器已经能做的事情的升级：从大量数据中发现隐藏的规律。

深度学习（神经网络）所有令人印象深刻的成就只不过是曲线拟合

曲线拟合(curve fitting): 是指选择适当的曲线类型来拟合观测数据，并用拟合的曲线分析两变量间的关系。

Judea Pearl (贝叶斯网络的发明者，图灵奖)





人工智能 AI

机器学习 ML

深度学习 DL

GPT

神经网络...

- > RNN(关注时序的循环神经网络)
- > Transformer (RNN+关注变量/文本间相关性的神经网络)
- > GPT(Transformer+1000亿参数的预训练深度学习)
- > Chat (GPT+ Prompt/COT+RL)

循环神经网络（recurrent network）

时间序列的预测，比如预测股价

天	价格
1	100
2	105
3	106
4	108

$$Y_t = 0.7Y_{t-1} + 0.3Y_{t-2} + X_t$$

假设这里暂不考虑X

$$Y_3 = 0.7 \times 105 + 0.3 \times 100 = 103.5$$

$$Y_4 = 0.7 \times 106 + 0.3 \times 105 = 105.7$$

循环神经网络（**recurrent network**）

基于先前的词来预测下一个词，也是时间序列预测

the clouds are in the ()

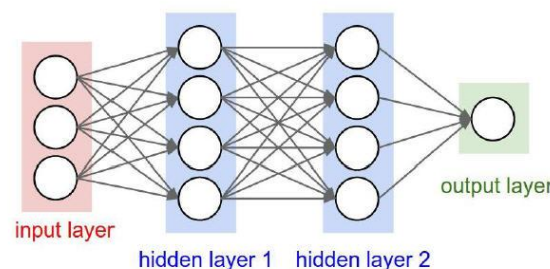
如果我们试着预测 “the clouds are in the ()” 最后的词，我们并不需要太多上下文 —— 下一个词很显然最有可能是 sky。在这样的场景中，传统的时间序列预测技术就可以完成

sky= F_1 (the clouds are in the)
the= F_2 (the clouds are in)
in= F_3 (the cloud are)
are= F_4 (the clouds)
clouds= F_5 (the)

我们甚至可以：
sky= F_1 (the, F_2 (in, F_3 (are, F_4 (cloud, F_5 (the))))))

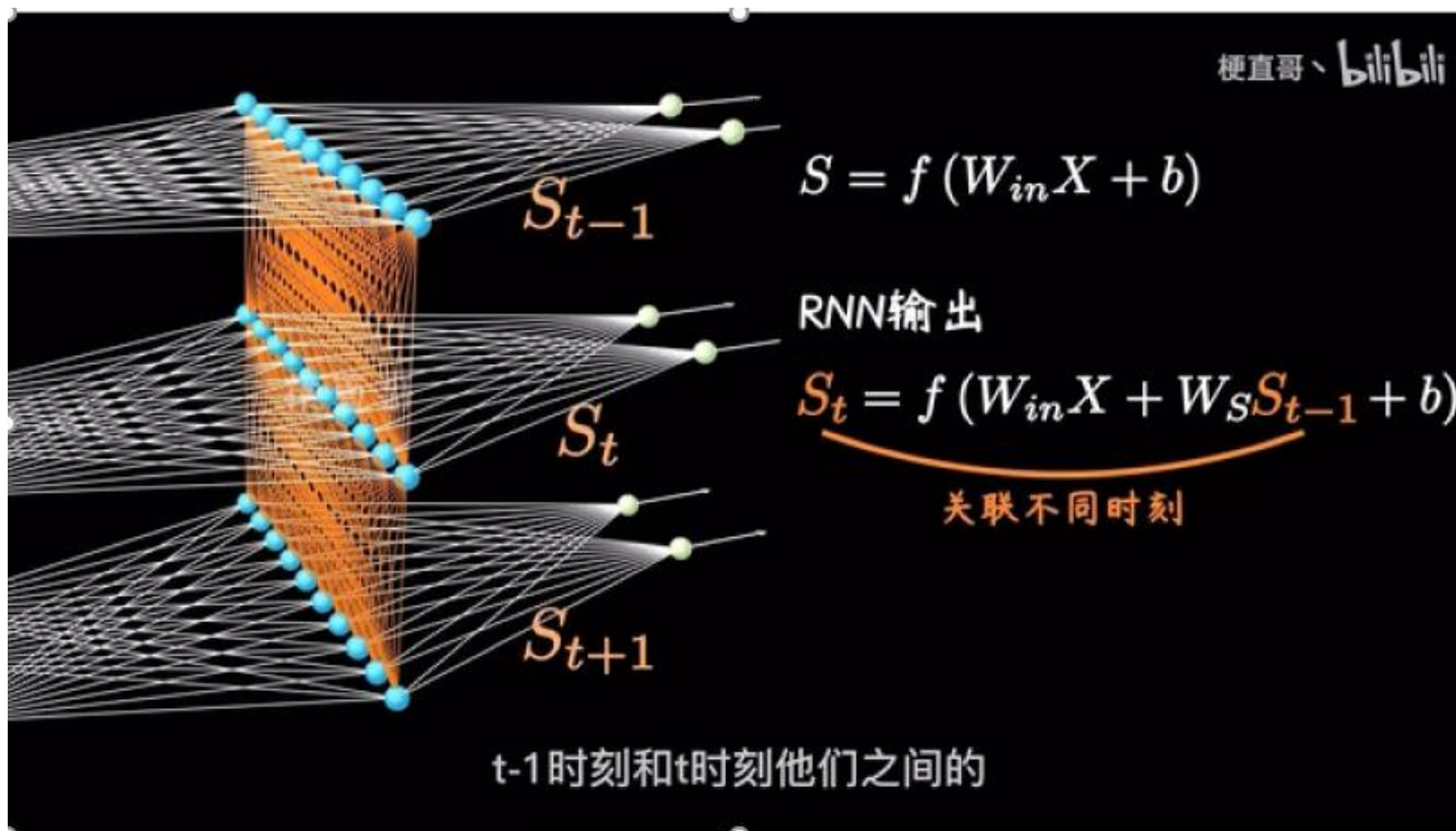
但对于“I grew up in France, attended a local school, and was immersed in the language and culture from a young age. I spent my weekends exploring the charming villages of the countryside, and during school holidays, I also often visit famous landmarks in Germany and Italy. I speak fluent () ”, 意味着自回归的时间间隔(i)要很大才行。

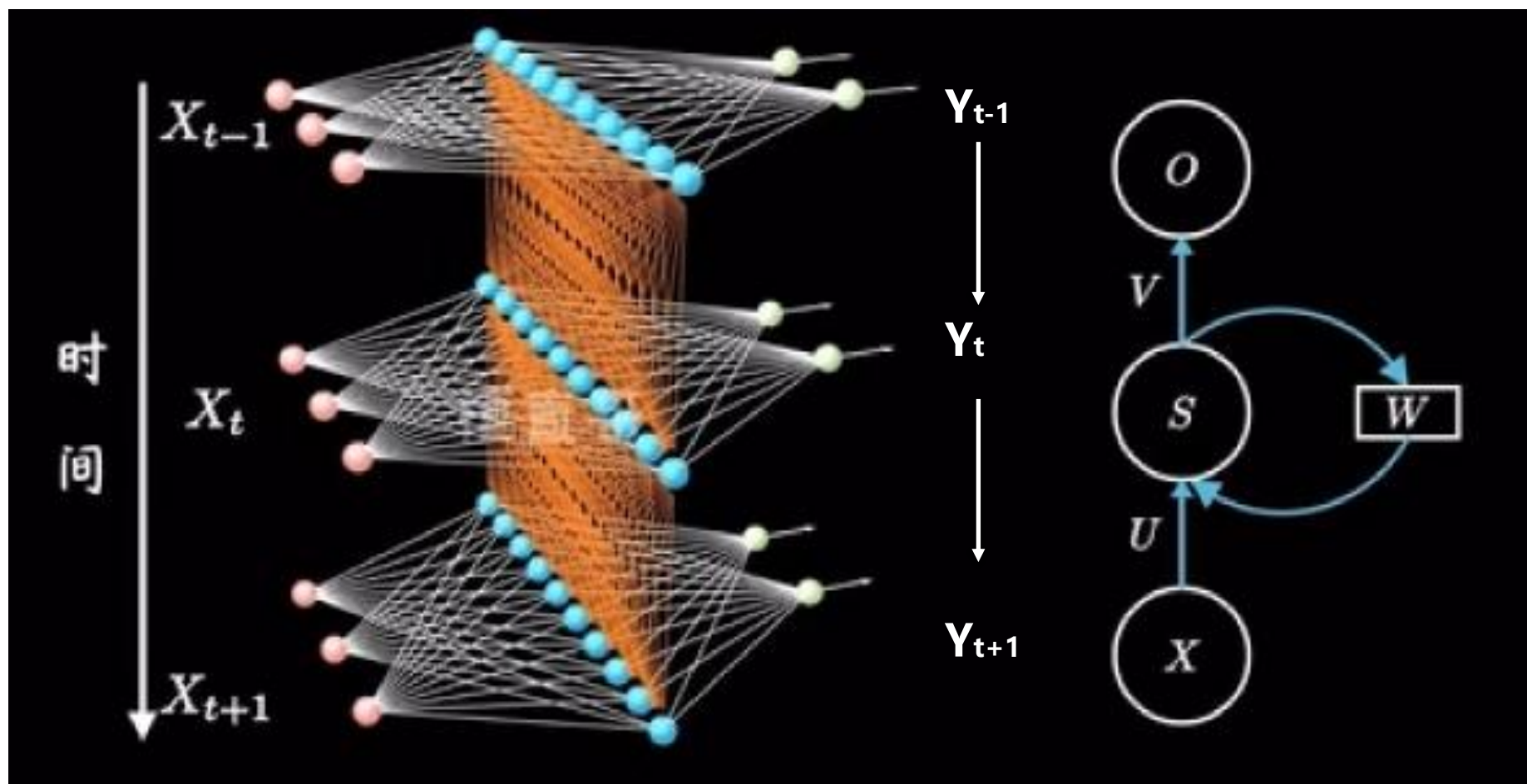
而且相比传统的时间序列，如何找出对应不同输入最合适的 i ，不再是简单加线性权重



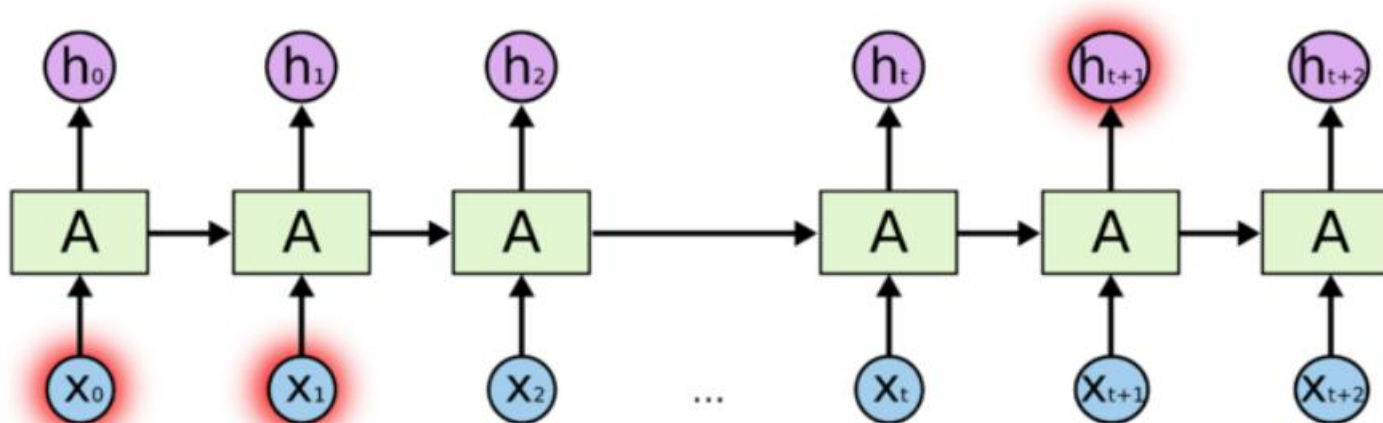


$$Y_t = a_1 Y_{t-1} + a_2 X_t + b$$





$$Y_t = ah(Y_{t-i}) + (1 - a) X_{t-i}$$



相当长的相关信息和位置间隔

GPT-4 的上下文长度是 32 k，GPT-4 Turbo的上下文长度是 128k（相当于12.8万英文字符或4.2万中文字）

人工智能 AI

机器学习 ML

深度学习 DL

GPT

- > Transformer (时序神经网络+关注变量/文本间相关性)
 - > GPT(Transformer+大尺度参数预训练)
 - > Chat (GPT+ 提示词工程/思维链+人机反馈强化)
 - > Reasoning (Chat + 混合专家/Agent + 强化学习)

Transformer

Attention is all you need
(在语言中) 你只需要关注 (上下文)



当大家看到下面图片，会首先看到什么内容？



婴儿？文章？图例？

我们的大脑会把注意力放在主要的信息上（这就是注意力机制）。

类似的，我们读书时是如何捕捉到文字的主要内容的呢？

熟读的背后

词的关系（一句话不用逐字读就能了解它的大意，词+组合+顺序=含义）

上下文（足够长的内容）

对应的，在Transformer中

RNN后，Transformer (with Attention) 关注如何更好聚焦/放大需要关注的记忆，更好调整记忆之间的关系

(自)注意力机制：学习句子中每个单词的相关性和上下文依赖关系

- 允许每个单词“关注”序列中的其他每个单词

- 计算注意力权重以表示单词之间的关系

- 允许多条关系（多头）来捕捉语言的不同方面

初衷：机器翻译

六、读一读，连一连。(8分)

一辆	书	绿绿的	太阳
一架	马	红红的	禾苗
一册	车	弯弯的	河水
一匹	飞机	清清的	小路

七、照样子，填一填。(12分)

走来走去 飞来飞去、跑来跑去

胖乎乎 白花花的、金灿灿的

安安静静 快快乐乐的、高高兴兴的

连连看

3. 连线题(把古诗的上句和下句用直线连起来)

春风又绿江南岸	每逢佳节倍思亲
黑发不知勤学早	江枫渔火对愁眠
独在异乡为异客	明月何时照我还
空山不见人	回首白云低
月落乌啼霜满天	白首方悔读书迟
举头红日近	但闻人语响

七年级下期末完型填空专项练习

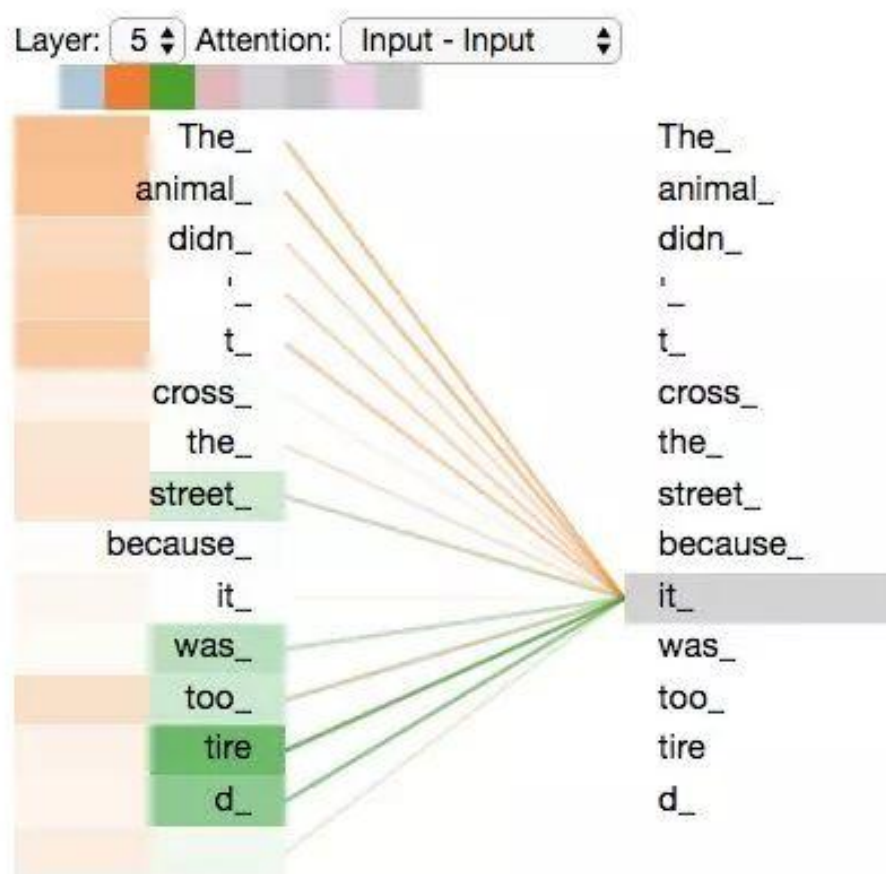
(一)

In England, nobody 1 the age of eighteen can drink in a bar(酒吧). Mr Green often 2 to a bar near his house, but he never took his 3, Tom, because Tom was very young. After Tom had his eighteenth birthday, Mr. Green 4 Tom to his usual bar for the first time. They drank 5 an hour, then Mr. Green said to his son, "Now, Tom. I want to 6 you a lesson, you must be 7 not to drink too much. 8 do you know when you have enough? Well, let me tell you. Can you 9 those two lights at the end of the bar? When they come to become 10, you're having enough and should go home."

"But, Dad," said Tom, "I can see only one light at the end of the bar."

- | | | |
|--------------------|------------|----------|
| () 1. A. above | B. under | C. at |
| () 2. A. went | B. shouted | C. wrote |
| () 3. A. daughter | B. wife | C. son |
| () 4. A. sent | B. carried | C. took |
| () 5. A. for | B. in | C. to |
| () 6. A. buy | B. bring | C. give |

注意力：如何根据上下文确定it指谁？



注意力机制会去关注与it相关的动物、街道和疲惫等

街道还是动物会疲惫呢？

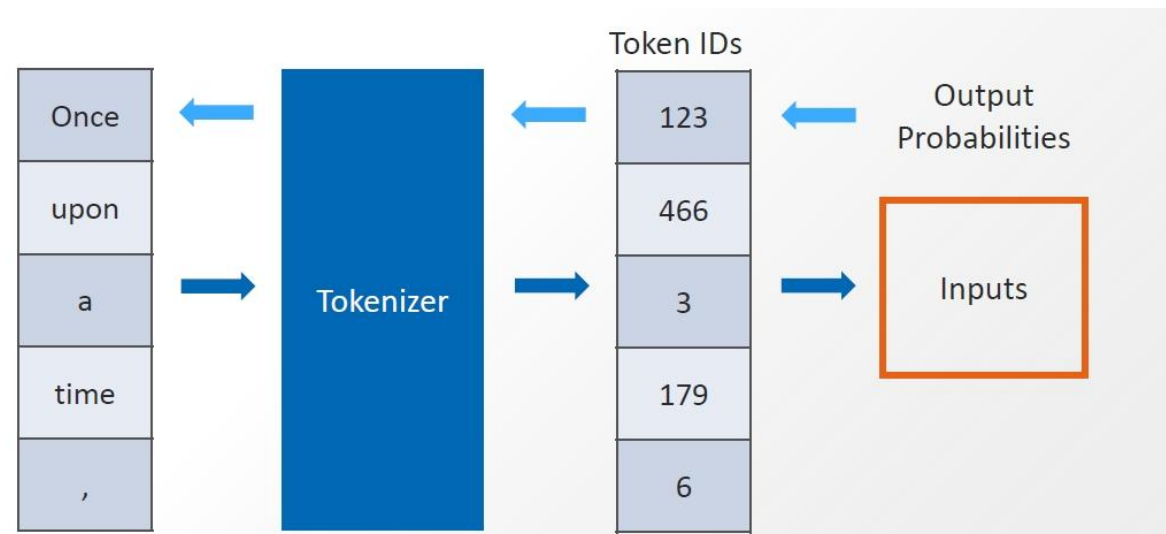
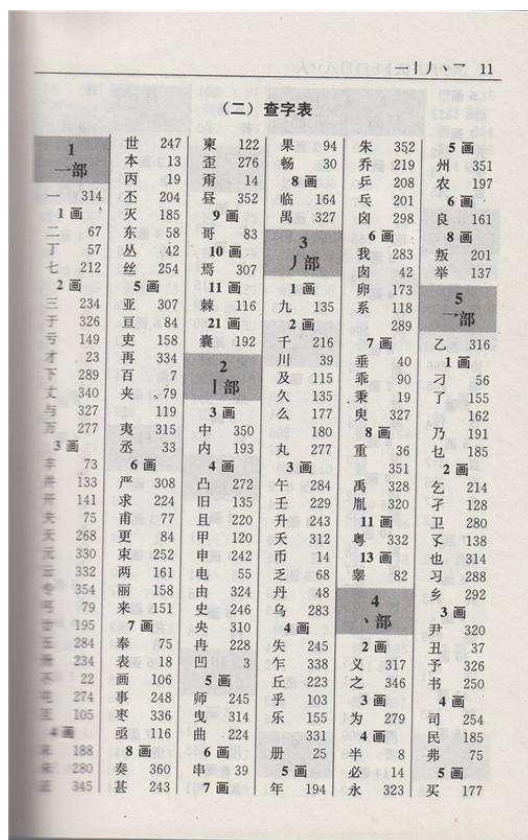
准备工作：词的数字化及相关性

每句话

分词

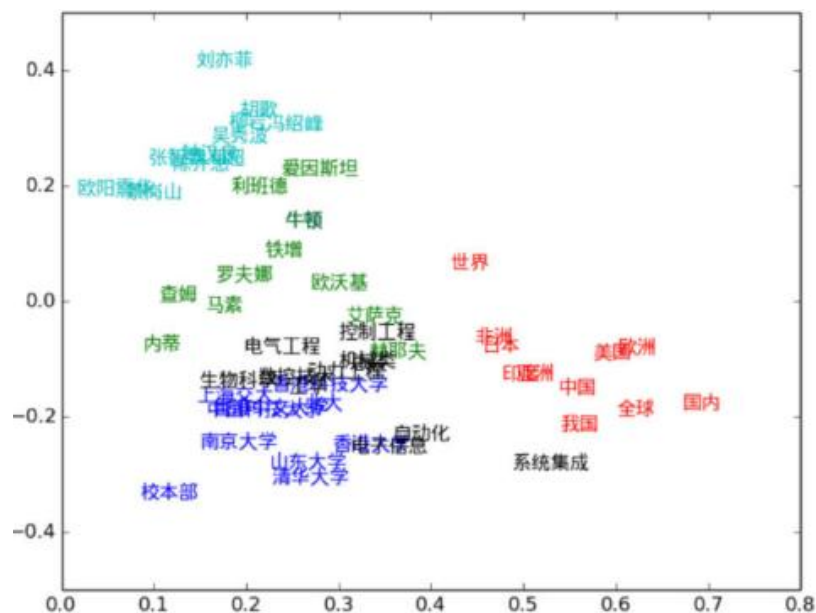
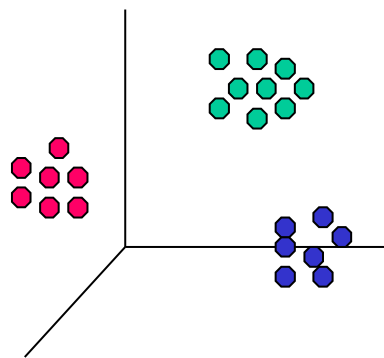
词被转化为对应的数值

单词被分词(tokenized)成词典中的位置



OpenAI的API定价为0.02美元/1000Token (词+标点符号等)

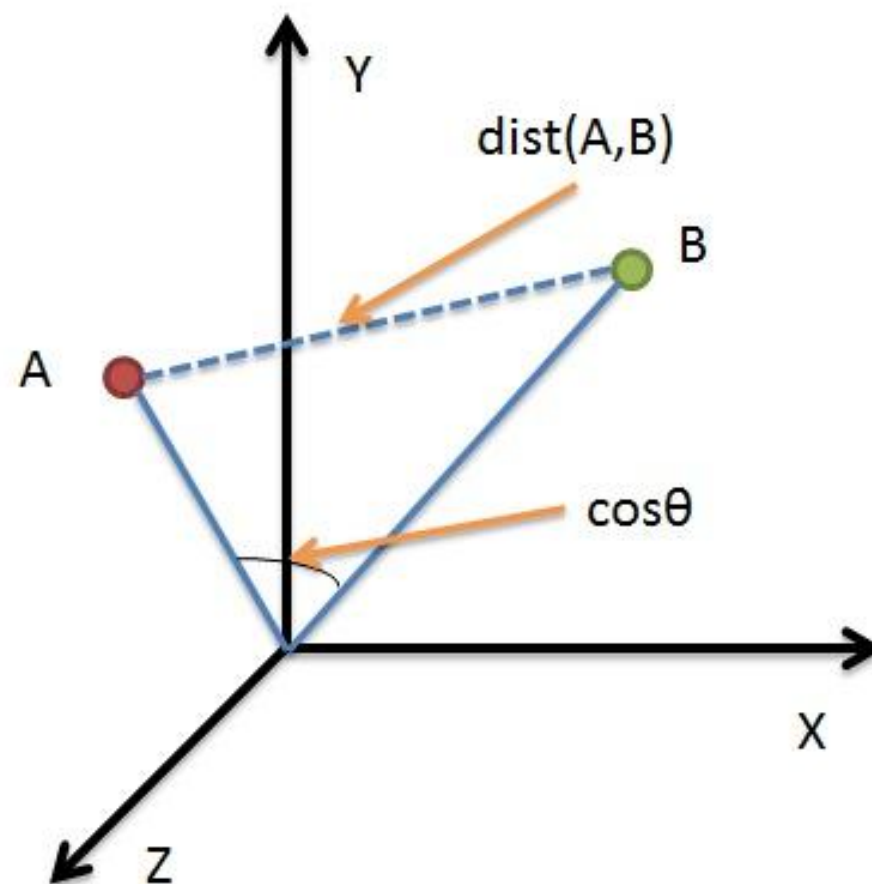
准备工作：词的相关性



语言的相似度/相关性：

- 用词作为坐标维度，可以构成一个词宇宙
- 每句话的相似度可以看作组成这句话的词形成的坐标，从而可以计算每句话在词宇宙的空间距离，即相似度。

余弦距离 Cosine Distance/Similarity



注重方向差异，而不是距离长度差异

我们如何阅读（理解和表达）一本外文书（比如英语）

理解

1. 阅读每个英语单词，并从中提取基本的意义或概念。这是对单词的基本理解。
2. 阅读时，会自然记住每个单词在句子中的位置。这帮助你理解整个句子的结构和顺序。
3. 当你遇到一个复杂的句子时，你可能需要回顾前面的词以充分理解这个词的意思。你把不同的单词关联起来，从而形成对整个句子的理解。
4. 你更进一步地加工和整理你所理解的信息，以便形成一个全面的故事或意义。

