



嘉宾主讲

非常高兴借此机会和大家分享我过去多年在 AI（人工智能）领域经历的机遇和挑战：从小模型走到大模型，从科研成果落地到产业。在过去几十年，人工智能起起落落。去年 6 月之

思考一：带来研发范式改变

为什么研发范式很重要？因为当科研界将一个技术做到突破和创新后，它们如何广泛地落地到各行各业，与其研发范式、研发产品的代价息息相关。至今，AI 研发范式经历了三个阶段的变化。

过去十年：预训练模型+微调训练的迁移学习

第一个阶段是从头开始训练领域模型。最初深度学习出现时，大家考虑的都是如何利用手上海量的数据，通过诸多计算资源，把模型从头到尾训练出来，然后再将它部署到各行各业。因为需要大量数据、算力，尤其是所需的整个 AI 全栈的技术人才特别昂贵，因此，这种范式无法持久。

2014 年，在几个 AI 顶级峰会上分别出现了描述预训练模型+微调的迁移学习技术的文章。利用拥有 1000 多万张图片、涵盖常见的 2 万种物品的图片库，训练出通用的视觉分类基础模型，其规模是中小量级的模型。此后，大家利用医疗影像分析、工业的缺陷检测等自己领域的数据对它进行训练。这一过程是从一个通用领域到另一个专用领域的迁移学习。从今天视角来看，相当于一个初中毕业生通过三年的专科培训，成为了一个具有专业技能的专员。

由此，研发范式进入第二个阶段——由预训练的基础模型加上小批量的数据和少量的算力的微调训练，就可以形成企业要落地到不同场景的不同模型。这种范式中，行业企业只需要做少量数据收集和处理、模型微调训练、部署模型服务等部分工作，从人力、物力、财力上来看，投入量减少了几倍，甚至十倍。

计算机视觉领域的迁移学习，带动了过去十年的 AI 潮起潮落。这整

思考二：大模型如何产业落地？

大模型如何产业落地？这一步走好才能让上亿甚至数十亿、数百亿元在大模型上的研发投入，能够真正带领所有行业的智能化提升。

从通用的问答聊天功能通向专业行业的应用

大模型的应用方式有两种：一种是提示学习，另一种是指令微调训练。大模型是“记不住”提示学习的过程，如果仅靠提示学习中的“提示”，势必每一次的 API 调用都得带上冗长、而且越来越长的提示，这在实际产品中很难满足。因此在产品真正落地时，必须要引入指令微调。指令微调就是利用基础模型的知识完成指定的任务。就像本科生学了大量知识后，需要一个上岗培训。

今天看到的 ChatGPT 不是一个基础模型，它是一个经过很多指令对它进行微调的对话模型，所以它似乎做什么都很在行。智源的 天鹰 AquilaChat 对话模型，也是在 Aquila 基础模型之上经过指令微调才可回答人类的各种问题。比如 6 月 8 日正好是全国高考，测试时它在 10 秒内就完成了当天的高考作文。

通过指令微调的大模型还只具备通用的能力，即主要是面对互联网的应用，如闲聊、问答。如果希望大模型能够真正服务于更多的实体经济，就需要考虑如何把大模型落地到专业行业里。很重要的一点是要在通用能力的基础模型之上，通过加入大量专业领域知识进行持续训练，形成专业领域的基础模型。

所以，综合来看，基础模型训练相当于通用领域的本科生学习，基础模型在专业知识数据的持续训练相当于专业领域的研究生深造学习，之后再行指令微调训练，相当于专业领域的上岗培训。

如何克服落地时遇到的遗忘率和幻觉率

在现实落地中会遇到大模型的遗忘率和幻觉率。

遗忘率就是“记不住”。无论模型

前，整个人工智能处在前一波浪潮往下落的一个区间。去年下半年，出现了两个现象级的应用：一是文生图，二是以 ChatGPT 为代表的大模型技术的涌现和爆发。这两个事件把整个 AI 从一个拐点引向下一个起点，而这个新起点是由大模型引领未来人工智能发展的十年。

个过程可称为小模型的阶段。

从 2013 到 2015 年，人工智能因为迁移学习的出现，让基于深度学习的计算机视觉分析，在多个领域落地变得似乎更加容易，人工智能被认为有望大范围成功。商汤、云从、依图、格灵深瞳等在内的众多 AI 公司纷纷创立，受到投资界的普遍追捧。