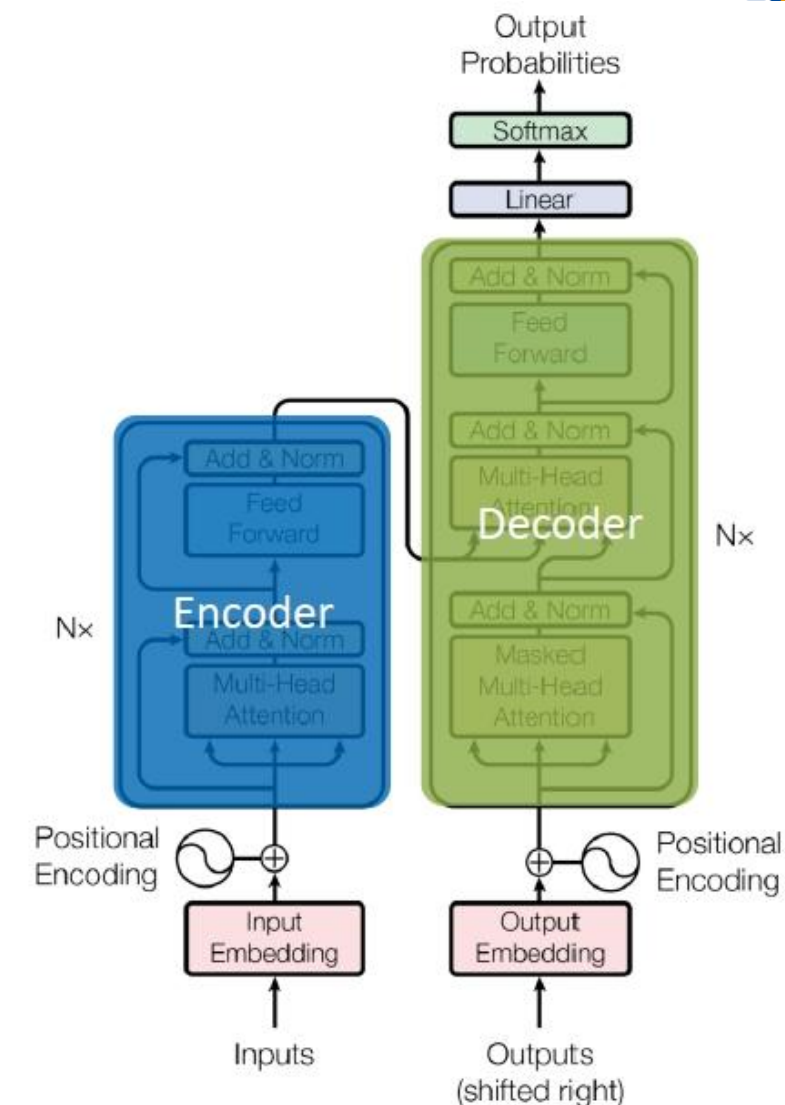
读完一部分后，由此形成了对这段文字的完整理解（字面和意思上的理解）

表达（把理解的内容（英语故事）用你自己的语言（中文）讲述出来）：

1. 你回忆起某个内容（的起点），开始用你自己的语言来表达。
2. 同时，根据记忆中每个词的位置信息，帮助你考虑语句表达的结构和顺序。
3. 在构思你要表达的新句子时，你不断参考之前已经用自己的语言表达出的话，以确保流畅和逻辑的衔接。
4. 在表述每个词时，你反复参考你对原文的理解，确保翻译准确，每个词都有所指。
5. 你对用自己的语言表达出的句子进行加工和调整，使其更自然和连贯。

■ 经过叙述的过程，你说出了一段流畅的中文语句，表述了你对英文原文的理解。

Transfomer技术架构



Images from **Attention Is All You Need**

<https://arxiv.org/abs/1706.03762>

理解

表达

如何理解编码和解码的每个过程：想象你正在阅读一本外文书（比如英语）



编码过程（理解）

1. 输入嵌入：

1. 这相当于你在阅读每个英语单词，并从中提取基本的意义或概念。这是对单词的基本理解。

2. 位置编码：

1. 当你阅读时，你会自然记住每个单词在句子中的位置。这帮助你理解整个句子的结构和顺序。

3. 自注意力机制：

1. 当你遇到一个复杂的句子时，你可能需要回顾前面的词以充分理解这个词的意思。你把不同的单词关联起来，从而形成对整个句子的理解。

4. 前馈网络：

1. 你更进一步地加工和整理你所理解的信息，以便形成一个全面的故事或意义。

经过编码过程，你形成了对这段文字的完整理解和内部表示

解码过程（表达：把理解的内容（英语故事）用你自己的语言（中文）讲述出来：

1. 输出嵌入：

这个阶段相当于你回忆起某个概念（你理解中的一个节点），并将其用你自己的语言来表达。

2. 位置编码：

使用位置编码来保留输出序列中每个词的位置信息。这帮助模型理解生成句子的结构和顺序。

3. 自注意力机制：

在构思你要表达的新句子时，你不断参考之前已经用自己的语言表达出的话，以确保流畅和逻辑的衔接。

通常解码器一次只生成一个词，并使用先前生成的词（加上位置编码）作为上下文，确保生成的新词与之前的词在逻辑上和语法上连贯。

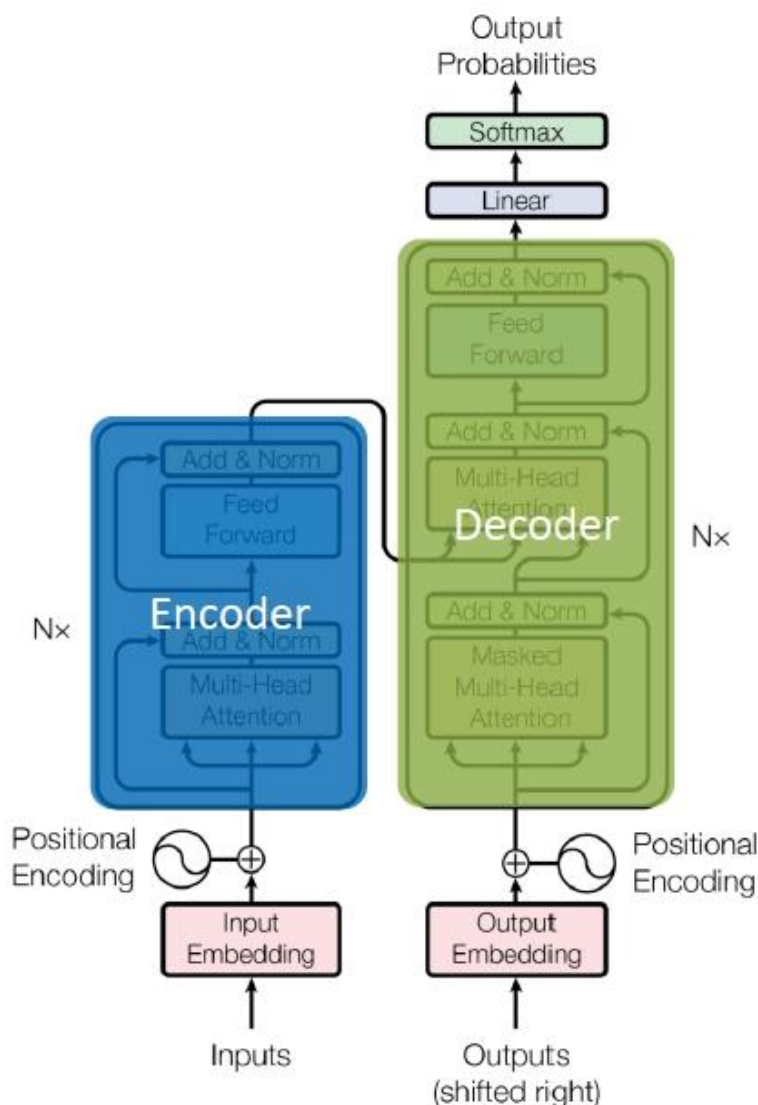
4. 编码器-解码器注意力：

在表述每个词时，你反复参考你对原文的理解，确保翻译准确，每个词都有所指。

5. 前馈网络：

你对用自己的语言表达出的句子进行加工和调整，使其更自然和连贯。

- 经过解码过程，你生成了一段流畅的中文文本，表述了你对英文原文的理解。



Images from **Attention Is All You Need**

<https://arxiv.org/abs/1706.03762>

编码（Encode）过程

1. 输入嵌入（Input Embeddings）：

1. 输入序列中的每个单词被转换成一个固定维度的向量，称为嵌入。这个过程通常依赖于一个预训练的词向量表

2. 位置编码（Positional Encoding）：

1. 因为 Transformer 没有像 RNN 那样的顺序处理能力，所以需要添加位置编码来为每个输入单词添加位置信息，帮助模型捕捉序列中的顺序。

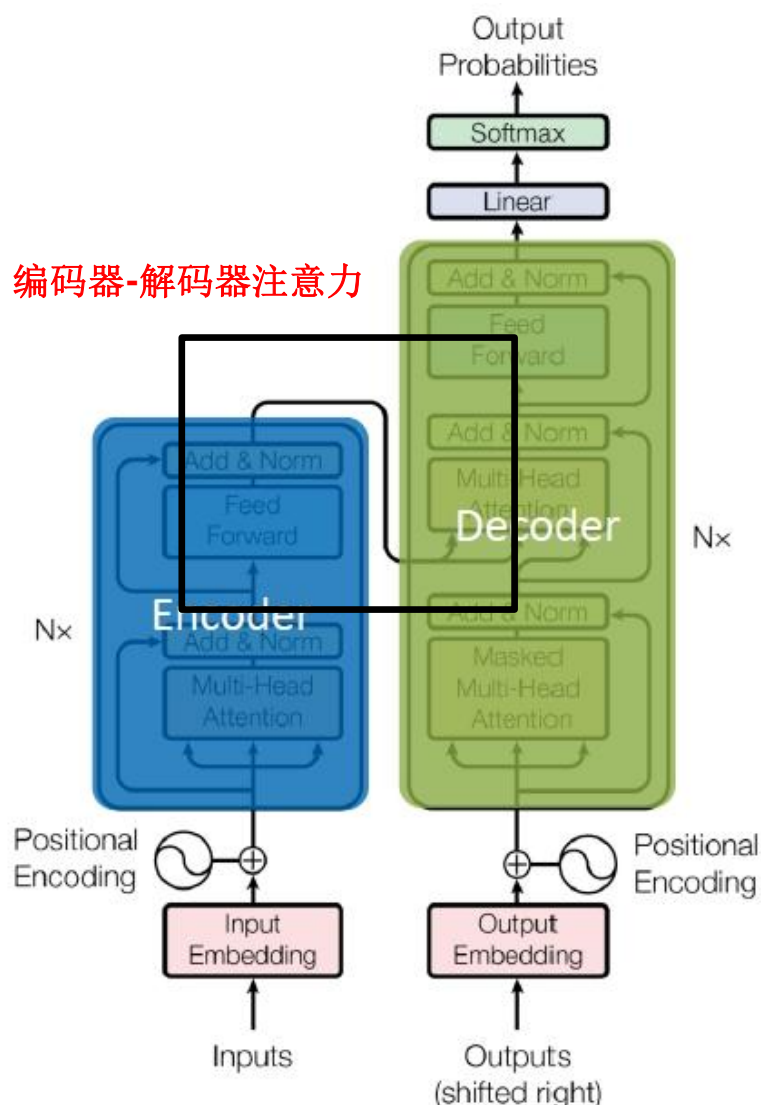
3. 多个编码层（Multiple Encoder Layers）：

1. 编码器由多个相同的子层组成，每个子层包括多头自注意力机制（Multi-Head Self-Attention）和前馈神经网络（Feed Forward Neural Networks）。
2. 自注意力机制：计算输入序列中每个词与序列中其他词的注意力得分，并进行加权求和，以强调相关单词的信息。
3. 前馈网络：对通过自注意力机制后得到的输出进行进一步的非线性转换。

4. 残差连接与层归一化（Residual Connections and Layer Normalization）：

1. 每个子层采用残差连接来保留输入信息，并通过层归一化（Layer Normalization）来稳定训练过程。

经过编码的结果是一个包含上下文信息的向量序列，可以传递给解码器。



Images from Attention Is All You Need

<https://arxiv.org/abs/1706.03762>

解码 (Decode) 过程

1. 输出嵌入 (Output Embeddings) :

1. 解码器接受输入时, 首先将目标序列中的词转换为嵌入

2. 位置编码 (Positional Encoding) :

1. 同样, 为解码器的输入词添加位置编码。

3. 多个解码层 (Multiple Decoder Layers) :

1. 与编码器类似, 解码器也由多个相同的子层组成。
2. 自注意力机制: 在目标序列中进行自注意力计算。

3. **编码器-解码器注意力 (Encoder-Decoder Attention)** : 将解码器的状态与编码器的输出结合以获取上下文信息。

4. 前馈网络: 与编码器中的前馈网络相似。

4. 生成输出 (Generating Output) :

1. 每个步骤根据解码器的输出来预测下一个词, 通过软最大 (Softmax) 函数输出概率分布, 并选择具有最高概率的词作为当前时刻的预测。
2. 解码器可以逐步生成完整的输出序列, 直到生成完成标志位置。

5. 残差连接与层归一化:

1. 同样, 解码器的每一层采用残差连接并进行层归一化

通过这种编码-解码机制, Transformer可以高效处理自然语言处理任务, 包括但不限于机器翻译、文本生成和问答。

2024年诺贝尔化学奖



David Baker



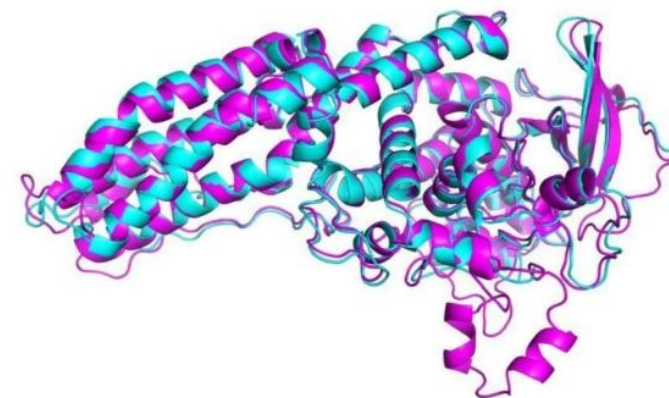
Demis Hassabis



John Jumper

David Baker
- 计算蛋白质设计

Demis Hassabis和John Jumper
- AlphaFolder在蛋白质结构预测方面的成就

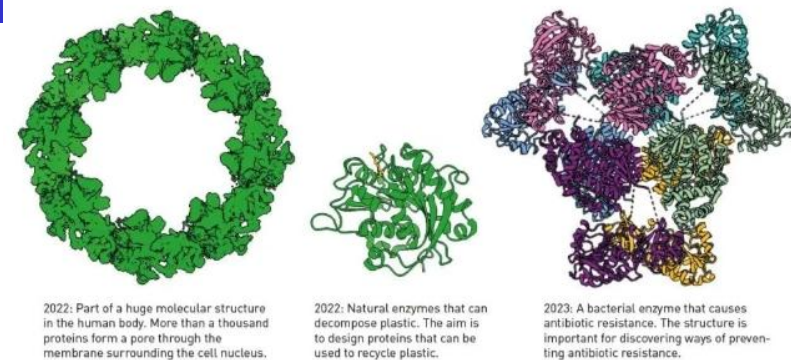


A DeepMind model of a protein from the Legionnaire's disease bacteria (Casp-14) (军团病细菌(Casp-14)蛋白

质的DeepMind模型)



AlphaFolder原理



使用 AlphaFold2 确定的蛋白质结构。

数据准备

蛋白质数据库（PDB）...

数据清洗：去除重复或错误的条目，标准化序列和结构的表示方式。

数据标注：将蛋白质的一级结构（氨基酸序列）与其对应的三维结构（折叠后的坐标）匹配好

特征提取:捕捉到蛋白质的折叠规则

氨基酸的物理和化学特性：如氨基酸的带电性、亲水性或疏水性等。编码为数值特征

共进化信息：通过序列比对技术（如多序列比对，MSA）提取在不同生物物种中共同变化的氨基酸残基

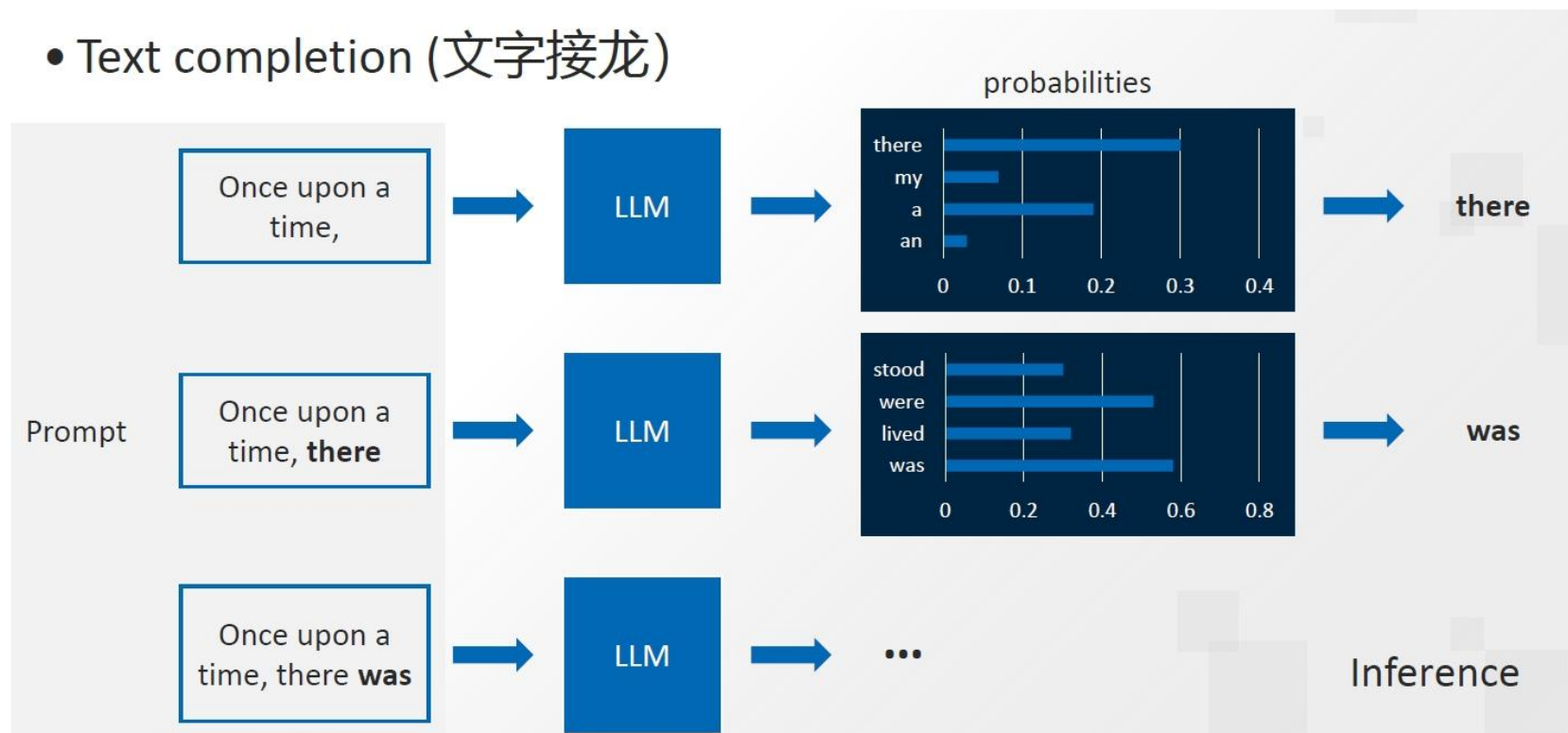
构建学习模型

Transformer：通过注意力机制捕捉序列中的长距离依赖性

CNN/GNN(图神经网络)：理解蛋白质的三维结构

GPT (Generative Pre-Trained Transformers) : 足够复杂的**预训练**后，用于**生成**...

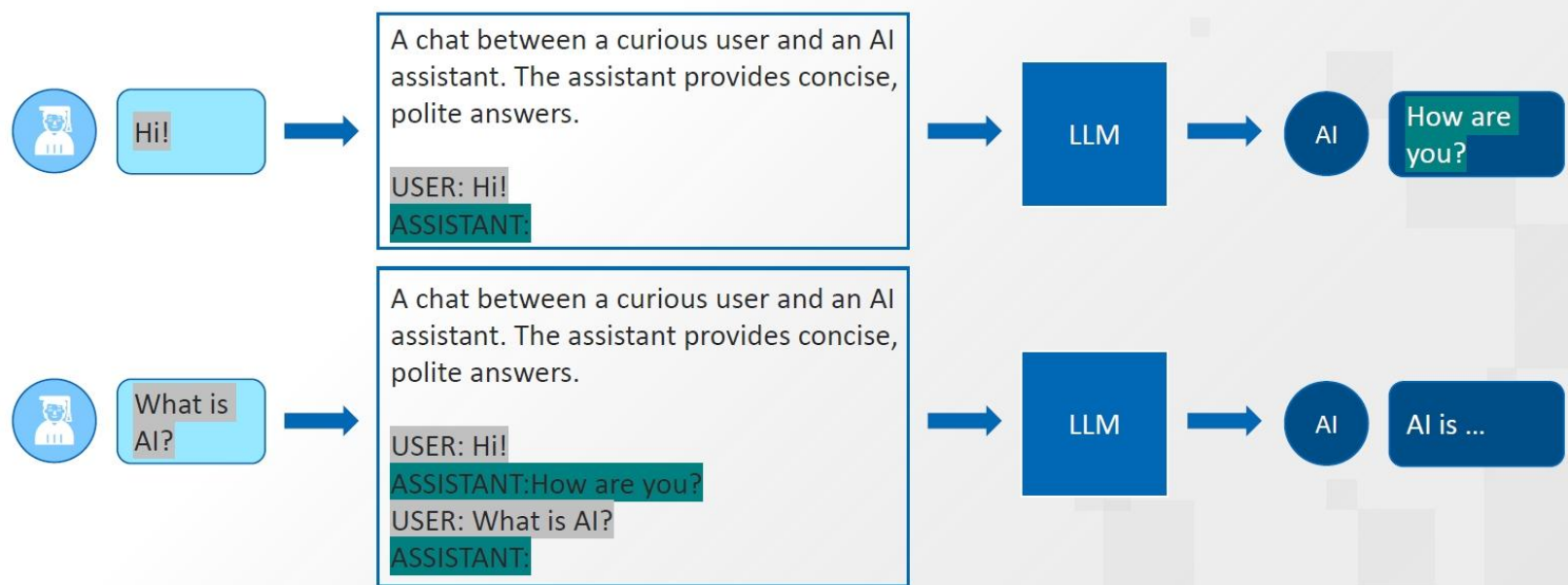
• Text completion (文字接龙)



如果大家还记得：
 $sky = F_1(\text{the clouds are in the})$
 $the = F_2(\text{the clouds are in})$
 $in = F_3(\text{the cloud are})$
 $are = F_4(\text{the clouds})$
 $clouds = F_5(\text{the})$

Chat: 如果把上轮交流看作上下文...

- For chat (multi-turn)



12:26

🌙 📶 11.0 KB/S 5G 66

✕ ...

早在 ChatGPT 面世之先，Sal Khan^Q（可汗学院创始人）就收到 OpenAI 的邀请，参加了一个视频会议，他甚至还必须为此签订了一份保密协议：会议上，他同当时还无人知晓的 GPT-4 一起做了一套生物试卷——这彻底改变了他对于 AI 在教育中价值的判断。

因为那份保密协议，那时的他还只能同极少数的人讨论他的想法，比如比尔·盖茨。“我们俩就像小孩子收获一个超酷玩具那般兴奋。”

而对于 OpenAI 这边，他们急需一个极好案例来证明这项技术的积极意义。所以他们找到 Sal Khan，希望 Khanmigo 能伴随 GPT-4 一起推出。

近日，果壳与可汗学院创始人 Sal Khan 进行了一场对话。得益于他新书（《教育新语》）的发布，他如今可以披露出更多当时同 OpenAI 协作的细节，以及他对于 AI 在教育中（尤其是未成年教育）应用的一些原则。

果壳

222 829 67 8

12:16

🌙 📶 5.00 KB/S 5G 65

✕ ...

Q: 所以你们的态度一下子变了，对于在此之前 AI 写的东西，你的评价是“毫无逻辑、毫无意义”。

A: 是的，GPT-3写的文章，你打眼一看好像还不错，但你真正读起来，你就会发现是挺有趣的，但没多大意义。

GPT-3.5的出现对很多人来说是一个惊喜。我记得 ChatGPT 发布的那天，是2022年11月30日，我给 Greg Brockman^Q发消息，我问他怎么回事，GPT-4理应还没推出，我们之前还签了保密协议，不能跟任何人提起，这才过几个月，你们就发布了。Greg 说，没有，我们只是在3.5模型的基础上放了一个聊天界面。这让几乎每个人都关注到了 OpenAI。

它表现比 GPT-3模型好得多，而且又有一个聊天界面让人们意识到大语言模型能做什么。之前人们不太知道，没有很好地利用大模型。

果壳

222 828 67 8

生成模型与推理大模型：生成内容 vs 生成逻辑

	生成模型（如GPT-4）	推理模型（如GPT-o1/o3）
能力定位	侧重通用自然语言处理和多模态能力，适合日常对话、内容生成、翻译以及图文、音频、视频等信息处理、生成、对话等。	侧重于复杂推理与逻辑能力，擅长数学、编程和自然语言推理任务。在生成模型的基础上通过RL等方法强化CoT能力而来
使用范围	支持文本、图像、音频乃至视频输入，可处理多种模态信息。	当前主要支持文本输入，不具备图像处理等多模态能力，幻觉通常比生成模型大。
应用场景	适合广泛通用任务，如对话、内容生成、多模态信息处理以及多种语言相互翻译和交流	适合需要复杂推理和逻辑分析的专业任务，如推导、编程和科学研究；在思路清晰度要求高的场景具有明显优势，比如采访大纲、方案梳理、解决方案。

OpenAI 的 Dota 2 人工智能智能体项目 OpenAI Five 已经经历了三年的发展。在 2019 年 4 月 13 日，OpenAI Five 成为了首个战胜了世界冠军战队的 AI 系统，但是当时 OpenAI 没有公开相关的论文和算法细节。近日，OpenAI 终于发布了描述该项目的论文《Dota 2 with Large Scale Deep Reinforcement Learning》。



OpenAI公开Dota2论文： 胜率99.4%

机器之心 2019-12-16

人工智能的长期目标是解决高难度的真实世界难题。为了实现这一目标，研究者在近几十年的时间里将游戏作为研究 AI 发展的基石。从双陆棋（1992）到国际象棋（1997）再到 Atari 游戏（2013），在 2016 年，AlphaGo 凭借深度强化学习和蒙特卡洛树搜索战胜了围棋世界冠军。近些年来，强化学习（RL）也在更多类型的任务上得到了应用，比如机器人操作、文本摘要以及《星际争霸》和《Minecraft》等视频游戏。

https://m.thepaper.cn/baijiahao_524769
3

OpenAI：以强化学习开局.....

成立于2015年12月美国旧金山，由小团队组成的非盈利性质的人工智能实验室，其目标是通过与其他机构和研究者的“自由合作”，向公众开放AI专利和研究成果。

大家猜猜是为了打破谁的技术垄断，从而获得马斯克、微软的认同？

Mastering the game of Go without human knowledge

- 2017.10.19

AlphaGo用到了上百万种人类专业选手的棋谱和基于自我对弈的强化学习。AlphaGo Zero则单纯基于与自己的对弈，而没有人类棋谱。通过近500万局自我对弈，AlphaGo Zero便能够超越人类并打败所有之前的AlphaGo版本

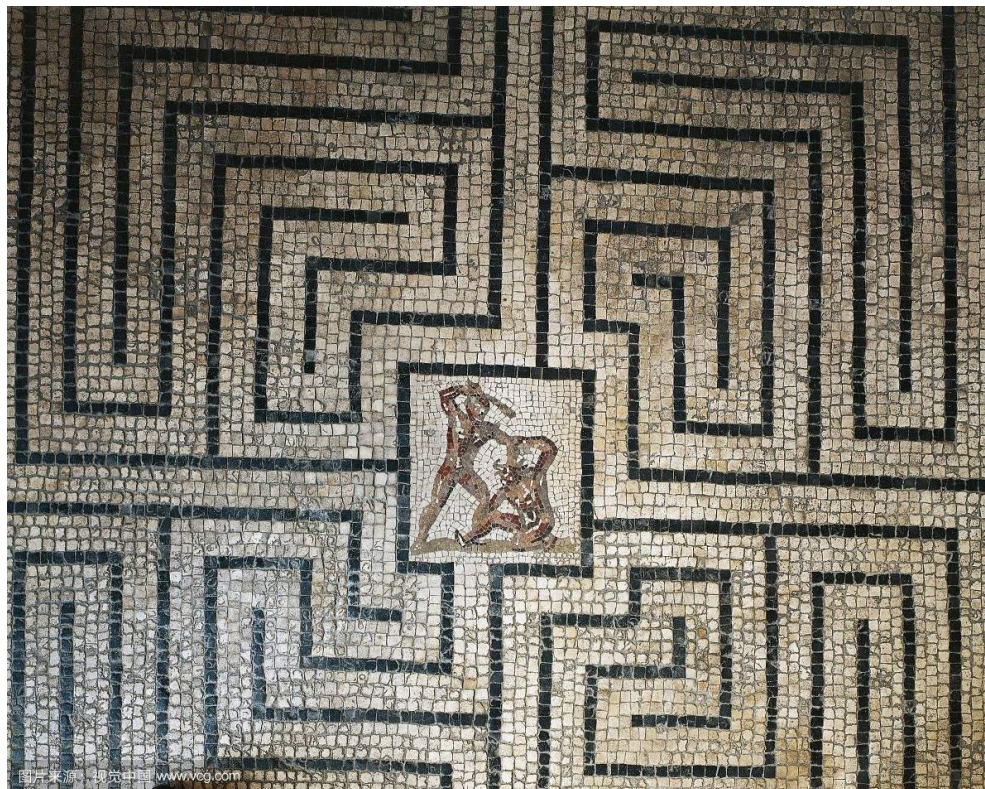
AlphaGo Zero 证明了即使在最具挑战的领域，纯强化学习的方法也是完全可行的：不需要人类的样例或指导，不提供基本规则以外的任何领域知识，使用**强化学习**能够实现超越人类的水平。



DeepMind, London, UK

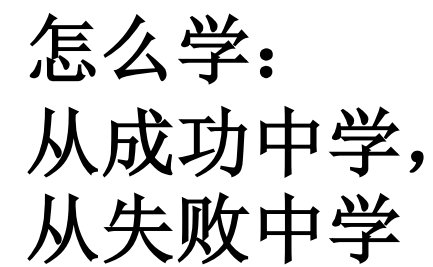
几十

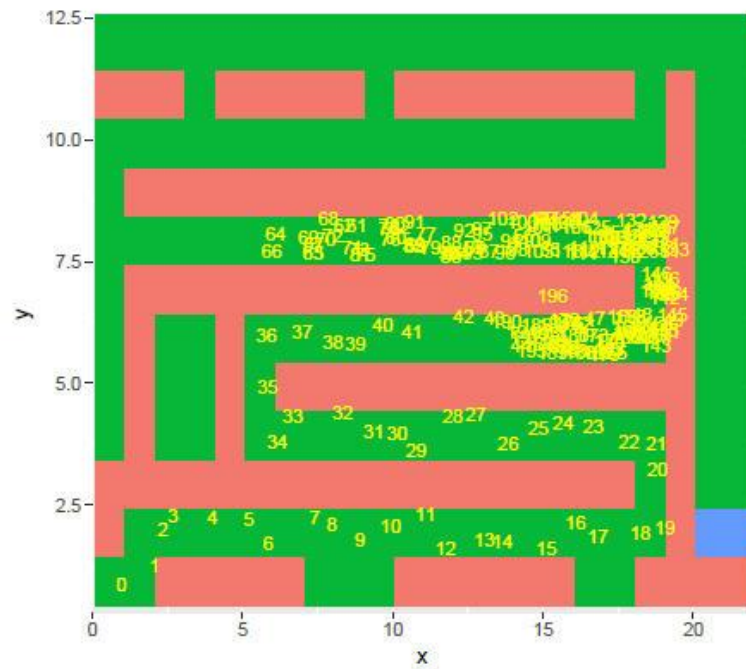
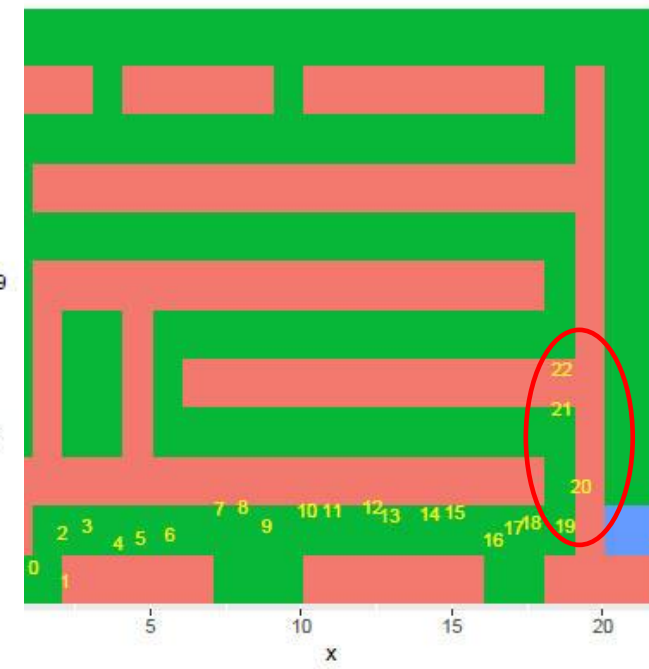
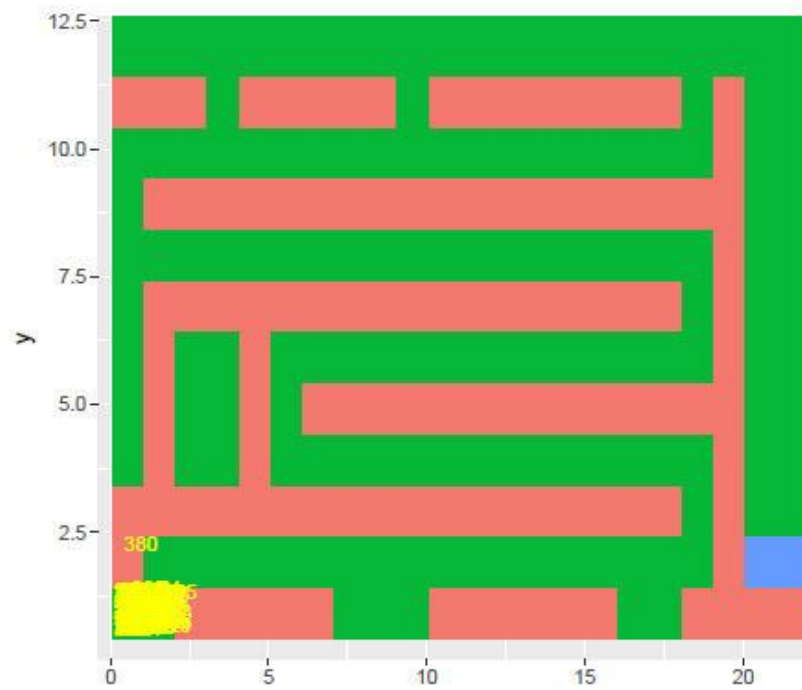
比如：米诺斯的迷宫



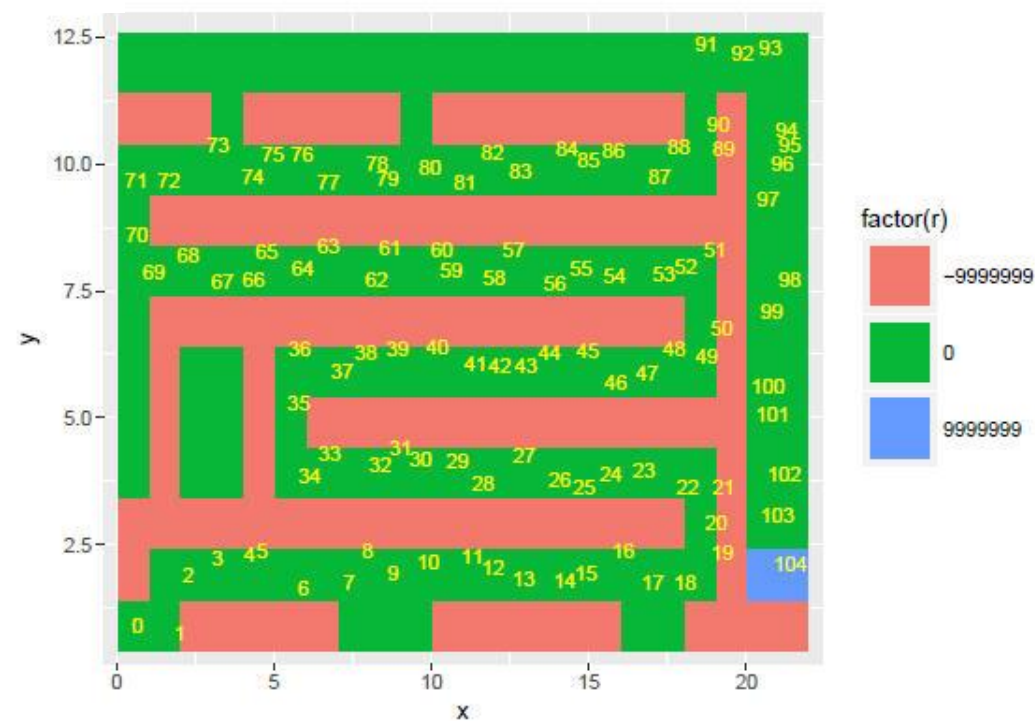
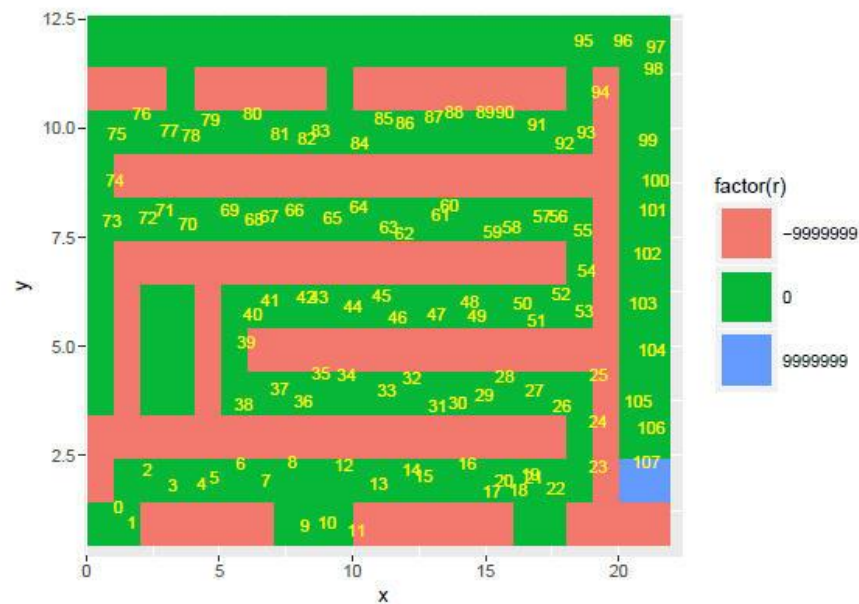
图片来源：视觉中国 www.veg.com

米诺斯是宙斯和欧罗巴的儿子，欧罗巴被天后赫拉排挤迫害，来到克里特岛，与当地国王结婚。之后，米诺斯成为克里特国王。米诺斯因为得罪了海神波塞冬，波塞冬便诱使米诺斯的妻子爱上一只公牛，生下了牛首人身的米诺陶洛斯。后来，因为雅典人杀死了米诺斯的另一个儿子，米诺斯求宙斯给雅典施加瘟疫，雅典被迫每年选送7对童男童女去供奉怪物米诺陶洛斯。

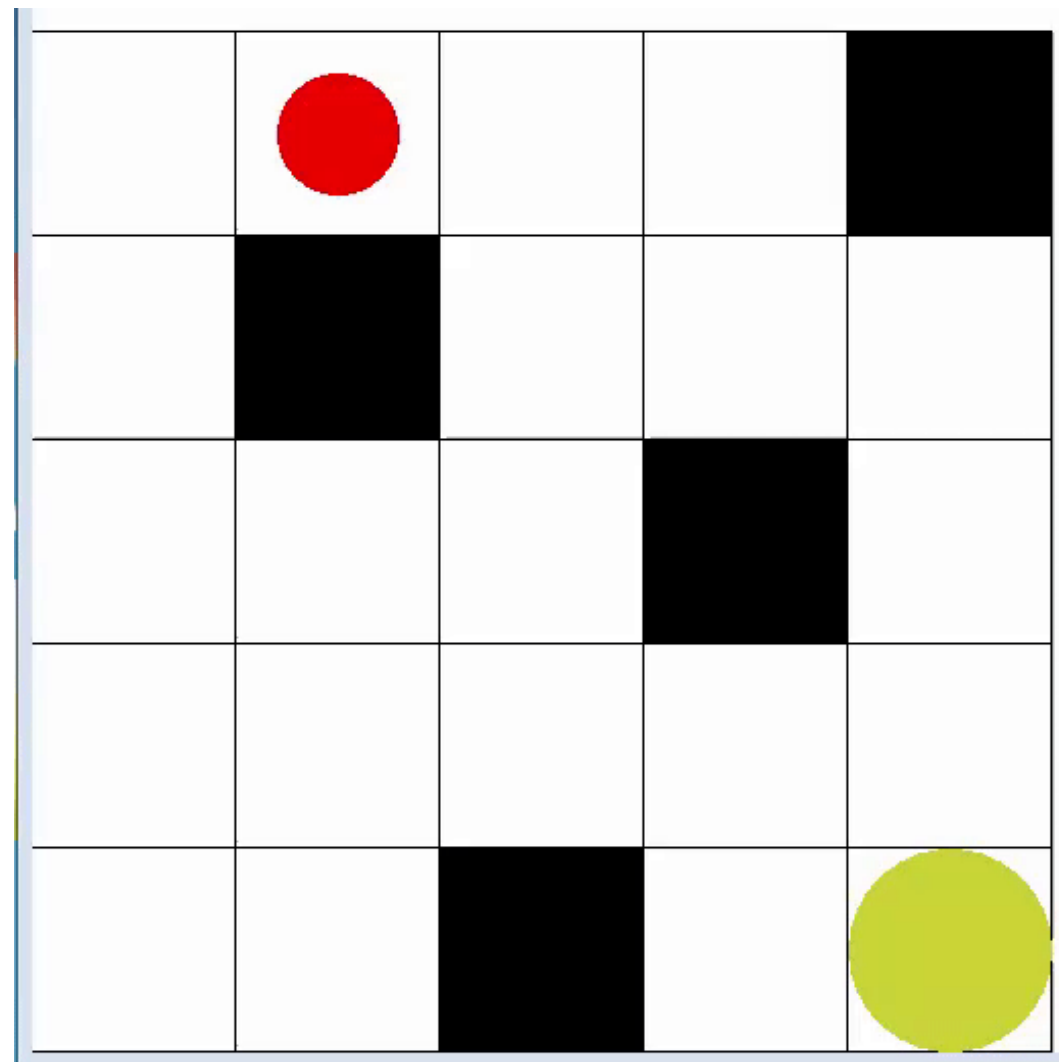
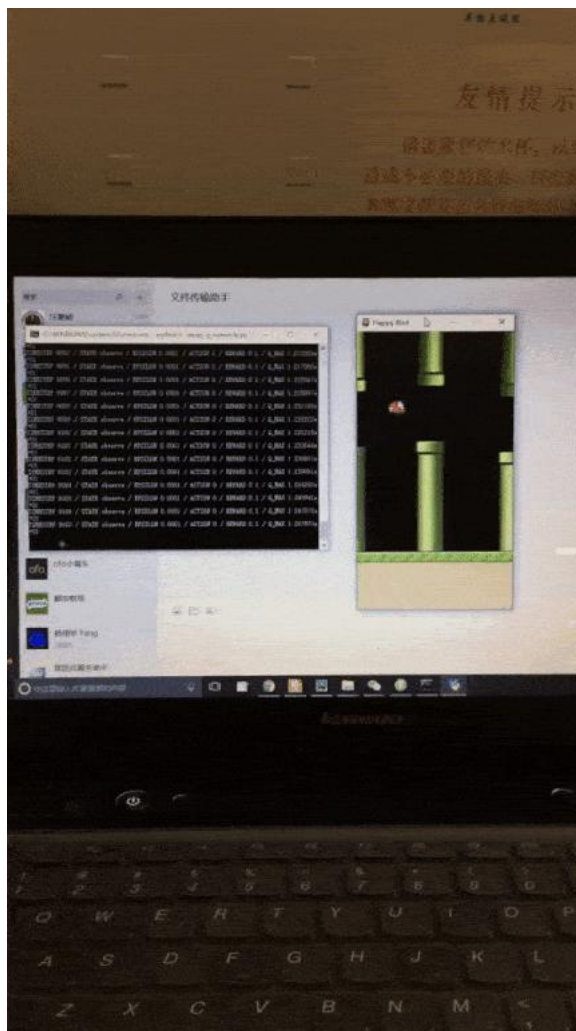
[illegible]



o o o o o o



愤怒的小鸟



如何理解强化学习思路？

两句话

不教答案，只告诉对不对

通过反复尝试学会怎么做

比如人类幼崽的牙牙学语、骑自行车

无法精确告诉他每一个动作，只能说：好 / 不好

他是怎么学会的

不断尝试，失败得到负反馈，成功得到正反馈

最终形成肌肉反应

所以

强化学习 = 允许犯错的学习

通过结果好坏调整行为

学的是选择，不是答案

试得足够多，就会变好（不断修正预期的过程）

什么时候用？

规则说不清，结果看得出，可以反复尝试

下棋、打游戏、走迷宫、训练宠物（包括人类幼崽）

不是所有学习都是强化学习

背单词、考试（GPT）

强化学习的投资逻辑

多期折现收益下的：

新收益 = 本期收益 + $\alpha \times (\text{未来的实际结果} - \text{原来预期})$

$$V_{\text{new}}(s) = V_{\text{old}}(s) + \alpha \left[\underbrace{r + \gamma V(s')}_{\text{未来的实际结果}} - \underbrace{V_{\text{old}}(s)}_{\text{原来预期}} \right]$$

未来的实际结果 = $r + \gamma V(s')$

r : 当期收益

$V(s')$: 未来长期可能拿到的收益

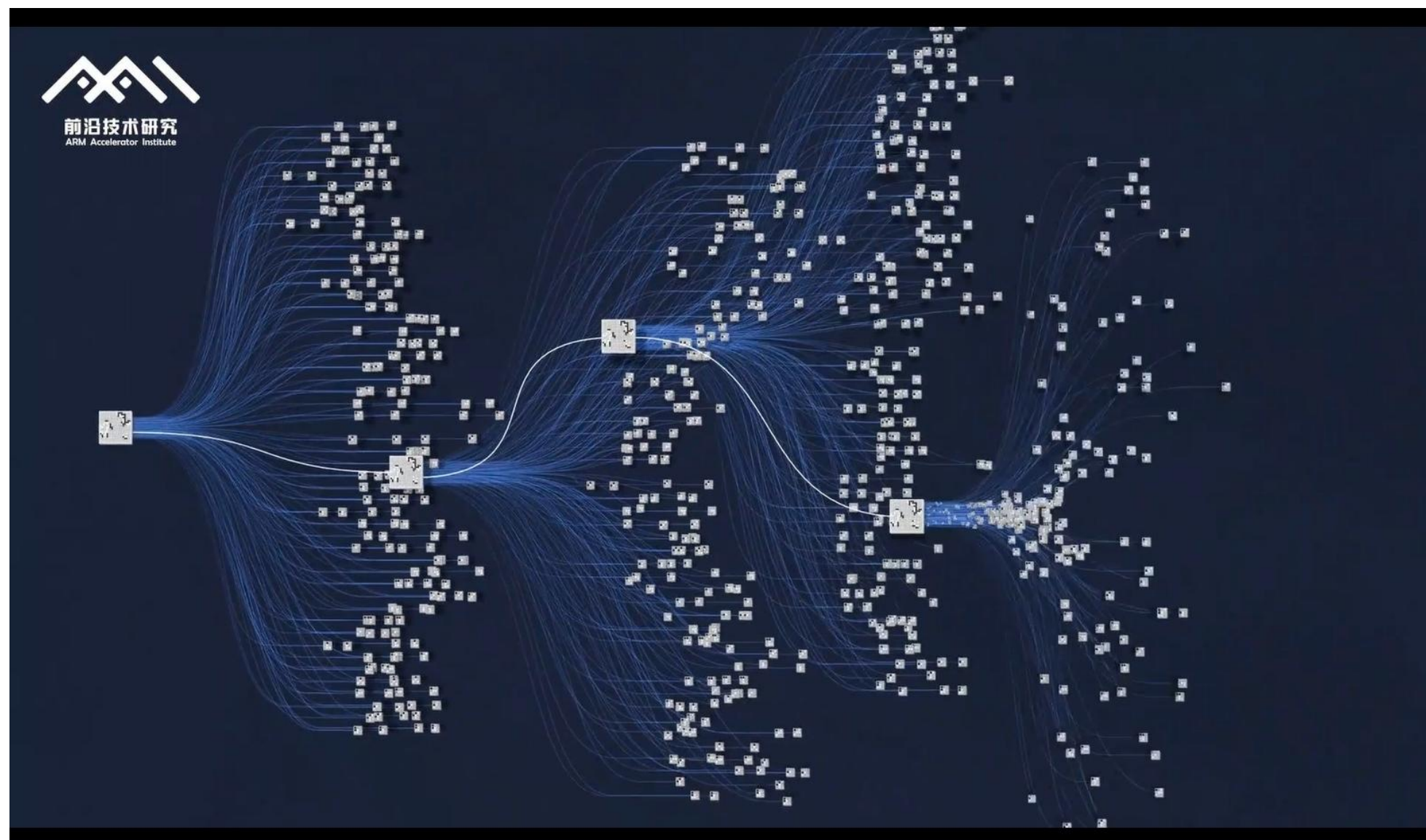
γ : 未来收益的折现因子

新收益	$V_{\text{new}}(s)$
本期收益	$V_{\text{old}}(s)$
未来的实际结果	$r + \gamma V(s')$
原来预期	$V_{\text{old}}(s)$
折现因子	γ
调整幅度	α (学习率)

对应到GPT

新看法/说法 = 旧看法/说法 + $\alpha \times (\text{不断生成的描述} - \text{原来预期的描述})$

真实结果：人类主人喜欢不；原来预期：我以为人类主人喜欢不



2019-3-19

Nvidia推出全新DRIVE Constellation自动驾驶汽车仿真平台。基于该平台，在云端就能虚拟仿真各种自动驾驶场景——不用再路测数百万公里了。从常规驾驶，到各种罕见的危机情况，都能在仿真中实现。

数据中心方案包括两个并排服务器：第一台服务器——DRIVE Constellation Simulator，从虚拟汽车生成传感器输出。第二台服务器——DRIVE Constellation Vehicle，包含DRIVE AGX Pegasus AI车载电脑。DRIVE AGX Pegasus接收传感器数据，做出决定，然后将车辆控制命令发送回模拟器。

与丰田开展基于仿真平台的合作



逆推是“不正常”的思考方式：人更习惯顺推，但由于没有足够强大的计算能力（意志力）而失败

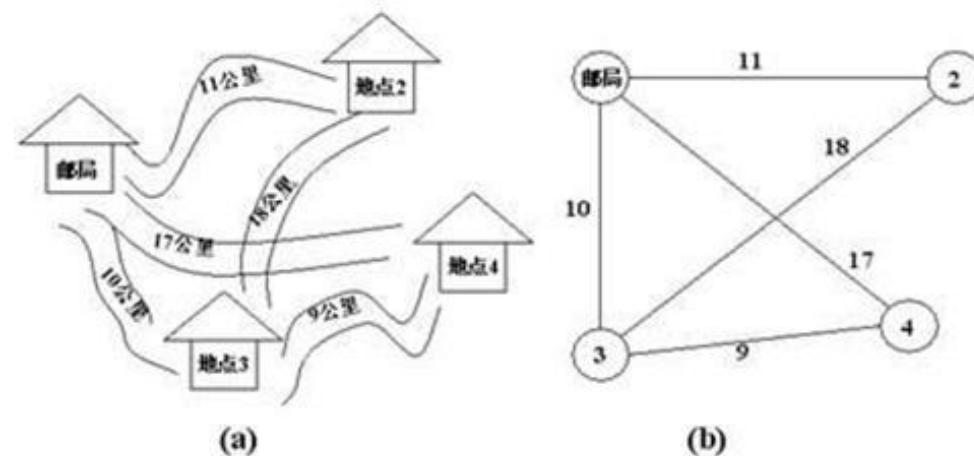


图1.3 邮递员送信问题

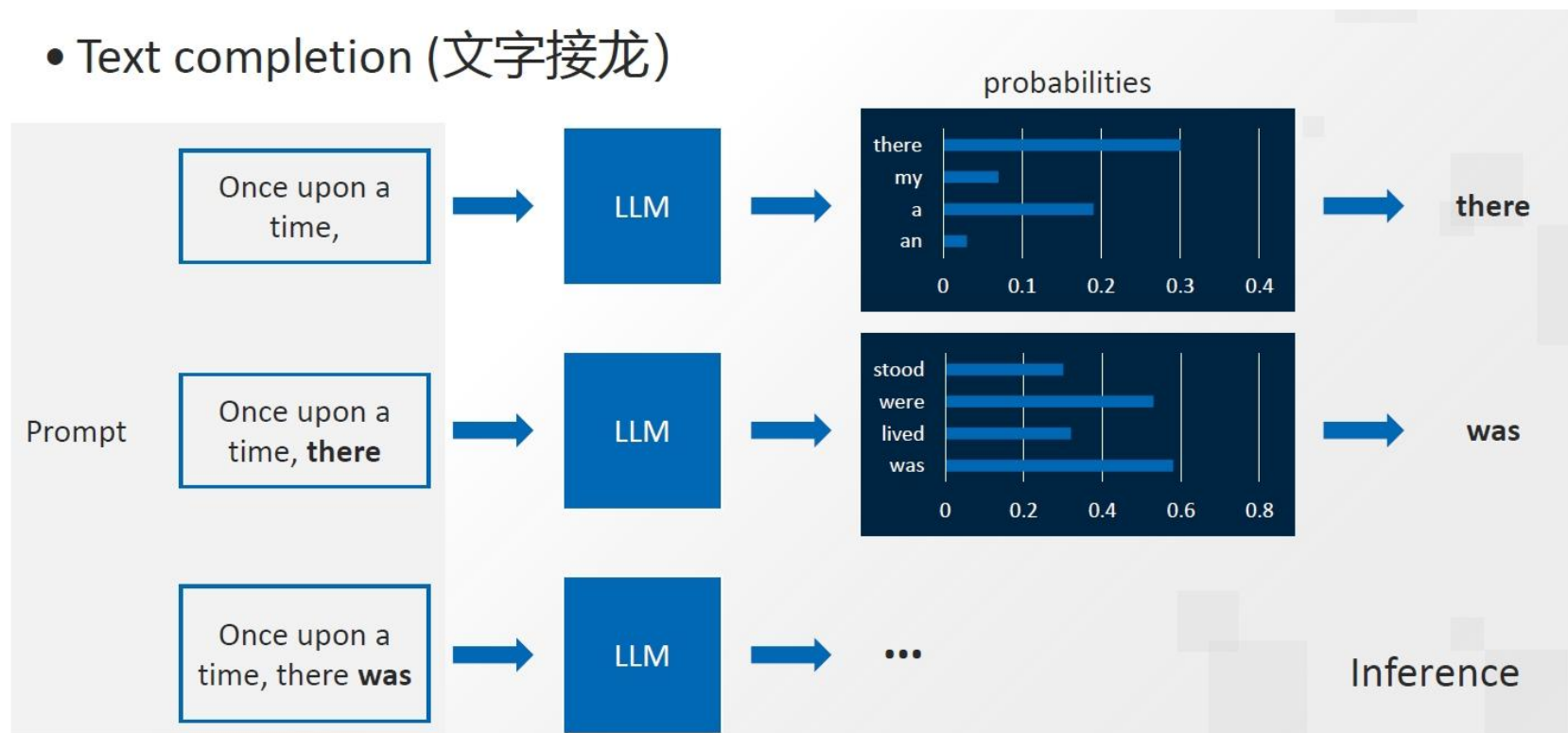
中国数学家管梅谷1962年提出中国邮递员问题



1990年至1995年任复旦大学运筹学系主任

增强学习如何用在语言学习?