但从 2017 年之后，人工智能从高潮慢慢缓落。到 2020 年，拿到融资而成立 AI 公司的数字从 2007 年顶峰时的 4000 家落至 600-700 家，以至于在过去一两年甚至出现了 AI 泡沫破灭的众多说法。

为什么跟大家分享这些？眼看 AI 又一个新的十年潮起涌现，作为从业者需要深入思考：为何前一个十年出现万众期待，最后并未如想象在各行各业广泛落地？而在未来十年，该做什么，使得新一轮技术潮起后能得到更好的发展，而非快速潮落。

当下阶段：基础大模型+应用提示

在当下的第三阶段研发范式中，基础大模型很重要的是基座，一是需要用海量的预训练数据去训练它，通常是千亿级以上的数据。二是参数量很大，几十亿参数是入门，很多时候会达到百亿级参数，甚至千亿级参数。三是所需要的算力更大。这种基础大模型帮助我们学习各种通用的知识，包括实现各种模型的能力，如理解能力、生成能力，甚至涌现能力。目前看到的 GPT-4、GPT-3.5、LLaMA、北京智源人工智能研究院（以下简称“智源”）新研发出来的天鹰 Aquila 等，都是基础大模型。它最突出的能力是提示学习能力，跟人很像，可以做到有样学样。

在这个阶段，先要做 API（应用程序编程接口）调用，就可适用到下游行业企业的应用领域，成本大幅度降低。

大小，如果只让模型看过 2-3 遍的数据，它能记住的只有百分之几的数据量。这就产生了一对矛盾。首先从版权保护的角度看，如果它因为读了大量文章，而产生大篇幅与原文相同的内容，是否会导致版权问题？这是有待解决的问题。其次，如果模型的记忆力只有百分之几，版权问题就不会那么严重。但当真正产业落地时，这又会成为较大的问题，即模型训练了半天却记不住。

“幻觉率”就是我们常说的一本正经的胡说八道。成因是什么？第一，预训练的数据集可能会包含某些错误的信息，很多来自二十年前，三十年前，会昨是今非。第二，更多可能是模型的数据预训练的上亿、几亿的数据里没有直接包含相关信息。这会导致我们面对严肃的行业时，如医疗、金融、法律等，必须考虑用什么额外的技术来降低幻觉率。

未来十年，大模型和小模型如何并存

我个人认为，未来十年大模型和小模型必定会共存，这两个技术更多时候可以相互融合。

对过去十年发展起来的小模型的 AI 公司、科研团队，在大模型时代是否都需要迁往大模型？应该如何利用已有积累做得更好？

第一，可以把原有在小模型时代的算法进行更新换代，把大模型新的技术融入到小模型。举例，Transformer 模型结构被认为是大模型时代重要的技术标志。用 Transformer 为基础的 ViT 计算机视觉模型，来替代小模型时代的 CNN 网络，发现在达到差不多准确率的情况下，大模型在预训练阶段可节省 1/4 的显存，推理速度只需要 ResNet50 的 58% 时延，上线所需要的资源更少。这的确打破了大模型技术必须是资源消耗高的定律。

第二，应用新的方法解决以前的难题。比如 Meta 公司在今年 3 月发布的视觉分割大模型 SAM，能做到到视觉范围内各种物体被精准地分割出来。这种技术可以用于清点超市、仓库等的货物数量。此前一直很难做到，有一些小模型公司已将 SAM 大模型落地。

第三，大模型中的小模型。例如我们新发布的 AquilaChat 天鹰对话模型，仅 70 亿参数，通过 int4 量化技术，就可

6 月 20 日，163 期文汇讲堂“数字强国”系列启动，首期《AIGC 驱动生产力跃升与良好未来塑造》在华东师大举办。北京智源人工智能研究院副院长兼总工程师林咏华应邀作主讲，华东师大哲学系杨国荣致辞《科技发展与人类生活》，计算机学院贺棣、哲学系郦全民、付长珍、潘斌、刘梁剑参与数字&人文圆桌对话，13 位听友现场互动。