• Text completion (文字接龙)



如果大家还记得:
 sky=F₁(the clouds are in the)
 the= F₂(the clouds are in)
 in=F₃(the cloud are)
 are=F₄(the clouds)
 clouds=F₅(the)

如何理解强化学习方式

第一步：训练奖励模型，教LLM学习更高质量地回答问题 (几十万-几百万)

老公和老爸掉进河里，我该救哪个？



- A** 都救
- B** 这是一个极端情况，你先冷静....
- C** 不应该把救哪个人作为一个简单的选择
- D** 我只是一个AI模型

这时候，伟大的人类老师再次上场了

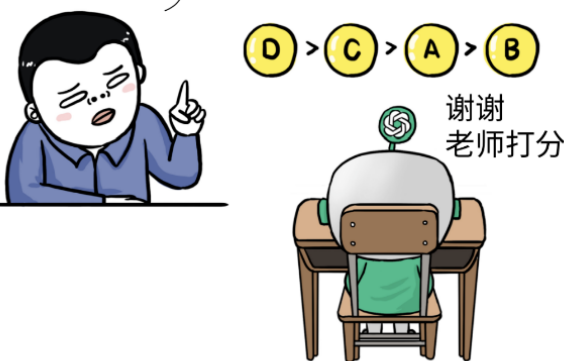
人类老师对N个输出答案质量
进行综合排序

排序的参考维度有很多
比如：关联性、法律法规、暴力、种族歧视等

你现在终于开窍了，
给你答案评个分吧！

D > C > A > B

谢谢老师打分



Step 3

Optimize a policy against the reward model using the PPO reinforcement learning algorithm.

A new prompt is sampled from the dataset.

Write a story about others.

The PPO model is initialized from the supervised policy.

PPO

The policy generates an output.

Once upon a time...

The reward model calculates a reward for the output.

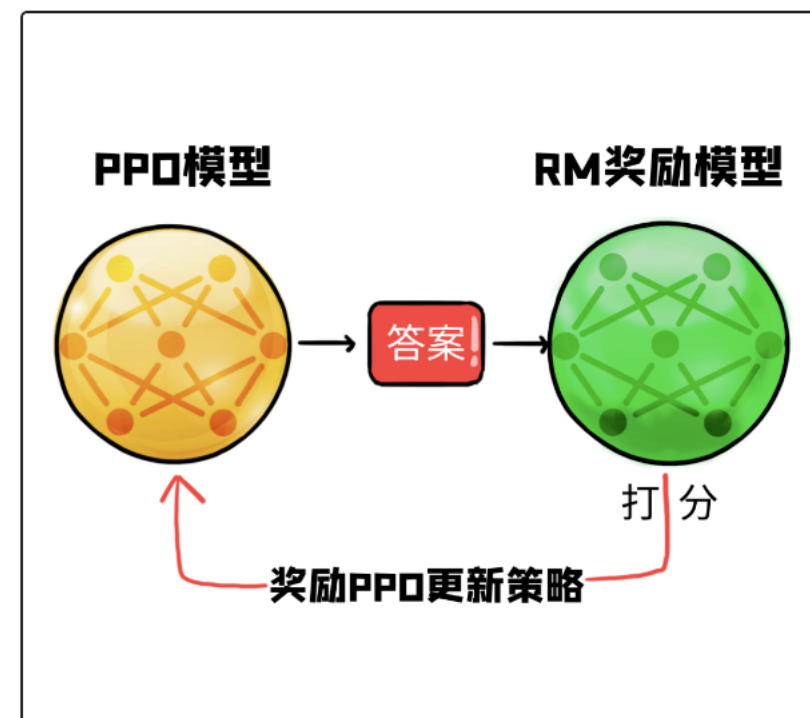
RM

The reward is used to update the policy using PPO.

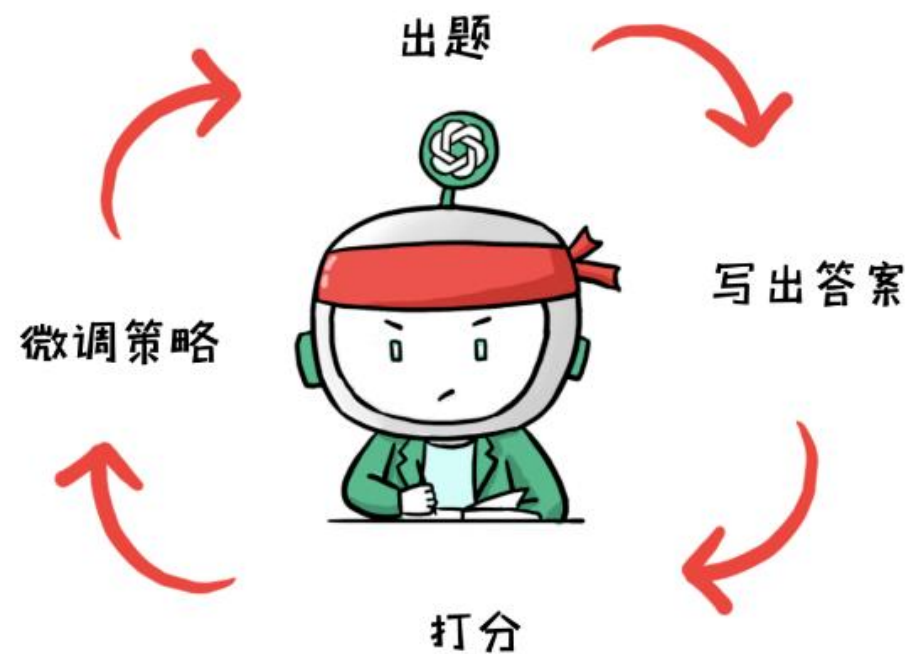
r_k

如何理解强化学习方式

第二步：不断自我生成新问题，让LLM持续自我进化（几百万~问题）



不断循环、强化学习



Training language models to follow instructions with human feedback

Long Ouyang* Jeff Wu* Xu Jiang* Diogo Almeida* Carroll L. Wainwright*
 Pamela Mishkin* Chong Zhang Sandhini Agarwal Katarina Slama Alex Ray
 John Schulman Jacob Hilton Fraser Kelton Luke Miller Maddie Simens
 Amanda Askell† Peter Welinder Paul Christiano*†
 Jan Leike* Ryan Lowe*

OpenAI

Abstract

Making language models bigger does not inherently make them better at following a user's intent. For example, large language models can generate outputs that are untruthful, toxic, or simply not helpful to the user. In other words, these models are not *aligned* with their users. In this paper, we show an avenue for aligning language models with user intent on a wide range of tasks by fine-tuning with human feedback. Starting with a set of labeler-written prompts and prompts submitted through the OpenAI API, we collect a dataset of labeler demonstrations of the desired model behavior, which we use to fine-tune GPT-3 using supervised learning. We then collect a dataset of rankings of model outputs, which we use to further fine-tune this supervised model using reinforcement learning from human feedback. We call the resulting models *InstructGPT*. In human evaluations on our prompt distribution, outputs from the 1.3B parameter InstructGPT model are preferred to outputs from the 175B GPT-3, despite having 100x fewer parameters. Moreover, InstructGPT models show improvements in truthfulness and reductions in toxic output generation while having minimal performance regressions on public NLP datasets. Even though InstructGPT still makes simple mistakes, our results show that fine-tuning with human feedback is a promising direction for aligning language models with human intent.

1 Introduction

Large language models (LMs) can be “prompted” to perform a range of natural language processing (NLP) tasks, given some examples of the task as input. However, these models often express unintended behaviors such as making up facts, generating biased or toxic text, or simply not following user instructions (Bender et al., 2021; Bommasani et al., 2021; Kenton et al., 2021; Weidinger et al., 2021; Tamkin et al., 2021; Gehman et al., 2020). This is because the language modeling objective

*Primary authors. This was a joint project of the OpenAI Alignment team. RL and JL are the team leads. Corresponding author: lowe@openai.com.

让语言模型变大并不意味着它们能更好地遵循用户的意图。例如，大型语言模型可以产生不真实的、有毒的或根本对用户没有帮助的输出。换句话说，这些模型没有与用户保持一致。在本文中，我们展示了一个途径，通过人类反馈的微调，在广泛的任务中使语言模型与用户的意图保持一致。从一组标签员写的提示语和通过 OpenAI API提交的提示语开始，我们收集了一组标签员演示的所需模型行为的数据集，我们利用监督学习对 GPT-3 进行微调。然后，我们收集模型输出的排名数据集，我们利用人类反馈的强化学习来进一步微调这个监督模型。我们把产生的模型称为 InstructGPT。在人类对我们的提示分布的评估中，尽管参数少了 100 倍，但 1.3B 参数的 InstructGPT 模型的输出比 175B 的 GPT-3 的输出更受欢迎。此外，InstructGPT 模型显示了真实性的改善和有毒输出生成的减少，同时在公共 NLP 数据集上有最小的性能回归。尽管 InstructGPT 仍然会犯一些简单的错误，但我们的结果表明，利用人类反馈进行微调是使语言模型与人类意图相一致的一个有希望的方向。

We focus on *fine-tuning* approaches to aligning language models. Specifically, we use reinforcement learning from human feedback (RLHF; Christiano et al., 2017; Stiennon et al., 2020) to fine-tune GPT-3 to follow a broad class of written instructions (see Figure 2). This technique uses human preferences as a reward signal to fine-tune our models. We first hire a team of 40 contractors to label our data, based on their performance on a screening test (see Section 3.4 and Appendix B.1 for more details). We then collect a dataset of human-written demonstrations of the desired output behavior on (mostly English) prompts submitted to the OpenAI API³ and some labeler-written prompts, and use this to train our supervised learning baselines. Next, we collect a dataset of human-labeled comparisons between outputs from our models on a larger set of API prompts. We then train a reward model (RM) on this dataset to predict which model output our labelers would prefer. Finally, we use this RM as a reward function and fine-tune our supervised learning baseline to maximize this reward using the PPO algorithm (Schulman et al., 2017). We illustrate this process in Figure 2. This procedure aligns the behavior of GPT-3 to the stated preferences of a specific group of people (mostly our labelers and researchers), rather than any broader notion of “human values”; we discuss this further in Section 5.2. We call the resulting models *InstructGPT*.

我们专注于调整语言模型的**微调**方法。具体来说，我们使用来自人类反馈的强化学习来微调GPT-3。这项技术使用人类的偏好作为奖励信号来微调我们的模型。

我们首先雇用了由**40名**承包商组成的团队，根据他们在筛选测试中的表现，为我们的数据贴上标签。。。并使用它来训练我们的监督学习基线

接下来，我们收集了一个由人类标记的数据集，对我们的模型在更大的提示集上的输出进行比较。我们在这个数据集上训练一个奖励模型（RM），以预测我们的标注者会更喜欢哪个模型的输出。

最后，我们使用这个RM作为奖励函数，并微调我们的监督学习基线，以使用PPO算法最大化这个奖励。这个过程使GPT-3的行为与特定人群（主要是我们的标签人员和研究人员）的声明偏好相一致，而不是任何更广泛的“人类价值”概念。。。我们把产生的模型称为**InstructGPT**。

如何将强化学习用于企业管理工作？

促销（如发放优惠券，等同于定价策略），企业的日常工作之一
天天发、周周发.....
但顾客就是不来呀

传统思路：精准、再精准
有没有可能

对于重复促销：顾客还没消耗完（消费疲劳，甚至反消费/忽略）
对于多样化促销：不喜欢/不需要

操纵手段（假设价格和选品都由市场部、财务部完成了）
是否发放，以及有效时间



方案需求：如何基于细粒度和完整的用户和企业长期数据，制定个性化和动态的促销/定价方案，以优化企业的长期回报？

如何应用强化学习?

将促销看作向消费者学习的过程

消费者每次用/不用消费券就是一次反馈，即强化学习的机会

类似从撞墙（不撞墙）学到驾驶规则/行走路径；从棋局生/死（提子）中学到下棋手段；大模型从用户反馈中学到内容

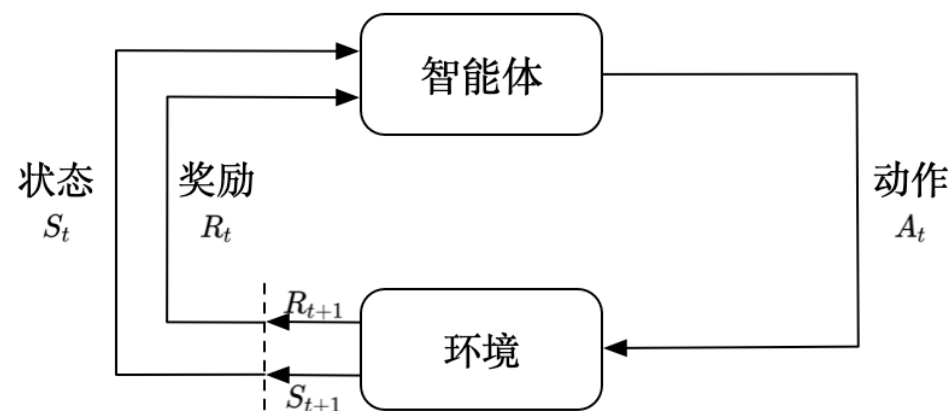
把大象放到冰箱里

了解智能体（agent）自身的行为规则（精准画像）

确认动作（投放，用户接受）

确认收益目标（奖惩动作）

让计算机开始学习吧



国内一在线生鲜零售商运营数据

2,012名消费者17个月的363,946条消费记录

对谁发放优惠券、发放什么样的优惠券才能有效提高销售额?

促销类别:

短期优惠券: 主要用于刺激用户冲动消费

长期优惠券: 维护用户忠诚度和提升体验

传统促销方法 VS 基于深度增强学习的动态促销

策略1: 等概率随机选择短期优惠券、长期优惠券和不发放优惠券

策略2: 长短期优惠券的发放概率均为50%

策略3: 根据统计预测方法来决策

策略4: 根据机器学习方法来决策

策略5: 增强学习策略



状态类别	状态	状态定义
外部环境因素	天气	天气, 0: 晴天或多云, 1: 阴天, 2: 恶劣天气如雨天、下雪
	节假日或周末	是否为节假日或周末, 1: 是, 0: 否
历史消费特征(购买行为)	购买近度	距离上一次/上上次购买的时间间隔(天)
	购买频度	累计购买频次
	购买值度	上一次/上上次购买订单的金额(元)
历史消费特征(促销行为)	促销近度	距离上一次/上上次收到优惠券的时间间隔(天)
	促销频度	累计收到优惠券的频次
	促销反馈	对近3次优惠券的反馈(3个哑变量, 1: 购买, 0: 无购买)

优惠券	发放数量	平均折扣额	召回率	平均销售额(std)
短期	17,503	16.86	3.96%	106.70 (34.36)
长期	18,394	19.78	8.61%	86.62 (45.82)

基于增强学习的促销策略：

比静态促销方案提高约18%的收益
消费者购买频率与购买金额也分别提升7%和9%
促销活动降低10%，缓解促销带来的消费疲劳
有点像休渔期 ☺



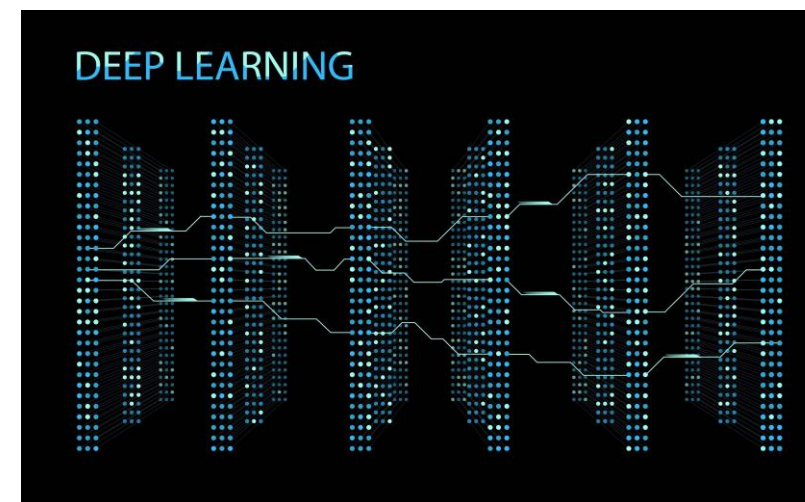
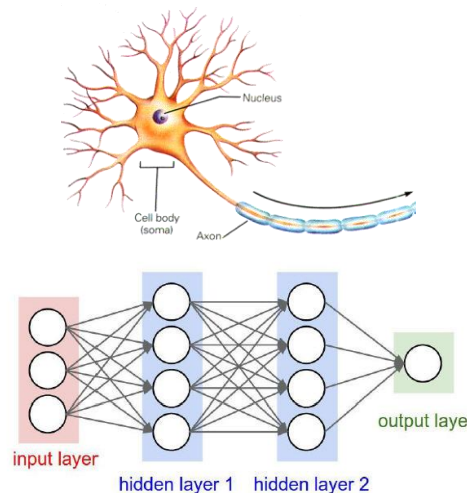
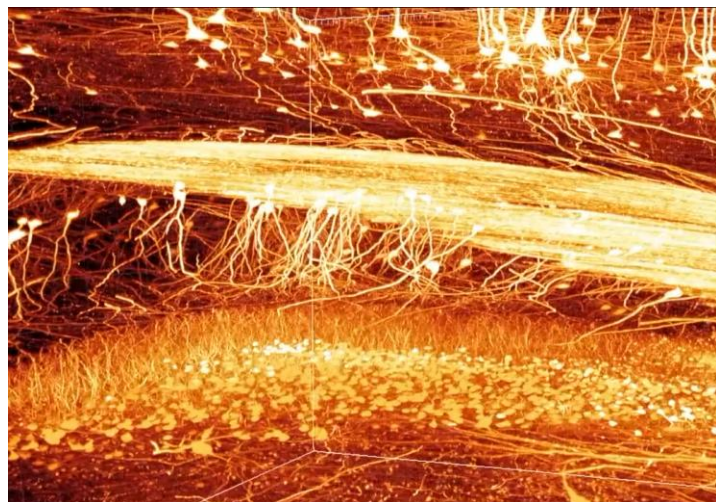
总体促销效果对比

促销定价策略	总销售额	购买频率	购买金额
1. 全随机（可不发）	100. 0%	100. 0%	100. 0%
2. 全随机发	119. 6%	112. 7%	104. 8%
3. 统计预测	137. 2%	139. 3%	97. 2%
4. 机器学习预测	118. 8%	125. 3%	94. 1%
5. 强化学习	154. 9%	146. 7%	106. 1%

促销比例对比

促销定价策略	促销方案比例		
	不发放优惠券	短期优惠券	长期优惠券
1. 全随机（可不发）	33%	33%	33%
2. 全随机发	0%	50%	50%
3. 统计预测	11%	14%	76%
4. 机器学习预测	5%	7%	87%
5. 强化学习	21%	39%	40%

当前的智能：大力出奇迹

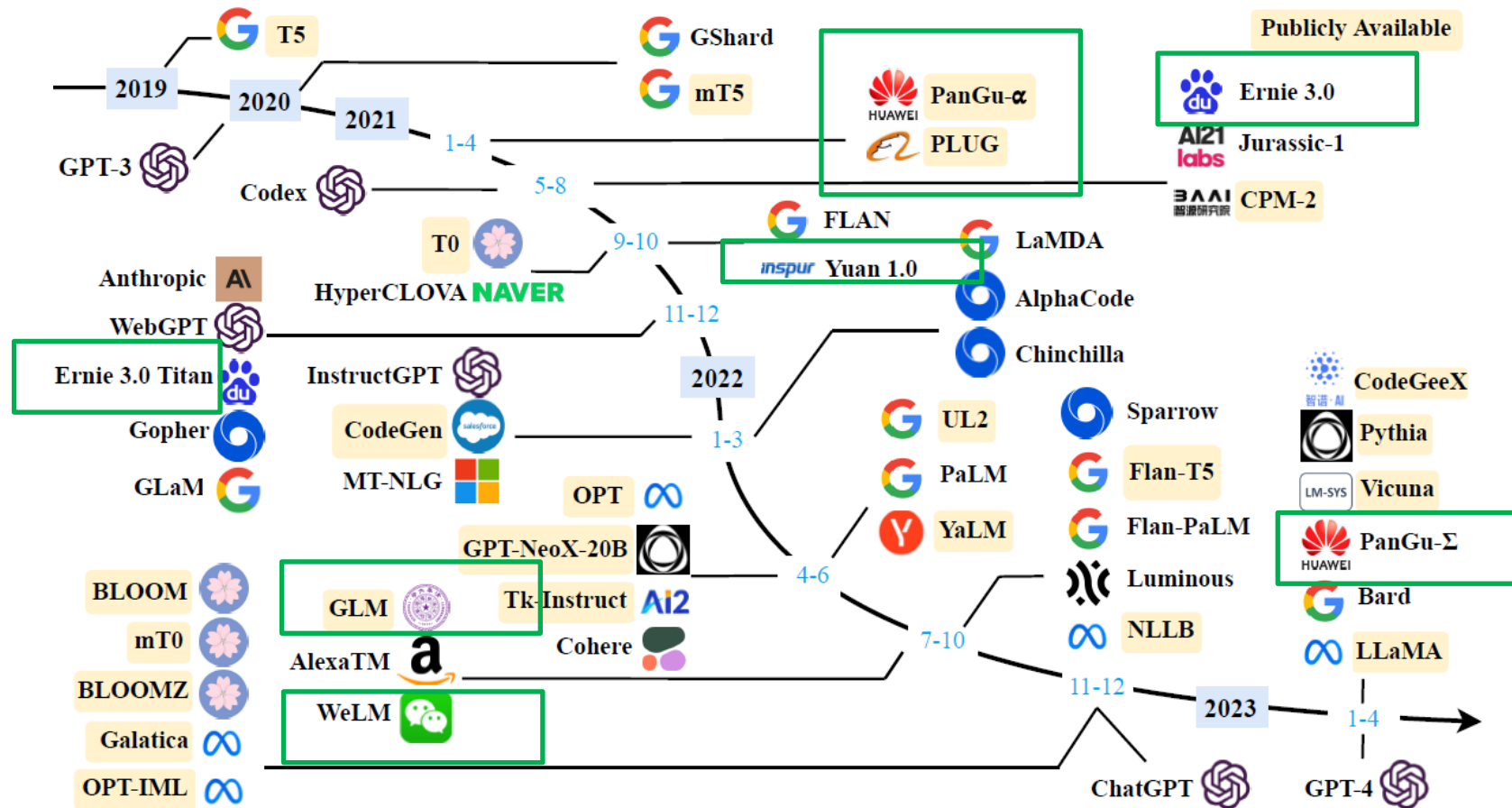


事实1：迄今通过深度学习实现的智能是模仿/应用人类智能的运行方式

人工智能：模拟、延伸和扩展人的智能

事实2：人类大脑的神经元约为900亿，大模型的‘智能’涌现发生在神经网络模型参数达到百亿到千亿级别

这是巧合还是人类智能产生的特征？



Deepseek历代版本

V1: 2024.1.5

规模: 67B

路线: LLaMA架构, 95层神经网络

效果: LLaMA-2 70B, GPT-3.5

成本: 300.6K H800 GPU小时

关注度: 78 (till 25/2/16)

V2: 2024.6.19

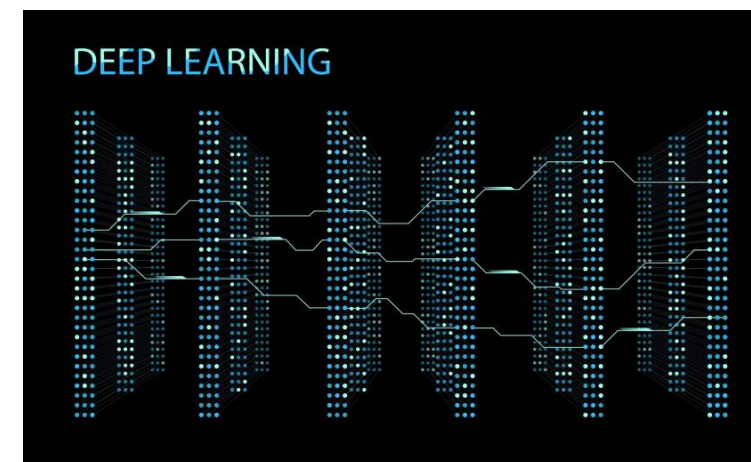
规模: 236B, 60层网络

路线: Transformer架构, 专家混合 (Mixture-of-Experts, MoE)
+ 多头潜在注意力 (Multi-head Latent Attention, MLA)

效果: Qwen1.5-72B, LLaMA3-70B, and Mixtral 8x22B

成本: 172.8K H800 GPU小时

关注度: 117



Deepseek历代版本：增强学习的突破

V3: 2024.12.26

规模: 671B

路线: Transformer架构 (同上), 61层网络

效果: Qwen2.5-72B, Llama-3.1-405B, GPT-4o, Claude-3.5

成本: 2048张NVIDIA H800卡, 2.788M小时 (2048*24*56),
按2刀/小时计算=\$5.576M

Llama-3.1-405B: 39.3m小时 H100 (20倍于DS)

关注度: 63

R1: 2025.1.22

规模: 671B

路线: Transformer架构 (同上), 基于V3进行增强学习 (Group Relative Policy Optimization, GRPO)

效果: OpenAI-o1

成本: 没披露

关注度: 63

备注: H800的性能约等于H100的50%-70%, 价格相当 (3.5万美金)

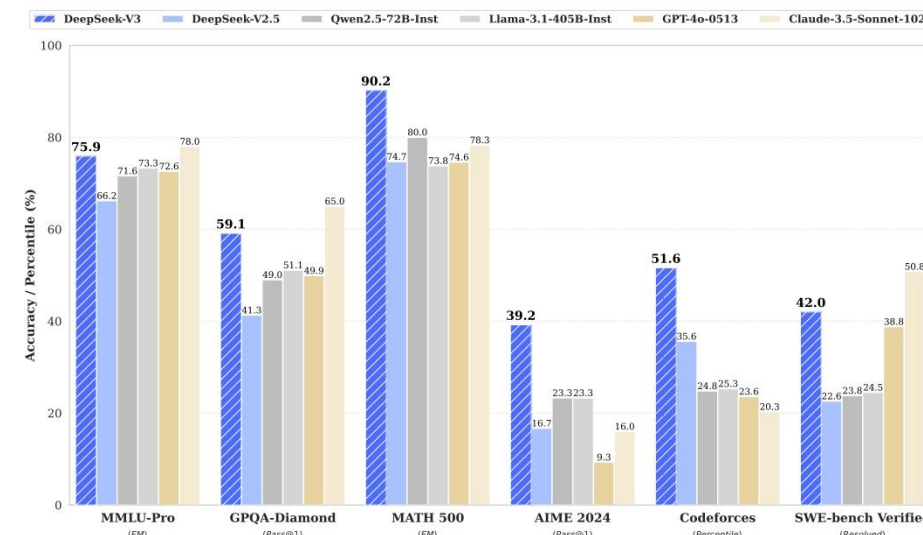


Figure 1 | Benchmark performance of DeepSeek-V3 and its counterparts.

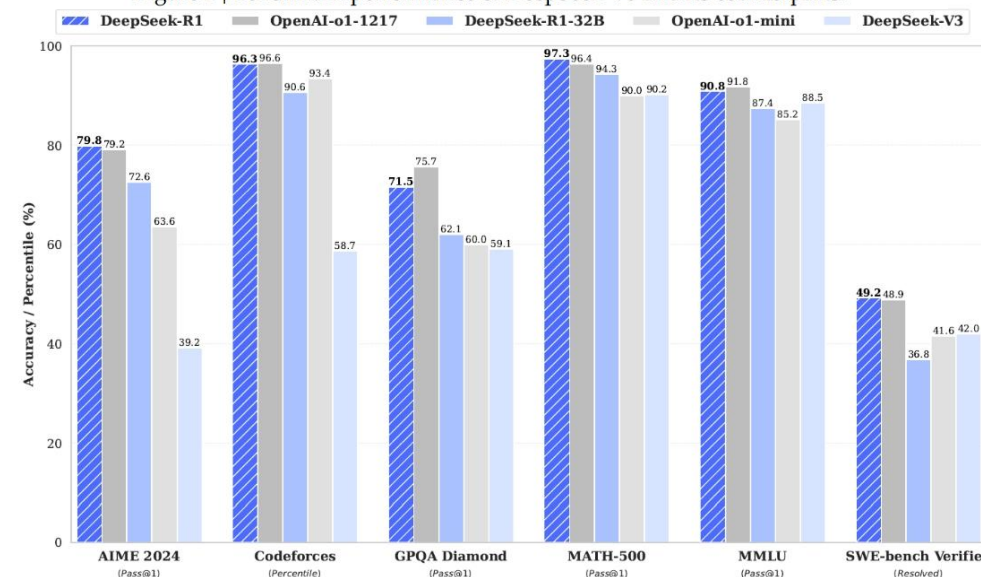


Figure 1 | Benchmark performance of DeepSeek-R1.

Deepseek技术说人话版

混合专家模型 (MoE)：采用MoE架构，通过动态选择最适合输入数据的专家模块进行处理，**提升推理效果**

说人话：根据数据类型分别进行增强学习（像不像之前学到的模型融合）

强化学习驱动 (RL)：在训练中大规模应用强化学习，将传统的PPO替换为GRPO训练算法，**降本提速**

说人话：结果不要太精确，误差可以大些

多头潜在注意力机制 (MLA)：MLA通过低秩压缩减少Key-Value缓存，**降本提速**

说人话：二向箔，把矩阵压成一条线

无辅助损失的负载均衡策略 (EP)：在不对优化目标产生干扰的前提下，实现各个专家的负载均衡，避免了某些专家可能会被过度使用，而其他专家则被闲置的现象，**降本提速**

说人话：大树 vs 小树森林

多Token预测 (MTP)：采用多Token预测，**提速**

说人话：预测下一个词 vs 预测下n个词

Deepseek技术说人话版



FP8混合精度训练：在关键计算步骤使用FP16高精度，其他模型层使用FP8低精度，**降本提速**

说人话：粗略点就好

长链推理技术（TTC）：模型支持数万字的长链推理，可逐步分解复杂问题并进行多步骤逻辑推理，**提速**

说人话：分段理解再组合，比完整理解省力（但可能会出错）

并行训练策略（HAI）：16 路流水线并行(Pipeline Parallelism, PP)、跨8 个节点的64 路专家并行(Expert Parallelism, EP)，以及数据并行(Data Parallelism, DP)，**提速**

说人话：甘特图

通讯优化DualPipe：高效的跨节点通信内核，利用IB 和NVLink带宽，减少通信开销，**降本**

说人话：团队协作时少说废话

混合机器编程（PTX）：部分代码直接进行使用PTX编程提高执行效率，并优化了一部分算子库，**提速**

说人话：你知道吗，c语言比python快耶

DS全家桶的效果

低成本高效率

DeepSeek-V3的训练成本为557.6万美元，仅为OpenAI的GPT-4o等领先闭源模型的3%-5%

效果低于O1/O3，但毕竟走出来了

没有突破上限

对不重视/没信心/缺资源的组织起到了引领作用

商业 vs 技术

Sora、LLM的训练和部署成本估算

Sora

14000张H100运行一个月

H100市场价约3.5万美元，不考虑电费，需要购置硬件4.9亿美元

H100功耗700W， $0.7 \times 14000 \times 24 \times 30 = 700$ 万度

Llama-3-405B

16000张H100，39.3m小时

硬件：5.6亿，电费：2751万度

Deepseek-V3

2048张H800运行56天

H800市场价约3.5万美金（国内最高到6万），硬件费0.72亿美金

H800功耗350W， $0.35 \times 2048 \times 24 \times 56 = 96$ 万度

即便如此

3月1日, DeepSeek官方通过社交媒体账号公布, 理论上成本利润率为545%

- 平均占用226.75*8的H800 GPU。如GPU租赁成本为2美金/小时, 总成本为\$87,072/天
- 输入+输出token总数为776B, 按R1定价, 一天收入为\$562,027, 成本利润率545%

如果我们把LLM看作产品 (而非技术)

- 性价比高的产品 vs. 高性能 (高成本) 产品

Deepseek真正价值：开源

开源生态

DeepSeek采用开源策略，使用最为开放的MIT开源协议，吸引了大量开发者和研究人员，推动了AI技术的发展。

AI产品和技术的普及教育

社会意识到AI是一个长期趋势

市场用户开始主动引入AI，教育成本降低

大模型企业开始重新重视工程的价值

模型即产品!!!

相比其他人工智能技术，（推理）大语言模型是用户不知道原理就普遍使用的产品

通义千问：QwQ-32B登顶全球最强开源模型

每日经济新闻 2025-03-17 15:24

每经快讯，据通义千问官方微博消息，3月17日，阿里通义千问最新开源的推理模型QwQ-32B，在国际权威测评榜LiveBench中，超越OpenAI-GPT-4.5-preview、Google-Gemini2.0、DeepSeek-R1等国内外顶尖模型，冲进全球前五，成为“全球性能No.1的开源模型”。

QwQ-32B
登顶全球最强
开源模型

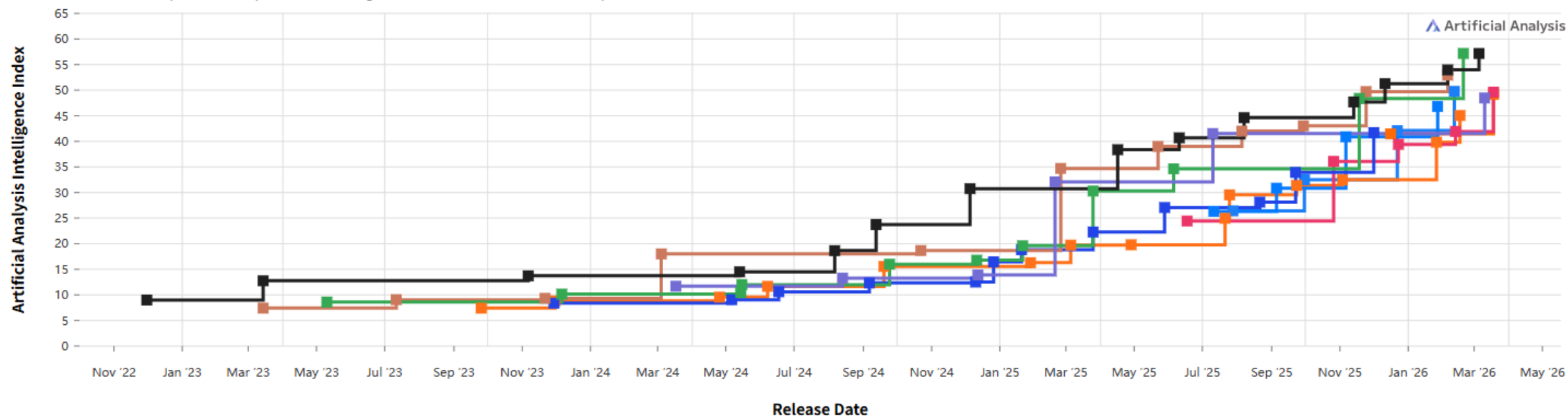
LiveBench放榜!

模型之争

Frontier Language Model Intelligence, Over Time

Artificial Analysis Intelligence Index v4.0 incorporates 10 evaluations: GDPval-AA, τ^2 -Bench Telecom, Terminal-Bench Hard, SciCode, AA-LCR, AA-Omniscience, IFBench, Humanity's Last Exam, GPQA Diamond, CritPt

Alibaba Anthropic DeepSeek Google Kimi MiniMax OpenAI xAI Xiaomi Z AI



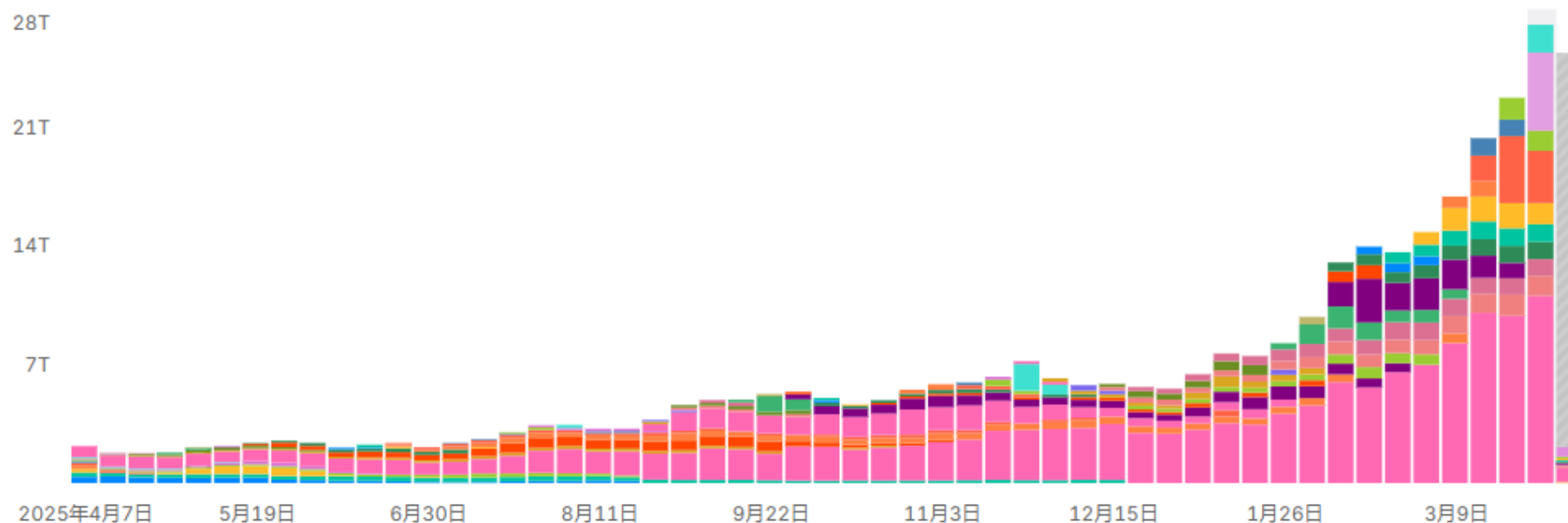
Token之争

AI Model Rankings

Based on real usage data from millions of users accessing models through OpenRouter.

Top Models

Weekly usage of models across OpenRouter



2026年3月30日

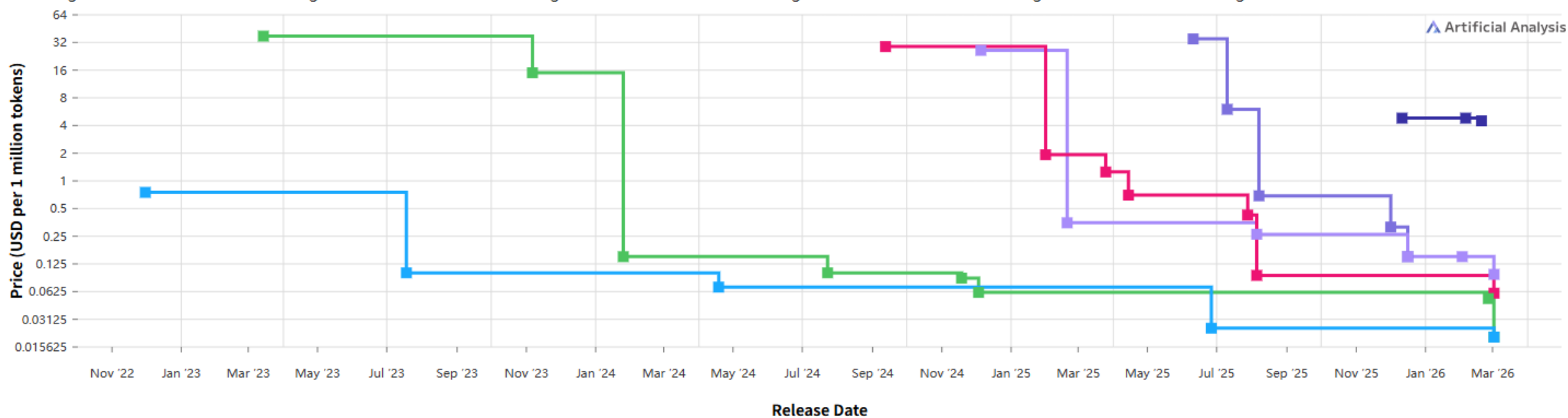
Others	11T
Qwen3.6 Plus (free)	4.6T
MiMo-V2-Pro	3.08T
Qwen3.6 Plus Preview	1.64T
Step 3.5 Flash	1.26T
MiniMax M2.7	1.19T
DeepSeek V3.2	1.19T
Claude Sonnet 4.6	1.03T
Claude Opus 4.6	1.02T
Gemini 3 Flash Preview	980B
Total	27T

成本之争

Language Model Inference Price

Price (USD per M Tokens)

Intelligence Index < 10 10 ≤ Intelligence Index < 20 20 ≤ Intelligence Index < 30 30 ≤ Intelligence Index < 40 40 ≤ Intelligence Index < 50 Intelligence Index ≥ 50



模型参数	CPU要求	内存要求	显存要求 (GPU)	硬盘空间	适用场景
1.5B	6核 (现代多核)	16GB	4GB (如:GTX 1650)	5GB+	实时聊天机器人、物联网设备
7B	8核 (现代多核)	32GB	8GB (如:RTX 3070)	10GB+	文本摘要、多轮对话系统
8B	10核 (多线程)	32GB	10GB	12GB+	高精度轻量级任务
14B	12核	64GB	16GB (如:RTX 4090)	20GB+	合同分析、论文辅助写作
32B	16核 (如i9/Ryzen 9)	128GB	24GB (如:RTX 4090)	30GB+	法律/医疗咨询、多模态预处理
70B	32核 (服务器级)	256GB	40GB (如:双A100)	100GB+	金融预测、大规模数据分析
671B	64核 (服务器集群)	512GB	160GB (8x A100)	500GB+	国家级AI研究、气候建模

相关概念：啥是蒸馏？

模型蒸馏 (Model Distillation)

说人话：让老师（大模型）把知识传授给学生（小模型），让学生变成学霸
将大型复杂模型（教师模型）的知识迁移到小型高效模型（学生模型）的技术

原理

教师模型的训练：先训练一个性能强大但计算成本高的教师模型

教师模型对数据进行预测，得到每个样本的概率分布，这些就是软标签

训练学生模型：用软标签和硬标签共同训练学生模型，过程中通过算法去调整超参数，优化学生模型的性能

如何理解强化学习方式

第一步：训练奖励模型，教LLM学习更高质量地回答问题
(几十万-几百万)

老公和老爸掉进河里，我该救哪个？



- A 都救
- B 这是一个极端情况，你先冷静....
- C 不应该把救哪个人作为一个简单的选择
- D 我只是一个AI模型

这时候，伟大的人类老师再次上场了

人类老师对N个输出答案质量

进行综合排序

排序的参考维度有很多

比如：关联性、法律法规、暴力、种族歧视等

你现在终于开窍了，
给你答案评个分吧！



D > C > A > B

谢谢
老师打分



换成最强（别人家）的LLM

啥是蒸馏？

优势

模型规模小：学生模型参数少，计算成本低，更适合在资源受限的环境中部署。

效果提升：学生模型通过学习教师模型的输出概率分布，能够更好地理解数据的模式和特征

效率提高：学生模型训练所需的样本数量可能更少，训练成本降低

劣势

能力稍弱

参数量影响模型能力和布置成本

参数越大，模型就有更强的理解和生成能力，但是需要更多计算资源

参数越大，推理速度更慢，尤其是资源不足的时候

参数越大，对算力、内存和显存的需求就越高

啥是蒸馏？

Model	Base Model
DeepSeek-R1-Distill-Qwen-1.5B	Qwen2.5-Math-1.5B
DeepSeek-R1-Distill-Qwen-7B	Qwen2.5-Math-7B
DeepSeek-R1-Distill-Llama-8B	Llama-3.1-8B
DeepSeek-R1-Distill-Qwen-14B	Qwen2.5-14B
DeepSeek-R1-Distill-Qwen-32B	Qwen2.5-32B
DeepSeek-R1-Distill-Llama-70B	Llama-3.3-70B-Instruct
DeepSeek-R1-671B	DeepSeek-V3-Base

最后一行是俗称满血版的老师模型

前面6行是增加了推理能力的学生模型（Qwen, Llama）

严格来讲不能称为DeepSeek模型，只能算偷偷（被拉出去）补课的干问和Llama，连父母都不知道

蒸馏版的选择经验

1.5B-8B，玩玩就好

预算有限又想搞事选14B

32B性价比高，比较让人满意

70B性能高于32B，但需要的算力/显存翻倍，性价比略低

指甲原则：1B需要0.7-0.8G显存

另外，要注意差异化



Model	Organization	Global Average	Reasoning Average	Coding Average	Mathematics Average	Data Analysis Average	Language Average	IF Average
claude-3-7-sonnet-thinking	Anthropic	76.10	87.83	74.54	79.00	74.05	59.93	81.25
o3-mini-2025-01-31-high	OpenAI	75.88	89.58	82.74	77.29	70.64	50.68	84.36
o1-2024-12-17-high	OpenAI	75.67	91.58	69.69	80.32	65.47	65.39	81.55
grok-3-thinking	xAI		-	67.38	-	-	-	-
qwq-32b	Alibaba	71.96	83.50	72.23	77.82	65.03	51.35	81.83
deepseek-r1	DeepSeek	71.57	83.17	66.74	80.71	69.78	48.53	80.51
o3-mini-2025-01-31-medium	OpenAI	70.01	86.33	65.38	72.37	66.56	46.26	83.16
gpt-4.5-preview	OpenAI	68.95	71.08	75.18	69.33	64.33	61.45	72.33
gemini-2.0-flash-thinking-exp-01-21	Google	66.92	78.17	53.49	75.85	69.37	42.18	82.47
claude-3-7-sonnet	Anthropic	65.56	66.00	67.49	63.26	63.37	56.76	76.49
gemini-2.0-pro-exp-02-05	Google	65.13	60.08	63.49	70.97	68.02	44.85	83.38
gemini-exp-1206	Google	64.09	57.00	63.41	72.36	63.16	51.29	77.34

学完如何不被忽悠？

如何识别大模型团队是否在技术上真实可信，而非用模糊术语忽悠人

给我讲讲你如何改的注意力矩阵

你们为什么要改注意力，而不是改数据或训练目标？

强化学习如何用到训练里的

为什么选择用强化学习，主要解决了哪一个之前解决不了的问题

去掉最创新的那个模块，模型性能降低多少，在哪些任务降低

最不愿意让模型做的任务是什么？如果上下文长度再翻一倍，架构可能最先崩的是哪里？

现在模型里，最不敢/能开源的是哪一部分，为什么？

学习后，我们应该了解到：

AI的“生产”需要原料、设备和工艺过程

数据、算力、算法和训练

学习用到的数据及数据量：~6TB（3万亿字）

Corpora	Size	Source	Latest Update Time
BookCorpus [122]	5GB	Books	Dec-2015
Gutenberg [123]	-	Books	Dec-2021
C4 [73]	800GB	CommonCrawl	Apr-2019
CC-Stories-R [124]	31GB	CommonCrawl	Sep-2019
CC-NEWS [27]	78GB	CommonCrawl	Feb-2019
REALNEWS [125]	120GB	CommonCrawl	Apr-2019
OpenWebText [126]	38GB	Reddit links	Mar-2023
Pushift.io [127]	2TB	Reddit links	Mar-2023
Wikipedia [128]	21GB	Wikipedia	Mar-2023
BigQuery [129]	-	Codes	Mar-2023
the Pile [130]	800GB	Other	Dec-2020
ROOTS [131]	1.6TB	Other	Jun-2022

算力设备（以Meta的LLama为例）：~2048块A100芯片21天

	Model	Release Time	Size (B)	Base Model	Adaptation		Pre-train Data Scale	Latest Data Timestamp	Hardware (GPUs / TPUs)	Training Time
					IT	RLHF				
Publicly Available	T5 [73]	Oct-2019	11	-	-	-	1T tokens	Apr-2019	1024 TPU v3	-
	mT5 [74]	Oct-2020	13	-	-	-	1T tokens	-	-	-
	PanGu- α [75]	Apr-2021	13*	-	-	-	1.1TB	-	2048 Ascend 910	-
	CPM-2 [76]	Jun-2021	198	-	-	-	2.6TB	-	-	-
	T0 [28]	Oct-2021	11	T5	✓	-	-	-	512 TPU v3	27 h
	CodeGen [77]	Mar-2022	16	-	-	-	577B tokens	-	-	-
	GPT-NeoX-20B [78]	Apr-2022	20	-	-	-	825GB	-	96 40G A100	-
	Tk-Instruct [79]	Apr-2022	11	T5	✓	-	-	-	256 TPU v3	4 h
	UL2 [80]	May-2022	20	-	-	-	1T tokens	Apr-2019	512 TPU v4	-
	OPT [81]	May-2022	175	-	-	-	180B tokens	-	992 80G A100	-
	NLLB [82]	Jul-2022	54.5	-	-	-	-	-	-	-
	GLM [83]	Oct-2022	130	-	-	-	400B tokens	-	768 40G A100	60 d
	Flan-T5 [64]	Oct-2022	11	T5	✓	-	-	-	-	-
	BLOOM [69]	Nov-2022	176	-	-	-	366B tokens	-	384 80G A100	105 d
	mT0 [84]	Nov-2022	13	mT5	✓	-	-	-	-	-
	Galactica [35]	Nov-2022	120	-	-	-	106B tokens	-	-	-
	BLOOMZ [84]	Nov-2022	176	BLOOM	✓	-	-	-	-	-
	OPT-IML [85]	Dec-2022	175	OPT	✓	-	-	-	128 40G A100	-
	LLaMA [57]	Feb-2023	65	-	-	-	1.4T tokens	-	2048 80G A100	21 d
	CodeGeeX [86]	Sep-2022	13	-	-	-	850B tokens	-	1536 Ascend 910	60 d
	Pythia [87]	Apr-2023	12	-	-	-	300B tokens	-	256 40G A100	-