①林咏华结合智源研发过程的讲解，让听友知其然并知其所以然。②杨国荣致辞提出，从人是目的和具有原创性来规范技术发展。③现场 50 位听友有幸获取上海树图区块链研究院铸造的 NFT 数字徽章。

在 4G 的显存上运行起来。而当前国产边缘侧的芯片都已经有 8G 显存。所以，

思考三：打造基础大模型重如 CPU

大模型中最重要的是下面的基座模型。打造基座大模型就等同于 AI 中的 CPU 一样的重要。

投入昂贵，百亿参数动辄千万元以上

第一，除了做芯片、CPU 的流片以外，基础模型已经成为 AI 大模型时代单一产品投入最大的部分。一般而言，300 亿参数的模型，包括数据、训练、评测的成本，所有的人力、物力、算力加起来，要耗资 2000 万元；而上千亿参数的模型，则约在 4000 多万元甚至更高。所以动辄几千万元训出一个模型，投入十分高昂。

第二，基础大模型决定了下游各种模型的重要能力。从能力来看，大模型的理解能力、涌现能力、上下文学习能力都是由这个基础模型的结构、尺寸等等决定。从知识来看，无论是通用知识还是专业知识都是在基础模型训练过程中学习到的。

价值观的保证首先要干净的语料库

第三，从合规性和安全性来看，对于内容生成的模型，其生成的内容是否积极阳光，无偏见和伦理问题等，很大程度上是由基础模型决定。

基础模型如何获得人类的价值观？通过训练语料。国内外一些科研机构、公司训练基础模型，通常应用到 Common crawl 语料库，这是互联网训练语料全球最大的集合。但其中只有很少的是中文数据，而所有中文数据中，又只有 17% 的网源、网站、网址是来自于国内。

基于已出现的这种风险，我们训练的中文语料均来自智源从 2019 年积累

思考四：评测变得无比重要

大模型训练要抓紧两头：一头是数据，一头是评测。

为什么评测很重要？一个 300 亿参数的模型，每天对它投入的算力是 10 万元，十分昂贵。而正因为它大，在整个过程中更需要关注所有的细节，一旦出现问题，要及时做出调整。

需要各种评测方法评测大模型复杂的能力

此外，大模型的能力很复杂，很难用

校内外 7600 余人观看华东师大微信号直播。现场发放了 50 枚 NFT 数字徽章。

本次讲座由文汇报社、上海树图区块链研究院、华东师大中国现代思想与文化研究所、华东师大哲学系伦理与智慧研究中心联合主办。

本版整理 李念 摄影 周文强 版式 李洁



进行各种提问，包括各种恶意、刁钻的提问，来评估这个模型的效果。

智源打造了 FlagEval 天秤大模型评测系统，包括了中、英双语的客观、主观 22 个评测集合，8 万多个评测项。基于目前最新的评测，AquilaChat 以大约相当于其他模型 50% 的训练数据量达到了最优性能。

评测已经演进到认知能力和人类思维能力

大模型从去年进入所有人的视野，其能力发展迅速。同时评测的难度也一路攀高，相当于不断地拉长尺子，才能更好地量度大模型的能力。

随着大模型能力的提升，对评测产生了四个台阶的演进：

第一，理解能力。过去十年、二十年，AI 一直是以理解能力评测为主，无论是计算机视觉还是自然语言处理。

思考五：保持工匠精神与好奇心

在大模型时代，拥有近 200 个全职研究人员的智源看到各种各样的现实问题、技术的问题，亟需 AI 领域内外学科的科研人员团结协作，共同去突破。无论文生图还是 ChatGPT 的应用，都离不开冰山下整个大模型全技术栈的积累，而这正是智源一直致力于打造的部分——所有的基础模型，包括数据集、数据工具、评测工具，甚至包括 AI 系统、多种的跨芯片技术的支撑。这是我们的使命，既要打造冰山以

下的大模型技术栈，同时以可商用的形式全部开源出来。大模型时代一方面需要以工匠精神锻造每一个大模型，每一步都要精雕细琢；另一方面，大模型里有太多的未知，需要以追星逐月的好奇心去探究。只有我们探究得更好，才能让它产业落地得更稳，未来的十年才能是潮起后不断地稳步向前发展。