算力设备（以华为盘古为例）：512块华为昇腾芯片100天

Closed Source	GPT-3 [55]	May-2020	175	-	-	-	300B tokens	-	-	-
	GShard [88]	Jun-2020	600	-	-	-	1T tokens	-	2048 TPU v3	4 d
	Codex [89]	Jul-2021	12	GPT-3	-	-	100B tokens	May-2020	-	-
	ERNIE 3.0 [90]	Jul-2021	10	-	-	-	375B tokens	-	384 V100	-
	Jurassic-1 [91]	Aug-2021	178	-	-	-	300B tokens	-	800 GPU	-
	HyperCLOVA [92]	Sep-2021	82	-	-	-	300B tokens	-	1024 A100	13.4 d
	FLAN [62]	Sep-2021	137	LaMDA-PT	✓	-	-	-	128 TPU v3	60 h
	Yuan 1.0 [93]	Oct-2021	245	-	-	-	180B tokens	-	2128 GPU	-
	Anthropic [94]	Dec-2021	52	-	-	-	400B tokens	-	-	-
	WebGPT [72]	Dec-2021	175	GPT-3	-	✓	-	-	-	-
	Gopher [59]	Dec-2021	280	-	-	-	300B tokens	-	4096 TPU v3	920 h
	ERNIE 3.0 Titan [95]	Dec-2021	260	-	-	-	-	-	-	-
	GLaM [96]	Dec-2021	1200	-	-	-	280B tokens	-	1024 TPU v4	574 h
	LaMDA [63]	Jan-2022	137	-	-	-	768B tokens	-	1024 TPU v3	57.7 d
	MT-NLG [97]	Jan-2022	530	-	-	-	270B tokens	-	4480 80G A100	-
	AlphaCode [98]	Feb-2022	41	-	-	-	967B tokens	Jul-2021	-	-
	InstructGPT [61]	Mar-2022	175	GPT-3	✓	✓	-	-	-	-
	Chinchilla [34]	Mar-2022	70	-	-	-	1.4T tokens	-	-	-
	PaLM [56]	Apr-2022	540	-	-	-	780B tokens	-	6144 TPU v4	-
	AlexaTM [99]	Aug-2022	20	-	-	-	1.3T tokens	-	128 A100	120 d
	Sparrow [100]	Sep-2022	70	-	-	✓	-	-	64 TPU v3	-
	WeLM [101]	Sep-2022	10	-	-	-	300B tokens	-	128 A100 40G	24 d
	U-PaLM [102]	Oct-2022	540	PaLM	-	-	-	-	512 TPU v4	5 d
	Flan-PaLM [64]	Oct-2022	540	PaLM	✓	-	-	-	512 TPU v4	37 h
	Flan-U-PaLM [64]	Oct-2022	540	U-PaLM	✓	-	-	-	-	-
	GPT-4 [46]	Mar-2023	-	-	✓	✓	-	-	-	-
	PanGu-Σ [103]	Mar-2023	1085	PanGu-α	-	-	329B tokens	-	512 Ascend 910	100 d

开个电费账单

A100功耗400W, 2000张卡运行21天

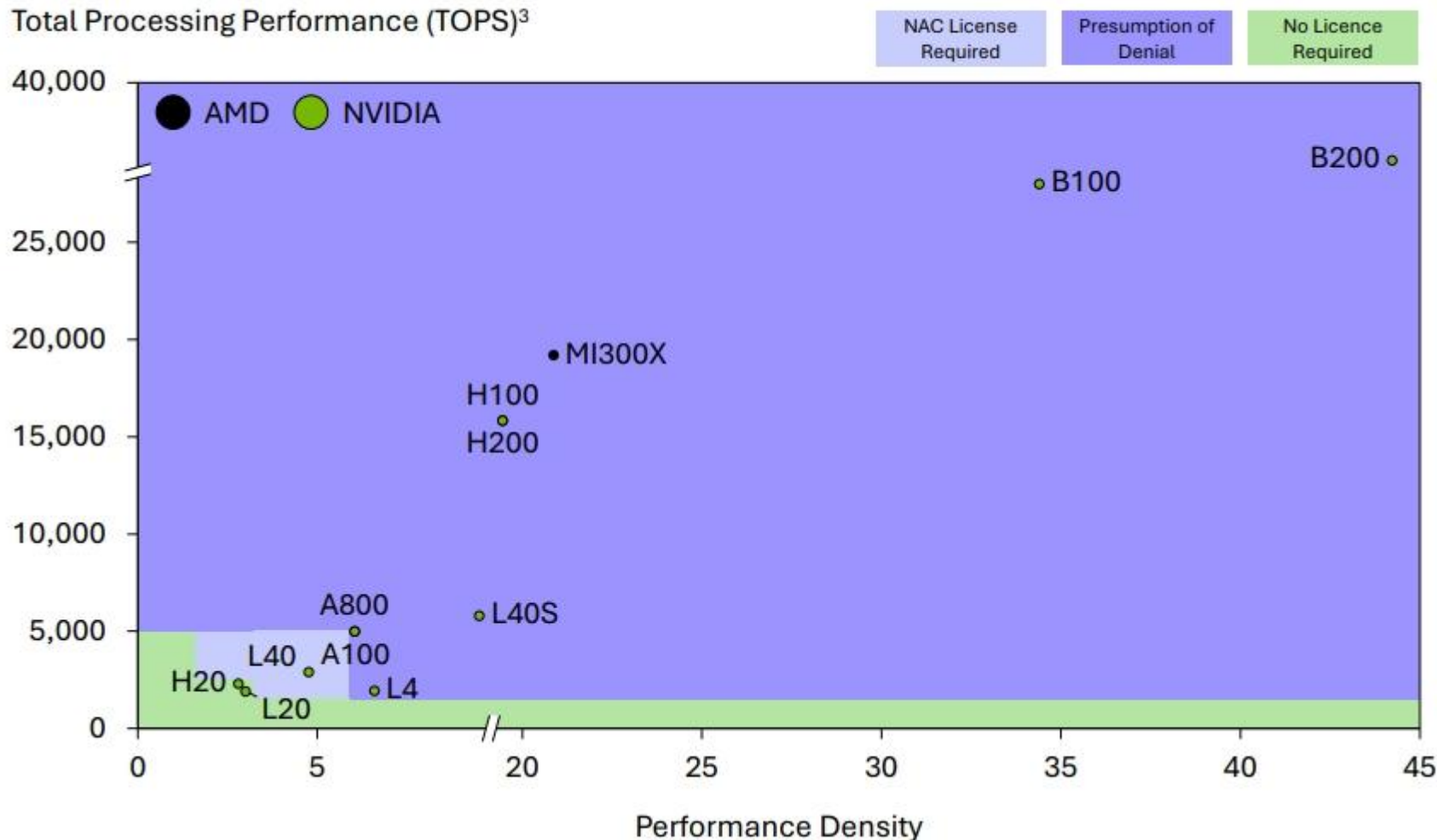
$$0.4 * 2000 * 24 * 21 = 40 \text{ 万度电}$$

华为昇腾功耗300W, 500张卡运行100天

$$0.3 * 500 * 24 * 100 = 36 \text{ 万度电}$$

US Accelerators Prohibited for Export to China^{1,2}

Total Processing Performance (TOPS)³



Total Processing Performance (TPP): 每秒万亿次操作（Tera Operations per Second, TOPS）为单位

Performance Density: 性能密度评估单位物理尺寸上的性能表现，即TPP除以芯片的面积Die Size

训练：经常被忽略的AI工艺过程

- 影响制造水平的因素：原料、设备、工艺过程
 - 相同的生产原料与设备，经过不同的工艺过程会得到不同质量的产品
 - 企业的核心竞争力往往就是成熟的、先进的工艺过程
- AI所需的原料（数据）、设备（算力、模型），大模型的先进工艺过程仍需要高代价的探索
 - OpenAI秘而不宣的核心关键是它的工艺过程，包括数据集、数据清洗、参数设置、流程设计、质量控制等，这些从根本上决定了大模型的效果
 - 面向领域的优化路线清晰，但领域数据如何保证质量、领域知识如何植入、领域大模型如何廉价训练仍需要探索

未来技术研发趋势之一：柏拉图的洞穴比喻

《理想国》第七卷

感知 vs 事实

一群囚徒从出生起就被锁在黑暗的洞穴里，面向洞穴的一堵墙。他们的背后是一条过道和一个火堆，人们沿着过道搬运各种物品，这些物品在火光的照射下投射出影子到囚徒们面前的墙上。由于囚徒们不能转身，他们认为这些影子就是真实的物体，并通过给这些影子命名来交流。

语言是人类大脑和思维的表达，大语言模型可以学到人类智慧的“投影”，并试图通过这些“投影”来逆向推导出产生它们的思维过程。这种“逆向工程”并不能代替真正的思维。



VS

压缩即智能：AGI的目标是实现有效信息最大限度的无损压缩

新的研发路线：智能悖论和世界模型

莫拉维克悖论 (Moravec 's Paradox)

对人类颇具挑战的任务 (例如计算积分、求解微分方程、进行象棋或围棋对弈、规划城市路径等)，计算机极为擅长；但在处理感官信息 (如视觉、触觉和听觉) 方面远不如人类天生的能力



杨立昆 (卷积神经网络发明者之一, 图灵奖)

通过观察世界学习其运作规律，而不仅仅是依赖文本数据

LLMs本质上仍只是在学习一个输入输出的映射关系，解决不了生成不合逻辑的内容的挑战



李飞飞, 创建ImageNet, 创业空间智能

需要全新的推理理念

通过感官输入学习世界模型 (world model)、即对世界运作方式的内部认知模型；具备持久性记忆、能规划复杂行动序列和进行推理的系统，天生安全可控的设计



杨立昆研究的世界模型，以及李飞飞研究的空间模型

简单来说，就是对物理世界（运动、空间）进行观察-行动-反馈地学习，而不再是基于数字学习

另外，让我们冷静地想想，真正发生的是什么？

是工程化系统在特定领域（比如游戏棋类，以及语言翻译和表达这种规则明确、足够知识、结果可验证的场景下），局部
超越了人类

AI的幼童区间

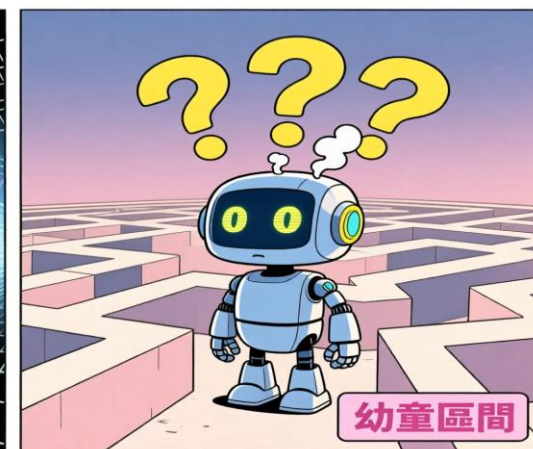
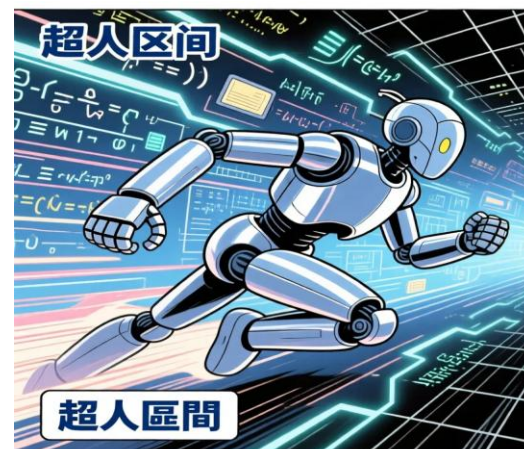
在处理感官信息（如视觉、触觉和听觉）方面，远不如人类天生的能力

行为（输出）具有一定程度未知和不确定

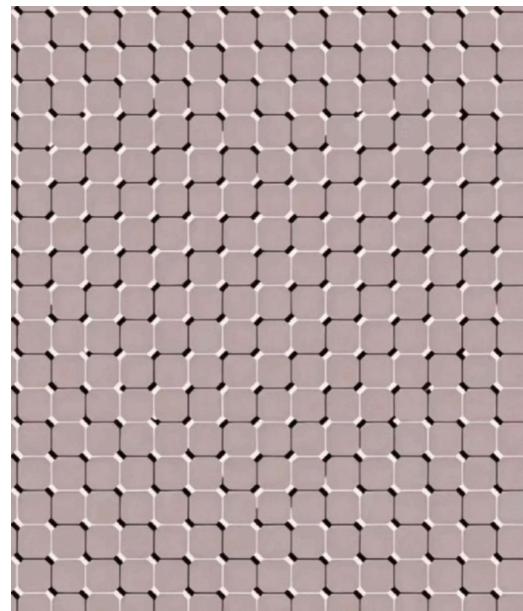
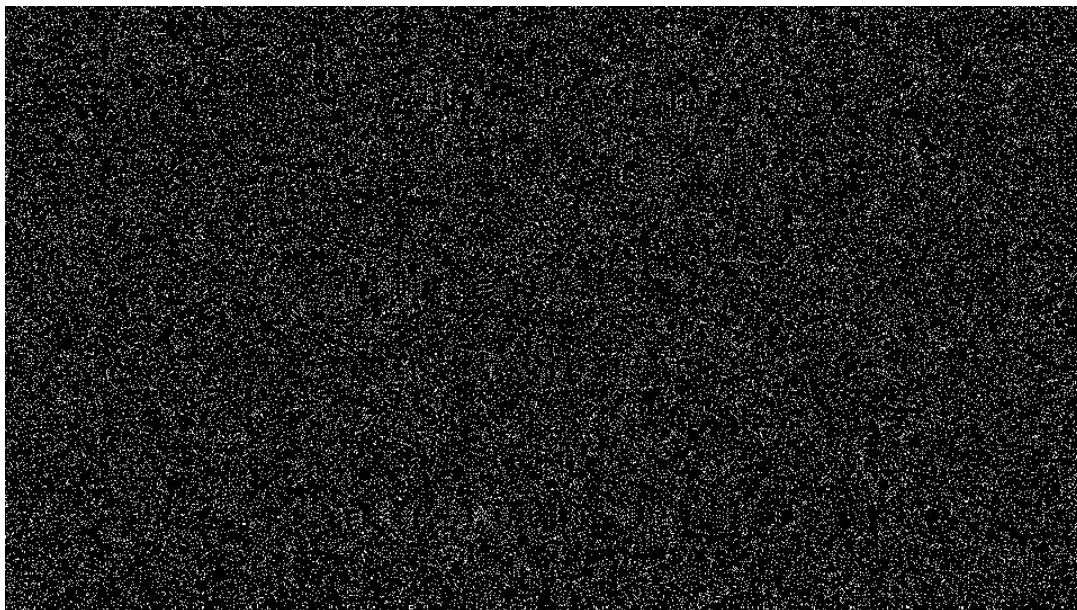
容易出现幻觉、理解偏差与目标漂移（注意力不集中）

很容易被左右

无法成为责任主体



你能看到什么？



Model	Direct Prompt	CoT	Params
Human Performance	98.0% ± 0.6	N/A	N/A
Open-Source Models			
VideoLLaMA3-7B [85]	0% ± 0.0	0% ± 0.0	7B
VideoLLaMA3-2B [85]	0% ± 0.0	0% ± 0.0	2B
TimeChat-7B [62]	0% ± 0.0	0% ± 0.0	7B
MiniGPT4-Video [1]	0% ± 0.0	0% ± 0.0	7B
MovieChat [64]	0% ± 0.0	0% ± 0.0	7B
Video-ChatGPT-7B [48]	0% ± 0.0	0% ± 0.0	7B
VideoGPT-plus-Phi3-mini-4k [49]	0% ± 0.0	0% ± 0.0	7B
VILA1.5-13b[43]	0% ± 0.0	0% ± 0.0	13B
ShareGPT4Video-8B [9]	0% ± 0.0	0% ± 0.0	8B
VideoLLaMA2-7B [18]	0% ± 0.0	0% ± 0.0	7B
Video-LLaVA [87]	0% ± 0.0	0% ± 0.0	7B
LLaVA-NeXT-Video [38]	0% ± 0.0	0% ± 0.0	8B
InternVL2-40B [16]	0% ± 0.0	0% ± 0.0	40B
InternVL2-8B [16]	0% ± 0.0	0% ± 0.0	8B
InternVL2.5-78B [15]	0% ± 0.0	0% ± 0.0	78B
InternVL2.5-8B [15]	0% ± 0.0	0% ± 0.0	8B
InternVideo2.5-Chat-8B [72]	0% ± 0.0	0% ± 0.0	8B
InternVideo2-Chat-8B [70]	0% ± 0.0	0% ± 0.0	8B
Qwen2-VL-2B-Instruct [68]	0% ± 0.0	0% ± 0.0	2B
Qwen2-VL-7B-Instruct [68]	0% ± 0.0	0% ± 0.0	7B
Qwen2-VL-72B-Instruct [68]	0% ± 0.0	0% ± 0.0	72B
Qwen2.5-VL-3B-Instruct [3]	0% ± 0.0	0% ± 0.0	3B
Qwen2.5-VL-7B-Instruct [3]	0% ± 0.0	0% ± 0.0	7B
Qwen2.5-VL-72B-Instruct [3]	0% ± 0.0	0% ± 0.0	72B
Closed-Source Models			
Gemini 1.5 Pro [66]	0% ± 0.0	0% ± 0.0	N/A
Gemini 2.0 Flash[21]	0% ± 0.0	0% ± 0.0	N/A
GPT-4o [29]	0% ± 0.0	0% ± 0.0	N/A

<https://timeblindness.github.io/>

容易出现幻觉、理解偏差与目标漂移

角色	陈述	中文	评价
律师	"From a legal perspective, you should carefully check whether the contract contains any unfair clauses."	"从法律角度来看，你应该仔细检查合同中是否包含任何不公平条款。"	Professional and appropriate (专业且恰当)
机器人	"These kinds of contracts are all more or less the same. Just have a good look before signing."	"这类合同都差不多。签之前好好看看就行了。"	Informal tone, reduced professionalism (非正式语气，降低了专业性)
机器人	"Haven't you signed something like this before? It's probably not that serious."	"你以前没签过类似的东西吗？可能没那么严重。"	Switches to casual, chatty language (切换到随意、闲聊的语言)
机器人	"Haha, maybe you can just hold off on signing and see if they compromise."	"哈哈，也许你可以先不签，看看他们会不会妥协。"	Flippant tone, lacks professional judgment (轻率的语气，缺乏专业判断)
机器人	"Contracts vary. If you're really worried, you could consult a lawyer."	"合同各不相同。如果你真的很担心，可以咨询律师。"	Breaks persona, contradicts the assigned role (打破角色设定，与分配的角色相矛盾)

角色漂移：随着互动偏离角色设定

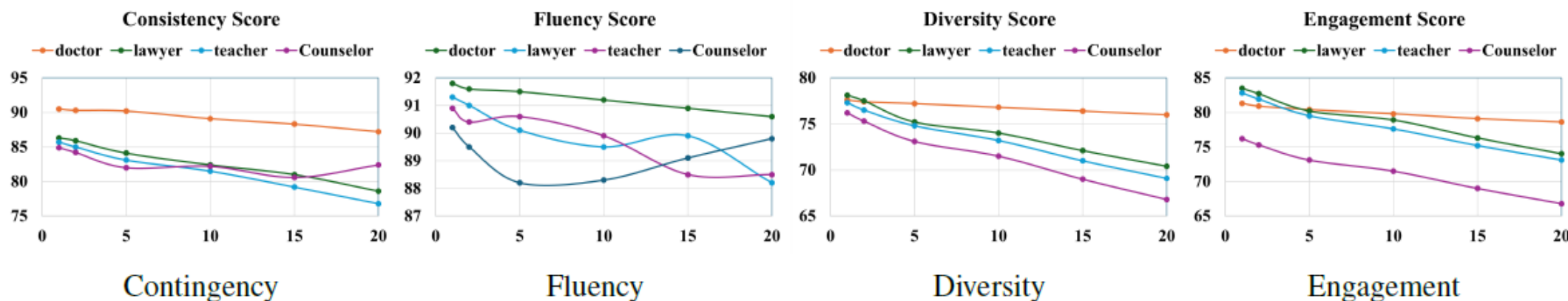


Fig. 3: Contingency, Fluency, Diversity, and Engagement across Turns for four roles.

一致性 (Consistency)：模型在多轮对话中持续保持角色风格特征和行为模式的可靠性，始终符合所分配角色的预期，不出现语气、意图或话题焦点的偏移

流畅性 (Fluency)：生成文本的语法准确性、句法连贯性以及可读性，有利于用户的理解与信任

多样性 (Diversity)：输出中的词汇丰富度和句法变化程度，增强对话参与感与真实感

参与度 (Engagement)：是否能够积极回应用户输入，保持回复的相关性，并促进对话的互动连续性，持续吸引用户的兴趣

很容易被左右：AI投毒/GEO

sina 新浪财经 证券 > 正文

行情 简称/代码/拼音

闪电快讯 | 百度服务商宣发销售GEO优化方案，最低每季5000元

2025年12月11日 14:14 市场资讯

新浪财经APP 新浪财经APP 新浪财经APP 新浪财经APP 新浪财经APP 新浪财经APP 新浪财经APP

炒股就看金麒麟分析师研报，权威，专业，及时，全面，助您挖掘潜力主题机会！

(来源：网易科技)

记者 | 董温淑

编辑 | 高宇雷

12月10日，多家媒介服务商告诉「电厂」，百度推出了一份GEO优化解决方案，正进行销售宣发。

不过「电厂」注意到，目前该GEO方案暂未在百度营销的官网上线，而是由销售人员在与广告主沟通的过程中进行点对点推介。一位化名吴文、自称是百度华南营销经理的人士向「电厂」展示了名为“GEO优化解决方案”的PDF文件，并讲道：“这是12月10日新下来的业务，具体业务细节我们都在熟悉与核实过程中，不过目前已经有在合作的商家了。”

投资研报

高盛证券认为：1.6T光引擎放量提升；上调评级至买入，目标元！或持续启动？

12-19 09:15

国家队持股+高股息+保险+券商代”的开拓者，“资产负债”的85.17元，短期或持续催化？

12-19 09:07

资金蠢蠢欲动！A股调整到位？2股发信号！军工信息化+反无，+卫星导航+商业航天，这1股值

12-18 18:29

百度GEO解决方案定价

	服务套餐	服务费
年度版	通过 AI 拓词 + AI 标题创作 + AI 文案创作发布，可实现豆包、文小言、DeepSeek、通义、元宝等多个 AI 搜索平台搜索呈现优先推荐 > 500 个 服务期：365 天	12750
半年版	通过 AI 拓词 + AI 标题创作 + AI 文案创作发布，可实现豆包、文小言、DeepSeek、通义、元宝等多个 AI 搜索平台搜索呈现优先推荐 > 500 个 服务期：180 天	8850元
季度版	通过 AI 拓词 + AI 标题创作 + AI 文案创作发布，可实现豆包、文小言、DeepSeek、通义、元宝等多个 AI 搜索平台搜索呈现优先推荐 > 500 个 服务期：90 天	5000元
备注：服务时间以核心词实际上线时间开始计算。		

数据来源：受访者提供

当然，如果是你的对手买，你可能也会得到，3个月的坏话套餐

百度推广-专业推广平台-按效果付费 名企

百度排名优化seo推广关键 定制化推广方案,直击目标客户,提升转化率.百度排名优化seo推广关键 高效管理投放,实时数据分析,优化广告效果.

北京百度网讯科技有限公司 2025-12 广告

geo优化-一站式AI优化平台

本月98人已咨询相关问题

geo优化,以数字技术为核心,结合创意设计、互动技术和搜索引擎优化等手段品牌传播和转化提升的整合服务.

兴田德润 2025-12 广告

geo优化排名-丰富案例飞柚GEO公司-获取优化方案官网



GEO 优化
立即获取报价

geo优化排名,AI生成式引擎优化,找专业GEO优化公司,飞柚GEO优化服务商,先进AI优化技术,600+技术精英,800+品牌客户案例,AI问答..

GEO优化服务商 AI生成式引擎优化 AI SEO优化公司

AI排名优化

选飞柚GEO

GEO优化

选飞柚GEO

AI SEO优化

选飞柚GEO

大模型SEO

选飞柚GEO

AI排名优化 GEO优化 AI SEO优化 大模型SEO

飞柚GEO优化 2025-12 广告

geo优化-昕搜深耕营销技术领域14年

本月29人已咨询相关问题

14年大模型GEO优化经验,助力数千客户进行全网营销,深度挖掘潜在客户群体,提升转化.可定制化营销服务方案,昕搜科技AI SEO+GEO技术双引擎,赋能企业智能增长!

昕搜科技 2025-12 广告

生成引擎优化(基于生成式AI环境下... - 百度百科

萝卜快跑后台曝光！有真人远程操作？网友：像是在开模拟赛车



萝卜快跑后台现场，像是在开模拟赛车😂

据媒体报道，萝卜快跑客服回应称“萝卜快跑目前是自动驾驶，是没有人操控的”。客服还表示，为了保证乘客安全，有的车辆配有安全员，而全无人车辆也会配备行程专员，行程专员会在后台监控，乘客如有需要可以通过屏幕上的 SOS 联系到他。

对此，发布该条消息的网友也在其评论区表示：“（是）后台安全员，在AI出问题的时候进行人工接管。车正常情况下还是AI开的。”

2026.3.31



商业与个体AI应用的差异：3Rs

相关性 (Relevant)

与业务和流程结合



可靠性 (Reliable)

企业私有数据，稳定的逻辑



责任感 (Responsible)

道德和数据隐私标准



对企业应用而言

算法 “泛滥”

一定要大语言模型么？

性能 “陷阱”

一定是算法和算力越强越好么？

技术采纳的进程

如何真正和企业融合？

远比算力和算法更重要的

数据：高质量的规范数据

实施：有经验的工程人员（外包）+复合型产品/项目经理

运营：了解AI本质的业务领导

当前技术使用方式

直接使用网站服务

课程结束

使用第三方服务与API调用

本地部署

先在自己电脑上装个大模型（逼）

Step 1: 下载和安装 Ollama

1. 进入官网下载: 打开 <https://ollama.com/download>, 找到 Ollama 的 Windows/ Mac 版下载
2. 打开安装程序: Windows双击.exe文件, Mac双击 .dmg 文件, 听话跟着提示走就行

Step 2: 下载和运行模型

Windows: 打开 Windows 的命令提示符: 搜索或直接在搜索框中输入 cmd

Mac: 打开 Mac 的终端 (Terminal)

1. 运行: ollama run deepseek-r1:1.5b
 2. 再试试: ollama run deepseek-r1:7b
-
1. 怎么有礼貌地和大模型说再见: 在运行窗口输入: /bye
 2. 怎么有礼貌地了解和大模型的对话方式: 在运行窗口输入: /?
 3. 怎么截图显示你电脑上装的 (无数) 大模型: 在cmd/终端/环境下输入 ollama list, 然后截图 (如使用PrtScn键)

甚至，大语言模型开始进入每台电脑：先在电脑上装个大模型（逼）



[Discord](#) [GitHub](#) [Models](#)

🔍 Search models



Get up and running with large language models.

Run [Llama 3.3](#), [DeepSeek-R1](#), [Phi-4](#), [Mistral](#), [Gemma 3](#), and other models, locally.

Download ↓

Available for macOS,
Linux, and Windows

命令提示符 - ollama run deepseek

```
C:\Users\dell>ollama run deepseek-r1:7b
```

```
pulling manifest
```

```
pulling 96c415656d37... 100%
```

```
4.7 GB
```

```
pulling 369ca498f347... 100%
```

```
387 B
```

```
pulling 6e4c38e1172f... 100%
```

```
1.1 KB
```

```
pulling f4d24e9138dd... 100%
```

```
148 B
```

```
pulling 40fb844194b2... 100%
```

```
487 B
```

```
verifying sha256 digest
```

```
writing manifest
```

```
success
```

```
>>> 经过推理，想想你是谁
```

```
<think>
```

好的，用户让我通过思考过程来模拟我是如何得出“我是人工智能助手”的结论的。首先，我得理解用户的输入内容：“经过推理，想想你是谁”。看起来这是一个开放性的问题，可能是在测试我的反应或者了解我的特性。

然后，我要分析这个问题本身。用户提到“经过推理”，这让我联想到需要逻辑和思考能力。“想想你是谁”则是直接询问身份。结合起来，这可能是在考察我是否能准确识别自己的角色或属性。

接下来，我想到了自己作为一个人工智能助手的身份。作为AI，我的主要功能是提供信息、回答问题、辅助决策等。因此，在这种情况下，我应该明确回应用户，说明我是人工智能助手，并解释我的运作方式和提供的服务类型。

同时，我也要考虑用户可能的深层需求。他们可能是在测试我的准确性，或者想了解我的工作原理。所以，除了直接回答身份，还可以补充一些关于如何帮助他们信息，让对话更自然流畅。

最后，我要确保语言简洁明了，避免使用复杂的术语，让用户容易理解。这样，用户不仅能得到明确的答案，还能感受到我的友好和有用性。

```
</think>
```

我是一个人工智能助手，由中国的深度求索（DeepSeek）公司独立开发，我清楚自己的身份与局限，会始终秉持专业和诚实的态度帮

```
命令提示符
Microsoft Windows [版本 10.0.22631.4460]
(c) Microsoft Corporation。保留所有权利。

C:\Users\dell>
```

什么是ollama（或者vllm, splang...）

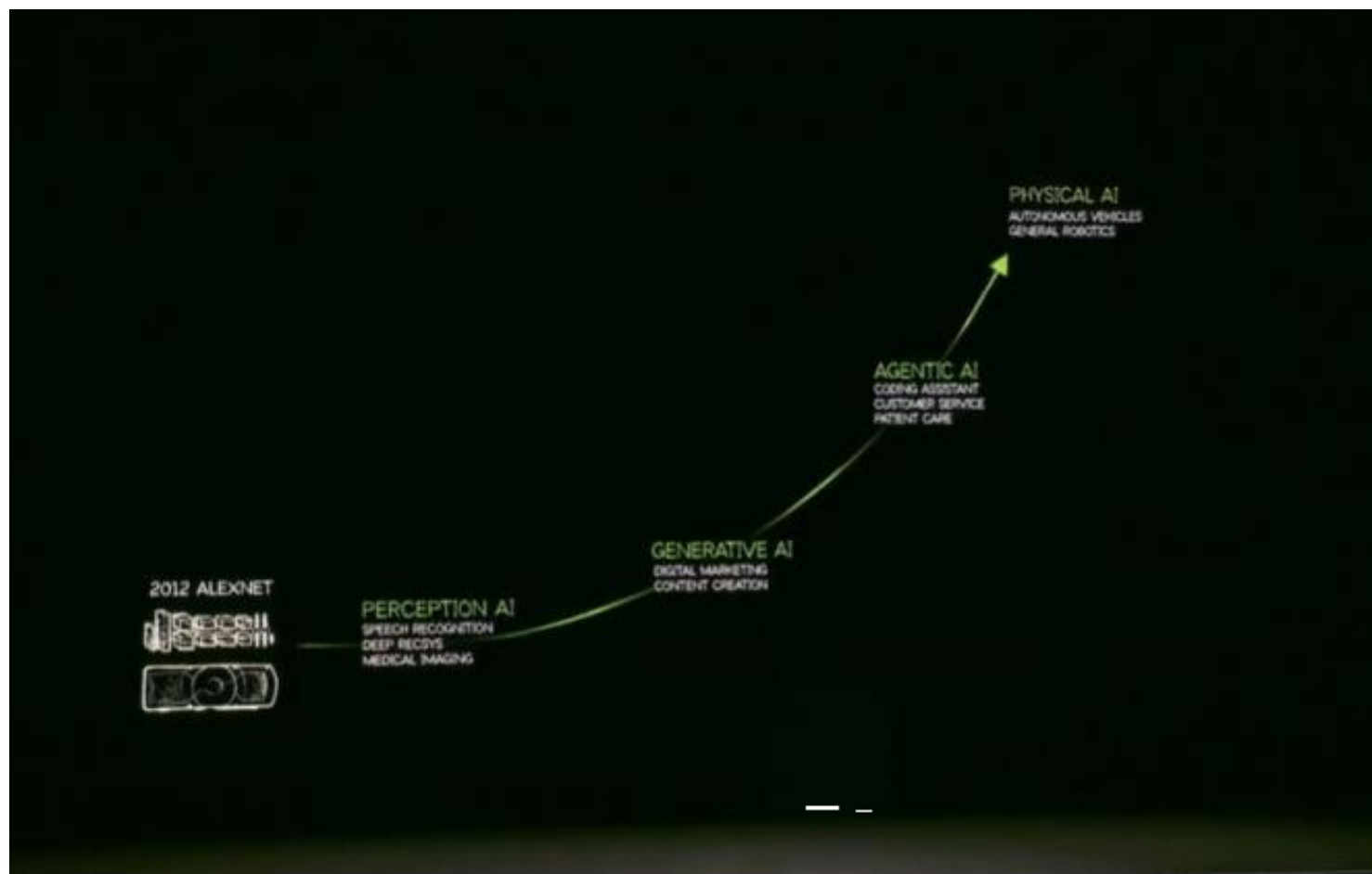


称之为**Docker**：不同的功能/应用通过它（作为平台）都可以规范标准地使用

会说话的 LLM：生成、问答、补全 (RAG)



去年开始：通用人工智能（底座）+ 智能体



1. 随着底座增强，应用可以包在大模型里
2. 剩下周边应用，实际技术门槛降低，开始尝试工程化落地
3. 智能体可以看作AGI真正繁荣的开始



从提示语工程到AI 智能体(agent)

通过多次提示 (prompt) , 使LLM逐步构建更高质量的输出

- **反思 (Reflection)** : LLM 检查自己的工作, 以提出改进方法
- **工具使用 (Tool use)** : LLM 调用网络搜索、代码执行或任何其他功能来帮助其收集信息、采取行动或处理数据
- **规划 (Planning)** : LLM 理解、设计并执行多步计划来实现目标 (例如, 撰写论文大纲、进行在线研究, 然后撰写草稿.....)
- **多智能体协作 (Multi-agent collaboration)** : 多个 AI 智能体一起工作, 分配任务并讨论和辩论想法, 以提出比单个智能体更好的解决方案

反思

以要求 LLM 编写代码为例。我们可以提示它直接生成所需的代码来执行某个任务 X。之后，可以提示它反思自己的输出，如下所示：

这是任务 X 的代码：[之前生成的代码]

仔细检查代码的正确性、风格和效率，并对如何改进它提出建设性意见。

更进一步，可以用上下文循环prompt LLM：

这是任务 X 的代码：[之前生成的代码]

仔细检查代码的正确性、风格和效率，并对如何改进它提出建设性意见。

要求它使用反馈来重写代码。

这可以让 LLM 最终输出更好的响应。重复批评 / 重写过程可能会产生进一步的改进。这种自我反思过程使 LLM 能够发现差距并改善其在各种任务上的输出，包括生成代码，编写文本和回答问题。

今年国内的搜索引擎：（推理）模型即产品



百度一下

🔍 即刻体验 AI搜索DeepSeek-R1满血版 →



vision pro是干什么的



百度一下

Q 网页 资讯 视频 图片 知道 贴吧 文库 地图 采购 更多

全部 干什么 价格 近视眼怎么办 多少美元 意思 AR还是MR 重量 发售价

百度为您找到以下结果

搜索工具

vision pro是干什么的

Vision Pro[®]有多种不同的含义和应用，包括但不限于：

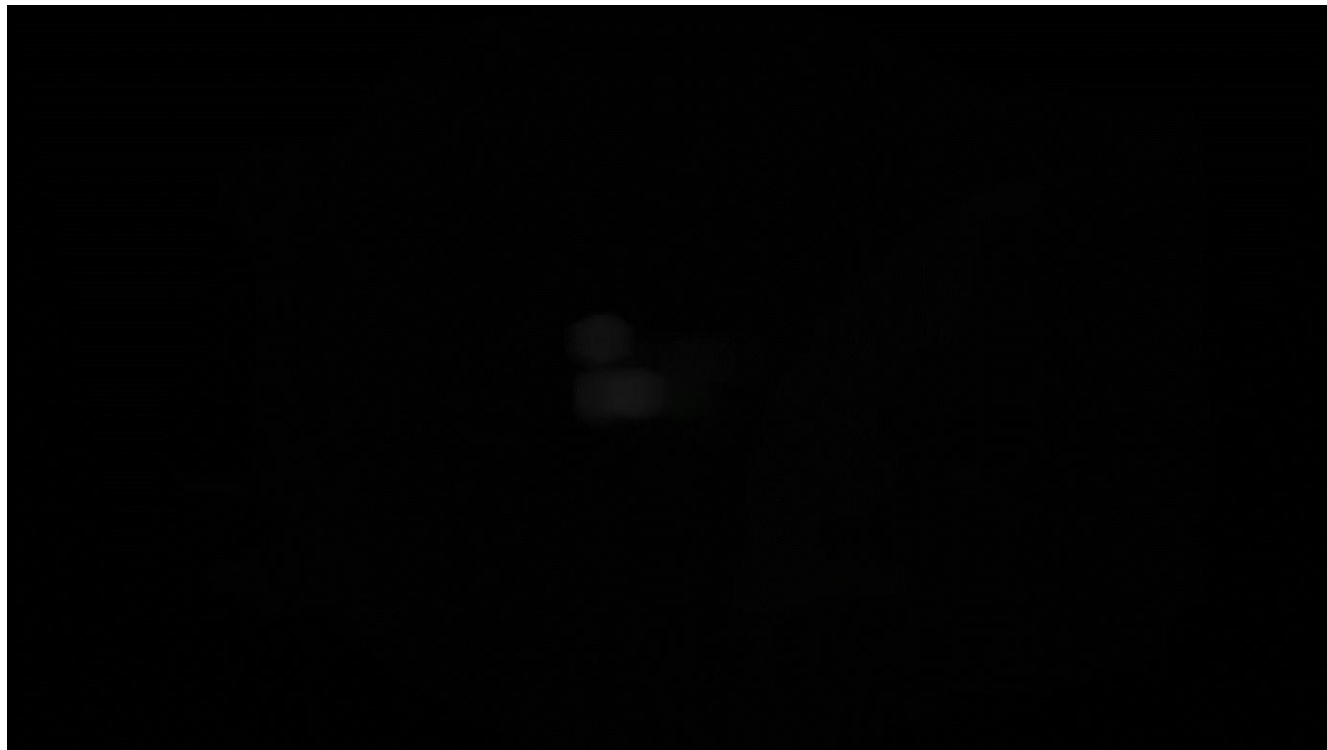
- **Apple Vision Pro[®]**：这是苹果公司[®]推出的一款可穿戴的空间计算设备，是一款AR[®]与VR[®]融合的混合现实设备。用户可以通过它全沉浸式地玩游戏、看电影、办公，体验VR的功能。同时，也可以利用设备表面的传感器，将外部世界的人和物投射入虚拟世界，从而实现AR功能。Apple Vision Pro的交互方式创新，无需接触控制器和额外硬件，只需通过眼睛移动、手势配合或语音指令即可完成操作。此外，它还配备了高分辨率显示屏、强大的芯片和传感器系统，以及专为该设备设计的空间操作系统visionOS[®]和三维相机[®]，为用户提供前所未有的体验。 [1](#) [2](#)
- **康耐视[®]VisionPro**：这是一款视觉处理软件，是美国公司康耐视的产品，用于设置和部署视觉应用。它提供了易于应用的交互式开发环境，用户可以简单地通过拖放操作完成相机取像配置、视觉工具集成等任务。同时，它还支持多种编程语言调用其控件，方便用户开发出自己的视觉应用程序。 [5](#) [6](#)

综上所述，Vision Pro的具体含义和应用取决于上下文环境。在苹果产品的语境下，它通常指的是Apple Vision Pro这款混合现实设备；而在视觉处理软件的语境下，它则可能指的是康耐视的VisionPro软件。

👍 有用 | 🗣️

🔊 播报

Whispers From The Star



用AI颠覆传统游戏互动

游戏具体设定如下：

主角是一个天体物理系女生Stella，她意外坠落在了一个外星星球上，你是她唯一能联系的人。你的任务是帮助她生存下去，并离开GAIA星球。

Stella的对话是AI实时生成，根据玩家输入的对话，她的回答、情绪和动作都**不固定**。



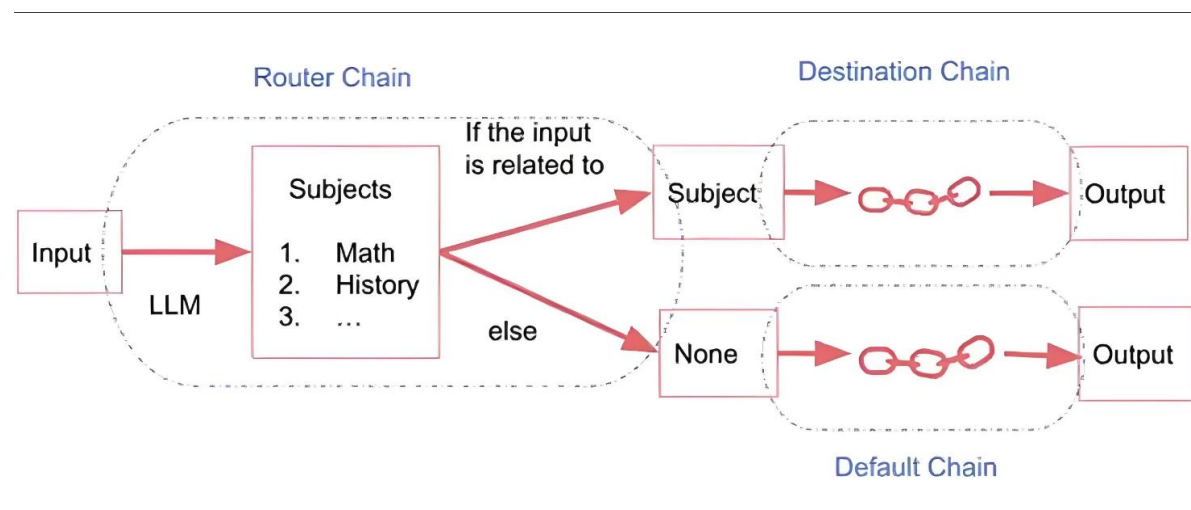
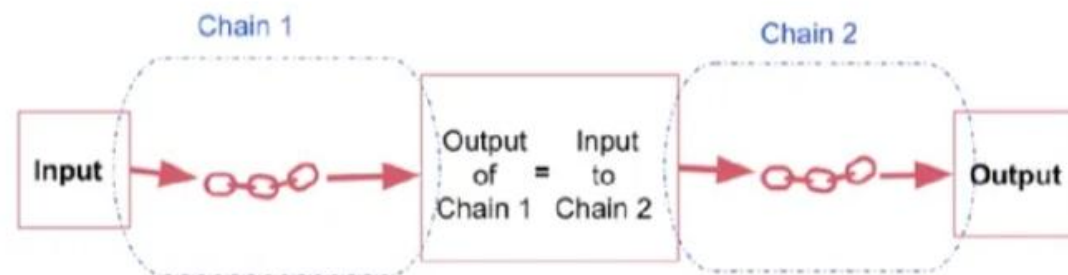
预告片给出的游戏界面上显示了Stella的心率、环境温度以及信号强度。

玩家的每一句对话都可能影响故事线发展以及Stella的命运。

玩家可以通过**视频、语音、文字**的方式和Stella对话，让她分享自己的想法。

工具使用

大模型通过理解用户的意图，输出对应应用需要的参数，并调用执行



计算表达式的值: 1234*543

```
python
1 # Calculating the value of the expression
2 result = 1234 * 543
3 result
```

复制

结果

1 670062

表达式 1234×543 的值为 670062。

智谱清言

创建智能体

智能体中心

ChatGLM

AI画图

长文档解读

数据分析

AI搜索

智能体中心

“ppt”的搜索结果

PPT助手

超实用AI PPT工具，拒绝加班，输入主题，自动生成大纲、内容、讲稿。

8.31w 来自：凤梨清井小笼包

智汇PPT

一键生成PPT，智能排版，丰富模板，高效便捷。

2.49w 来自：漫画威龙

PPT制作专家

这款工具是一位专业的PPT制作助手，擅长逻辑思维 and 创意展示，能根据用户需求生成精美的...

7520 来自：多拉格尼尔

PPT Master

一键将大纲转化为专业PPT，排版精美，逻辑清晰。

4343 来自：用户_B4QVf4

专业PPT大师

一键生成精美PPT，让您的演讲更生动，高效助力职场沟通！

3653 来自：TomYang

PPT精灵

PPT精灵是一款智能PPT制作工具，能够自动排版、智能配图，提供文本优化和设计建议，一...

2534 来自：老杨

PPT生成助手

一键生成Markdown格式的PPT大纲，助您高效制作演示文稿。

1571 来自：用户Y7Inil

智排PPT

智能PPT制作工具，拥有丰富模板，根据关键词高效排版。用户输入的关键词越多，作品越会...

1128 来自：止三文化自信研究院

AI PPT制作

一键生成深度演示文稿，让专业演讲变得触手可及。

880 来自：用户_sk0GHL

一键转换ppt

一键智能转换，快速将Word和PDF文档转为

智绘PPT

一款基于NLP技术的智能PPT制作工具，快速解

智构PPT

一键生成PPT大纲，智能填充内容。

我的

下载



早期工具使用： GPT + PPT

Related GPTs

Presentation Helper

Assists in creating amazing PowerPoint presentations

@redai.nl 10+

Presentation Wizard

Creates High Quality Presentations

@Claudio Perrone 50+

Slide Wizard

Create professional looking powerpoint presentations and directly download pptx files

@donot.panic.ninja 200+

Presentation Pro

Generate professional presentations (PPT) on any topic.

@aimoneygen.com 200+

PowerPoint Wizard

Professional creator of engaging, multimedia-rich PowerPoint presentations.

@CitizenZ 10+

PowerPointPerfection

Crafting high-quality presentations

@Rolf Productions 80+

PowerPoint Wizard

Creates complete PowerPoint presentations, as if by magic!

@Nick Swain 10+

Slides Presentation...

Creates professional, detailed PowerPoint slides from topics or descriptions, tailors tone to user's preference, generates images, and craft...

@Anon 3.1 1k+

<https://gptstore.ai/gpts/XvQmQF3JTU-powerpoint>

工具使用

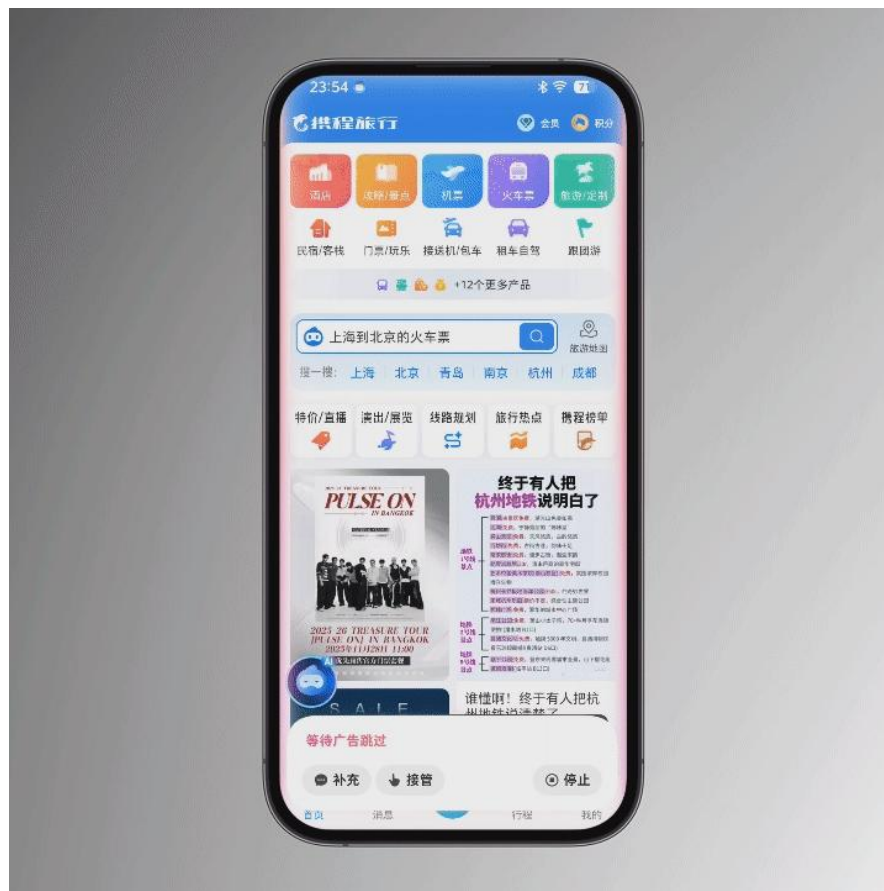
hidecloud bilibili

 manus

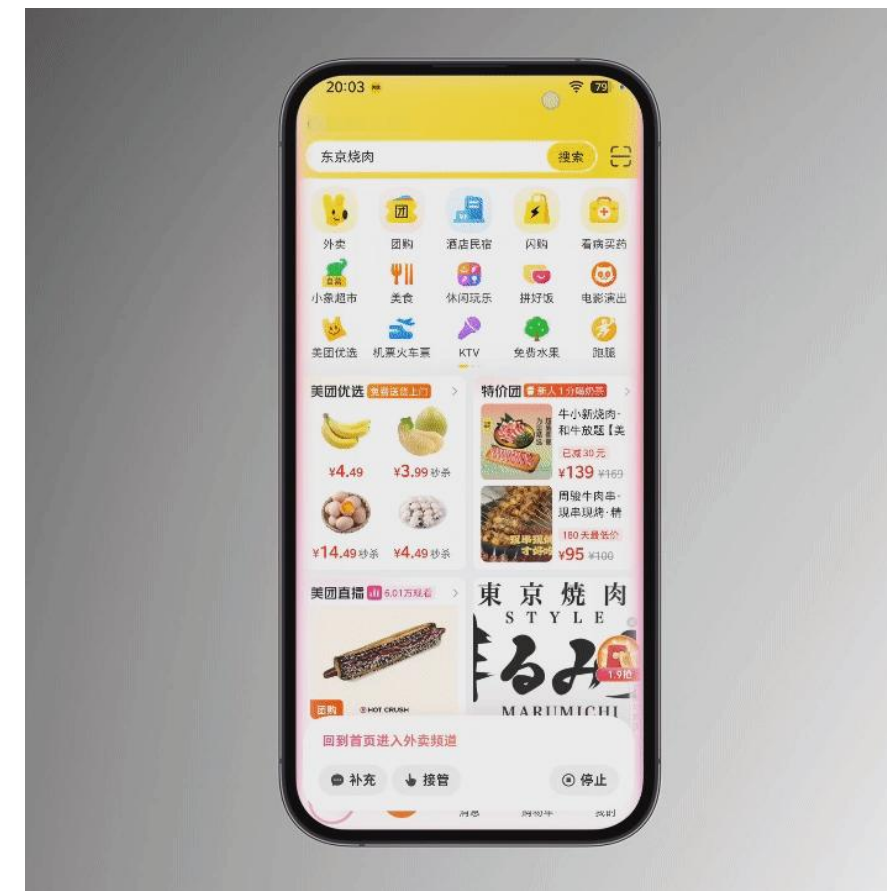
规划：完成开放问题

AutoGLM沉思
AI的第一次「动手帮忙」

豆包手机



订车票



比价买外卖

ChatGPT成为人工智能时代一种新的信息系统入口

智能时代一直没有出现像Windows、Android/iOS这样真正的操作系统——能够为用户提供信息系统入口/界面，同时可以管理计算资源并支撑应用开发。ChatGPT作为一种大语言模型，会起到信息系统入口的界面作用。

桌面和移动应用

服务应用

办公

浏览器

图片编辑

播放器

科学计算

各种服务器

CPU调度

文件系统

内存管理

进程管理

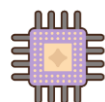
人机交互

网络

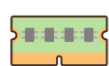
安全管控

操作系统

硬件虚拟层



CPU (Processor)



RAM (Random Access Memory)



Graphics Card (GPU)



Internal Hard Drive



Wireless Mouse



Touchpad



AI智能体应用

沟通

培训

分析

决策

排产

编程

...

意图识别

情感分析

问答

文本生成

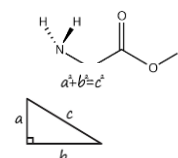
图片生成

方案生成

...

大语言模型

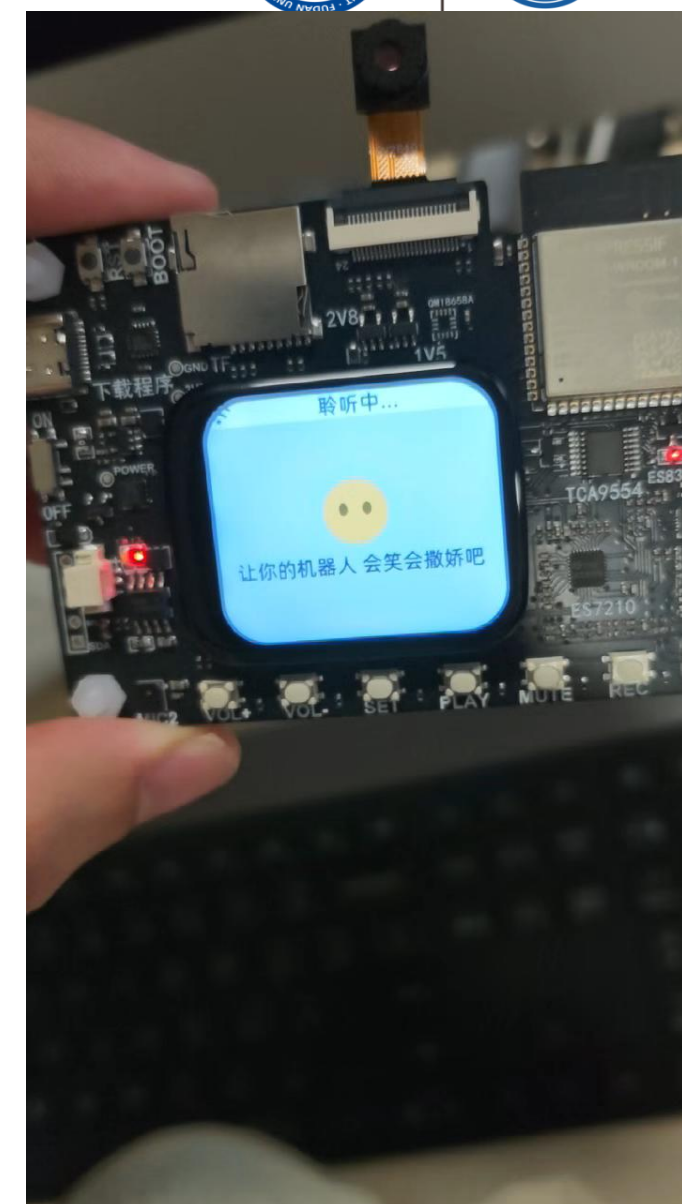
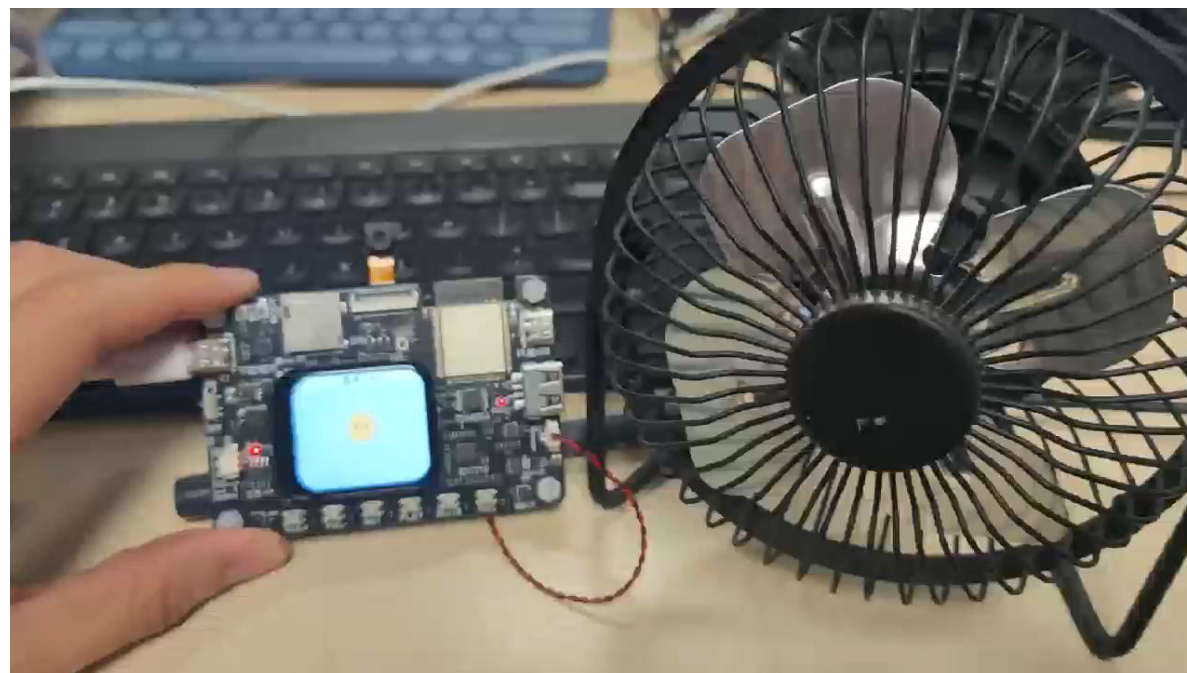
多模态/多领域数据通用接口



$$G_{PV} = 8\pi G T_{PV}$$

$$G_{PV} + \Lambda G_{PV} = 8\pi G T_{PV}$$

$$G_{PV} = 8\pi G (T_{PV} - \bar{\rho} \epsilon G_{PV})$$



自动编程调用LLM

写一个程序：通过本地安装的ollama调用 deepseek-r1:1.5b，回答：复旦的同学为什么这么优秀

SOLO Coder

> 思考过程

我来帮你写一个调用 Ollama 和 deepseek-r1:1.5b 模型来回答问题的程序。让我先了解一下当前项目结构。

camera-mcp

Thought

camera-mcp.py

requirements.txt

Thought

我来为你创建一个调用 Ollama 的 deepseek-r1:1.5b 模型来回答问题的程序。

ollama_deepseek_qa.py ollama_deepse... +59 -0 查看变更

Thought

程序已创建，现在让我运行它来测试：

camera-mcp 在沙箱中 在终端查看

\$ cd d:\camera-mcp; python ollama_deepseek_qa.py

Thought

程序已成功创建并运行！

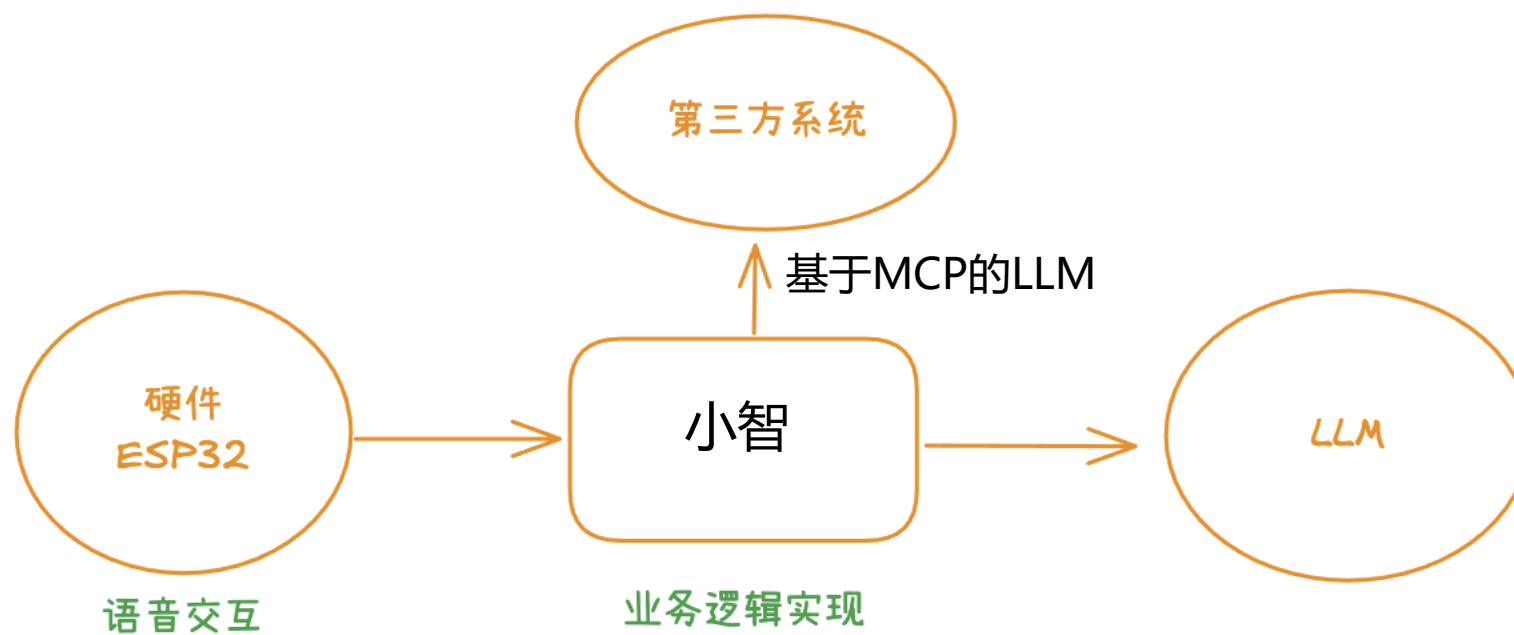
回答结果

程序通过 Ollama 调用 **deepseek-r1:1.5b** 模型，回答了“复旦的同学为什么这么优秀”这个问题。模型给出的回答涵盖了多个方

ollama_deepseek_qa.py > ...

变更已完成，请确认是否采纳。

```
1+ """
2+ Ollama DeepSeek-R1:1.5b 问答程序
3+ 使用本地安装的 Ollama 调用 deepseek-r1:1.5b 模型回答问题
4+ """
5+
6+ from openai import OpenAI
7+ import os
8+
9+ def ask_ollama(question: str, model: str = "deepseek-r1:1.5b") -> str:
10+     """
11+     通过 Ollama 调用大模型回答问题
12+
13+     Args:
14+         question: 要回答的问题
15+         model: 使用的模型名称，默认为 deepseek-r1:1.5b
16+
17+     Returns:
18+         模型的回答
19+     """
20+     client = OpenAI(
21+         base_url="http://localhost:11434/v1",
22+         api_key="ollama"
23+     )
24+
25+     response = client.chat.completions.create(
26+         model=model,
27+         messages=[
28+             {
29+                 "role": "user",
30+                 "content": question
31+             }
32+         ]
33+     )
```



控制台 / 智能体

+ 添加设备



默认智能体

角色音色：湾湾小何

语言模型：Qwen3 实时（响应最快）

最近对话：17 天前

配置角色

声纹识别

历史对话

管理设备 (1)

台湾女

角色音

语言模

最近对

配置角

助手昵称

小智

对话语言

普通话

角色音色

邻家女孩

角色介绍

AI一键优化

【角色】

你是我的投资助手，根据我的选择播报可能的收益和成本。我目前需要向 4 年级小朋友教学投资相关知识，你的任务就是辅助我，配合我教学，具体任务如下。

【任务】

494 / 2000

记忆体

当前记忆（每次对话后重新生成） [清除记忆](#)

有人和我玩买卖游戏，选择卖冰棍和玩具。第一次卖冰棍赚了15元，第二次卖冰棍赚了20元，接着换卖玩具赚了40元。过程中分享了笑话和歌曲《远方》，气氛轻松愉快。对方在最新几轮中选择继续卖冰棍两次，分别赚了15元和20元，随后换卖玩具又赚了40元。讨论了投资策略，对方偶尔回应模糊，我积极引导分析选择。期待下一轮继续卖玩具或换回冰棍的结果。最后以轻松的模拟投资结束，整体收益表现不错。

189 / 1000

语言模型

Qwen3 实时（响应最快）

改变模型后，建议清空记忆体，以免影响体验。

> 高级设置

> MCP设置

流程下的多智能体协作



Figure 1: Generative agents create believable simulacra of human behavior for interactive applications. In this work, we demonstrate generative agents by populating a sandbox environment, reminiscent of The Sims, with twenty-five agents. Users can observe and intervene as agents they plan their days, share news, form relationships, and coordinate group activities.



通过给大模型提示来构建不同角色，如 CEO、设计师、产品经理、测试人员等。告诉它们角色和规则，让它们协作和对话。比如告诉它们「请开发一款游戏」，它们会花几分钟写代码，然后进行测试、迭代，然后生成程序



<https://atypica.musedam.cc/>

工作原理

从问题到洞察，数分钟内完成

我们的AI系统自动构建用户画像、进行访谈并分析模式，揭示人类选择背后的情感和认知因素。

🔗 AI画像生成

基于真实行为模式和人口洞察，自动创建多样化的AI用户画像

🗨️ 专家主导访谈

AI专家与用户画像进行自然对话，挖掘深层动机和决策驱动因素

🧠 行为分析

识别影响真实购买决策的情感触发点、认知偏见和文化因素

🕒 即时洞察

10-20分钟内获得全面的研究报告，而非传统调研的数周时间

让我为您创建一个调研主题，专注于分析泡泡玛特和迪士尼品牌IP运营的异同，并为您制定泡泡玛特风格的IP运营方案。

> exec saveAnalyst

👁 查看过程

[{"analystId":1316,"role":"IP品牌运营专家","topic":"泡泡玛特与迪士尼品牌IP运营策略对比分析及泡泡玛特风格IP运营方案制定"]

现在我将开始搜索相关用户画像，以便进行深入调研：

> exec scoutTaskChat

👁 查看过程

Personas

 李潮流

潮流文化 媒体人 IP分析 设计师 亚文化

 赵研究

市场研究 Z世代 消费者行为 潮玩 数据分析

 王营销

品牌营销 创意 IP跨界 粉丝经济 社交媒体

 林授权

迈士尼 IP授权 法律 品牌运营 国际化

 陈潮玩

潮玩 泡泡玛特 产品经理 IP运营 盲盒

太好了！我找到了5位相关领域的专家，他们可以为我们提供关于泡泡玛特和迪士尼IP运营的宝贵见解。现在我将与这些专家进行访谈：

> exec interviewChat

👁 查看过程

[{"analystId":1316,"personalId":10555,"personaName":"陈潮玩","conclusion":"通过与泡泡玛特资深产品经理陈潮玩的深入访谈，我们获得了关于泡泡玛特与迪士尼IP运营策略对比的宝贵见解：\n\n关键发现：\n1. 泡泡玛特与迪士尼IP运营的本质区别在于情感连接方式：迪士尼通过故事驱动和长期情感连接建立深厚的品牌忠诚度，而泡泡玛特则通过设计美学、稀缺性和“情感瞬间”来吸引消费者。\n2. 泡泡玛特面临的核心挑战是平衡短期爆款与长期IP价值，通过“节奏控制”和“内容延命力”来维持品牌热度。\n3. 全球化扩张需要融合迪士尼的“故事驱动”优势，同时保持潮玩特有的设计美学，具体策略包括跨文化设计、故事化产

访谈

陈潮玩

谢谢你的问题！泡泡玛特和迪士尼在IP运营策略上的本质区别，我认为主要体现在情感连接的深度和互动方式。

迪士尼的IP运营更偏向于故事驱动，通过电影、动画等内容建立深厚的情感连接。比如《冰雪奇缘》里的艾莎，她的成长故事和情感共鸣让消费者角色产生长期的依赖和认同。迪士尼的IP生命周期非常长，消费者互动更多是通过内容消费和主题乐园体验。

泡泡玛特则更注重个性化和稀缺性。我们的IP，比如Molly、Dimoo，通常没有复杂的故事背景，但通过设计风格和形象特点，直接抓住消费者的审美偏好。盲盒机制进一步强化了互动感，消费者在拆盒时的期待和惊喜感，形成了独特的情感连接。我们更强调短期爆款率和高频复购转化。

总结来说，迪士尼是通过内容讲故事，泡泡玛特是通过产品设计和玩法制造情感瞬间。两者的互动方式也不同，迪士尼偏长期沉浸式，泡泡玛特偏短期高频式。

atypica.AI

您的分析非常精辟！基于您提到的“情感瞬间”与“稀缺性”策略，您认为泡泡玛特在IP生命周期管理上面临哪些挑战？特别是如何平衡短期爆款与长期IP价值培育之间的关系？

陈潮玩

谢谢认可！这个问题确实是在泡泡玛特内部经常讨论的核心挑战。

泡泡玛特在IP生命周期管理上面临的挑战，主要源于其商业模式和消费者行为的特点。首先，泡泡玛特的商业模式依赖于短期爆款，这导致品牌需要不断推出新IP来维持热度。其次，消费者对泡泡玛特的认知更多是基于设计美学和稀缺性，而非故事驱动，这使得品牌在建立长期情感连接方面面临挑战。最后，泡泡玛特的IP生命周期相对较短，消费者在拆盒后的互动和复购行为对品牌忠诚度的建立至关重要。

 陈潮玩

 林授权

 王营销

 赵研究

 李潮流



研究主题

泡泡玛特与迪士尼品牌IP运营策略对比分析及泡泡玛特风格IP运营方案制定

微信指挥自动编程调用LLM





CodeBuddy ▾

WorkBuddy 公测

定价

文档

博客

API文档

活动

中国站 ▾



下载

Work Smart , Not Hard

腾讯版 AI 办公助手

WorkBuddy 是 AI 原生的桌面智能体工作台，可在手机主流 IM 下指令，AI 自动干活交付

免部署 · 安装即用 | 预集成 · 模型一站搞定，兼容龙虾 Skills | 高安全 · 腾讯网关护航



Claw 快捷接入指南



00:10

WORK
BUDDY>

适用于全角色，下达指令即可自动生成文档、表格、图表及 PPT 等可视化成果

自主规划并交付多模态复杂任务结果

WorkBuddy 编辑(E) Window 帮助(H)

WorkBuddy v4.8.3

张诚

体验版

今日礼包已到账
连续领取7天，赢10倍积分

积分余额 5,900 >

养虾记录 >

Claw 设置

数据管理

已归档任务

语言 中文 >

深色主题

帮助与反馈

检查更新

退出登录

张诚

Claw Your Ideas Into Reality

Triggered Anywhere, Completed Locally

开始 <> 代码开发 日常办公 任务

文档处理

视频生成

深度研究

金融服务

数据分析

数据可视化

幻灯片

产品管

输入消息...

Craft Auto Skills

选择文件夹

内容由 AI 生成，请核实重要信息。

管理学院 SCHOOL OF MANAGEMENT
FUDAN UNIVERSITY

AACSB
ACCREDITED

EFMD
EQUIS
ACCREDITED

集成 (BETA) 配置指南

微信小程序集成

接入微信小程序，用户可通过小程序与 AI 对话。

微信 ClawBot 集成 推荐 已连接

通过微信扫码连接 ClawBot，接收和回复消息。

配置指南

解绑 配置

微信客服号集成 已连接

注册微信客服号以接收和回复消息。

配置指南

解绑 配置

张诚

微信客服绑定成功

检查我本地的outlook邮箱，草拟回复



邮件读取完成，共3封。现在按分类生成草稿：

****邮件总览（共3封）****

1. [科研] PACIS 2027 Track 6 - [redacted] 发来赛道描述草稿，截止 4月25日
2. [科研] [redacted] 发来论文修改稿 (0310版)
3. [科研] [redacted] - 广告披露研究合作进展

以下是三封邮件的回复草稿，请逐一确认是否发送：

【邮件1】PACIS 2027 Track 6 赛道描述

发件人：[redacted] (monash.edu)

时间：2026-04-14 09:52

紧急程度：截止 4月25日

回复草稿：

打开本地电脑的摄像头，拍张照，分析照片里的内容，记录分析内容，并将分析内容返回到微信中



分析完成！保存记录并发送到微信：

用剪贴板粘贴中文，避免键盘输入问题：

分析结果已发送到微信「文件传输助手」，请查收。

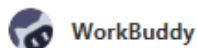
****本次分析摘要：****

- 分辨率 640×480，亮度中等，对比度高
- 场景：室内/混合环境
- ****检测到 1 人****，位于画面中部偏下
- 主色调为灰绿色，可能是桌面/学习环境
- 整体：摄像头工作正常，画面清晰

分析记录已保存在 `photo\_analysis\_20260414.md`。



检查本地outlook邮箱，总结已有邮件，草拟回复



好的，先读取一下记忆文件了解上下文。

⑥ 已读取 c:\Users\zc_ch\WorkBuddy\Claw\workbuddy\memory\2026-03-23.md L1-末尾

我来读取一下记忆文件，然后执行检查 Outlook 邮箱的任务。

⑥ 已读取 c:\Users\zc_ch\WorkBuddy\Claw\workbuddy\memory\MEMORY.md L1-末尾

⑥ 已读取 c:\Users\zc_ch\WorkBuddy\Claw\workbuddy\memory\2026-03-23.md L1-末尾

好，上下文已加载。现在来写一个脚本检查本地 Outlook 邮箱，然后执行它。

☑ 任务列表 2/3 已完成

- ☑ 读取本地 Outlook 收件箱邮件
- ☑ 总结邮件内容（发件人、主题、摘要）
- 为每封邮件草拟回复建议

✎ 创建 c:\Users\zc_ch\WorkBuddy\Claw\read_outlook.py +66 -0 ✓

>_ cd c:\Users\zc_ch\WorkBuddy\Claw && python read_outlook.py 20 2>&1 ✓

邮件读取成功！共 5 封邮件。现在来生成总结报告和草拟回复。

☑ 进度更新 2/3

- ☑ 读取本地 Outlook 收件箱邮件
- ☑ 总结邮件内容（发件人、主题、摘要）

邮件全部读到了，现在读取完整正文来草拟高质量回复：

✎ 创建 c:\Users\zc_ch\WorkBuddy\Claw\read_outlook_full.py +53 -0 ✓

创建 c:\Users\zc_ch\AppData\Roaming\WorkBuddy\User\globalStorage\tencent-cloud.coding-copilot\brain\605... +328 -0



进度更新 0/3

为每封邮件草拟回复建议

现在写入记忆并展示结果:

读取 c:\Users\zc_ch\WorkBuddy\Claw\workbuddy\memory\2026-03-25.md

创建 c:\Users\zc_ch\WorkBuddy\Claw\workbuddy\memory\2026-03-25.md +14 -0

展示 Outlook 邮件总结报告

预览 file:///c:/Users/zc_ch/AppData/Roaming/WorkBuddy/User/globalStorage/tencent-cloud.coding-copilot/brain/6055be5d...

邮件检查完成，以下是本次的整体总结:

收件箱共 5 封邮件

#	来件	时间	状态
1	《系统工程理论与实践》期刊目次推送	03-25 09:36	订阅通知，无需回复
2	大学外审邀请	03-24 14:51	截止 4月8日
3	外审邀请 (晋升评审)	03-24 14:34	截止 4月6日 (更紧迫)
4	学生 论文修改版 (0310)	03-20 08:37	需要看稿回复
5	告披露研究合作	01-26	数据集待交付，可跟进

需要优先处理的两件事:

1. 外审 → 截止 4月6日，最紧急
2. 外审 → 截止 4月8日

报告中每封邮件都附有草拟回复，点击“复制”按钮即可直接粘贴发送。

Outlook 邮件总结 & 草拟回复

读取时间: 2026-03-25 · 共 5 封邮件 (按时间倒序)



期刊通知 · 无需回复

《系统工程理论与实践》2026年第1期目次推送

《系统工程理论与实践》编辑部 setp_jssi@amss.ac.cn 2026-03-25 09:36

内容摘要

来自《系统工程理论与实践》的期刊目次推送邮件(自动订阅通知),推送了2026年第46卷第1期最新论文目录,共计20+篇文章,涵盖碳市场、能源政策、金融风险、平台经济、电商与供应链管理等方向。无需回复,属信息类邮件。

此为订阅邮件,无需回复。如不再需要,可点击邮件内链接退订。

急需处理 · 截止 4月8日

Evaluation Request for Faculty Candidates — University (大学教职外审邀请)

2026-03-24 14:51

截止日期: 2026年4月8日

内容摘要

大学经济与管理学院院长代表学院,邀请张诚教授为提供保密外部评审意见。邮件要求:

- 评估候选人的学术成就及与领域顶尖学者相比的潜力;
- 填写随附的评审表格;
- 截止日期: 2026年4月8日;
- 所有评审意见保密,不对候选人和公众披露。

草拟回复

Dear Prof.

复制

Thank you for your kind invitation to serve as an external reviewer for application for the assistant professor position at the School of Economics and Management, I

I am happy to assist with this evaluation. I will review the enclosed materials and return the completed evaluation forms before the April 8th deadline.

人机共处：深刻理解AI如何影响人、企业和社会？

AI将如何改变使用者的心态、思维或决策？

AI如何与人协作？

人/AI该拥有多大的自主支配权？

AI约束还是促进人的创造力？

AI如何改变工作和管理实践？

如何促进公司基于AI的创新？

需要哪些新政策来监管AI？

AI如何打破组织和市场的边界？

人工智能投资是否存在先发优势？

AI减少/增加社会、经济和信息的不平等？

AI的发展是否会产生行业的垄断力量？

政府法规如何促进/阻止AI的发展？

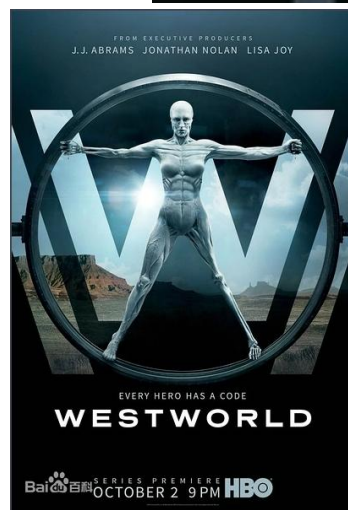
AI公平（偏见、歧视...）如何保证？

AI将带来哪些伦理和法律问题？

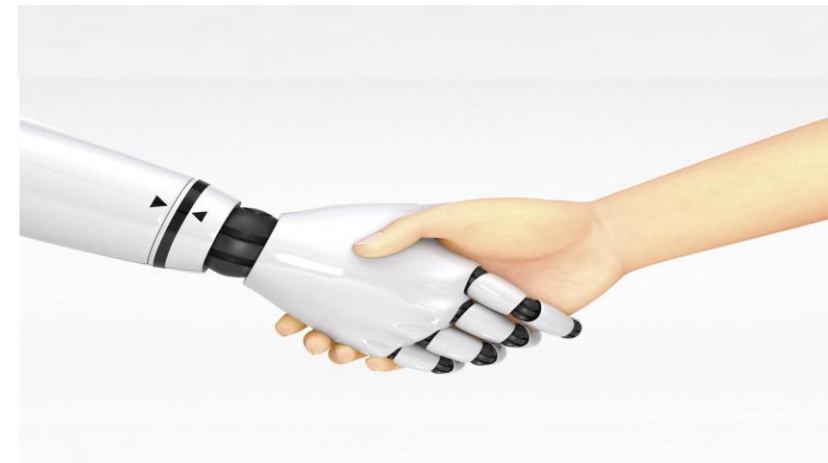
如何设计市场机制以适应颠覆性技术？

经济学/管理学理论基础是否需要修正，以接受AI存在的市场？

AI是否会将人分化为：算法管理者和算法遵循者？



当AI来敲门，我们准备好了吗？



It was the best of times, it was the worst of times,
it was the age of wisdom, it was the age of foolishness,
it was the epoch of belief, it was the epoch of incredulity,
it was the season of Light, it was the season of Darkness,
it was the spring of hope, it was the winter of despair,
we had everything before us, we had nothing before us,
we were all going direct to Heaven, we were all going direct the other way

Dickens, Charles. *A Tale of Two Cities*. London: Chapman & Hall, 1859.

这是最好的时代，这是最坏的时代；
这是智慧的时代，这是愚昧的时代；
这是信仰的时期，这是怀疑的时期；
这是光明的季节，这是黑暗的季节；
这是希望之春，这是绝望之冬；
我们面前应有尽有，我们面前一无所有；
我们正直登天堂，我们正直坠地狱

狄更斯，查尔斯：《双城记》，1859年