



综述

以 ChatGPT 为代表的生成式 AI 对通信行业的影响和应对思考

张嗣宏, 张健

(中兴通讯股份有限公司, 江苏 南京 210012)

摘要: ChatGPT 的发布引发生成式 AI 热潮, 代表着通用人工智能奇点时刻的到来, 信息产业生态系统极有可能重构, 国内产业界纷纷加强以 ChatGPT 为代表的智算领域研究, 运营商作为算网基础设施建设的主力军, 迎来智算发展的新机遇。详细剖析了生成式 AI 的能力、发展情况和应用前景, 思考了生成式 AI 背后的技术要素、对于算网资源的诉求、对通信行业带来的影响, 以及运营商在生成式 AI 发展浪潮中面临的机遇和挑战, 最后探讨了运营商的定位和应对策略。

关键词: ChatGPT; 生成式 AI; 通信行业; 智算

中图分类号: TP393

文献标志码: A

doi: 10.11959/j.issn.1000-0801.2023116

Impact and countermeasures of generative AI represented by ChatGPT on the telecom industry

ZHANG Sihong, ZHANG Jian

ZTE Corporation, Nanjing 210012, China

Abstract: The release of ChatGPT, sparked a wave of generative AI, representing the arrival of the singularity moment of general artificial intelligence and highly likely restructuring the information industry ecosystem. The domestic industry has strengthened the research in the field of intelligent computing represented by ChatGPT, and operators have become the main force in the construction of computing network infrastructure, ushering in new opportunities for the development of intelligent computing. A detailed analysis of the capabilities, development status, and application prospects of generative AI were provided, and the technical elements behind generative AI, the demands for computing network resources, the impact on the communication industry, the opportunities and challenges faced by operators in the development wave of generative AI were thought about. Finally, the positioning and response strategies of operators were discussed.

Key words: ChatGPT, generative AI, telecom industry, intelligent computing

0 引言

2022 年 11 月月底, OpenAI 公司发布人工智

能对话应用基于生成式预训练转换模型的对话程序 (chat generative pre-trained transformer, ChatGPT), 其展现出的意图理解能力、语言流



畅性、持续对话能力，一扫传统自然语言处理（natural language processing, NLP）应用能力较弱的形象，成为继 2016 年 AlphaGo 击败李世石之后的又一个爆点事件，在全球 AI 产业界掀起大模型、生成式 AI 的热潮。截至 2023 年 4 月国内外先后发布了十多个类似的大模型。

通信行业从业者也在思考如何看待以大模型为基础的生成式 AI，电信运营商的算网基础设施需要如何应对。本文将基于生成式 AI 发展趋势、关键技术要素、未来应用场景分析，提出生成式 AI 对于算力、网络、数据中心的关键需求，进而提出电信运营商应对这一轮 AI 发展热潮的建议。

由于生成式 AI 处于快速发展时期，颠覆人类认知的事件不断发生，基于当前情况的分析与判断也可能存在不足之处，所以这个阶段需要对任何可能性都保持开放态度，本文观点供读者参考。

1 ChatGPT 为代表的生成式 AI 的发展趋势

1.1 传统 AI 面临的问题及生成式 AI 带来的提升

近年来，深度学习成为推动 AI 发展的重要驱动力，但随着深度学习技术进入瓶颈期，产业发展面临多个挑战。首先，传统深度学习多有监督训练，对标注数据依赖性大，非常耗费人力成本；其次，模型的领域特性强，跨领域迁移能力较弱，通用性不强。此外，目前深度学习对人类感情的理解还停留在浅层次的语义层面，不具备良好的逻辑推理能力，无法真正理解用户意图。

以 ChatGPT 为代表的生成式 AI 有大算力、大数据、大模型的典型特点，首先，其基于大规模无标注数据进行预训练，再通过少量标注微调，大幅降低了对数据标注要求；其次，生成式 AI 的预训练大模型具有多模态（文字、图片、程序、视频等）、跨模态（“文生文”“文生图”“图生图”）内容生成能力，有较强的通用性和跨场景使用能力。此外，与传统 AI 相比，生成式 AI

能够更有效地捕捉用户的意图，以理解上下文，使对话更流畅自然，具有更强的逻辑和组织能力^[1-2]。

1.2 当前国内外生成式 AI 的发展现状

2017 年谷歌提出了 Transformer 结构，开创了自然语言预训练的新模式。2018 年 OpenAI 发布了 GPT-1，谷歌发布 BERT，随后微软、脸书（Facebook）等巨头陆续发布了一系列大模型。2022 年 11 月，基于 GPT-3.5 的 ChatGPT 正式发布，标志着生成式 AI 从量变走向质变。2023 年 3 月，ChatGPT 进一步升级至 GPT-4，在美国律师执照统考、研究生入学考试、医学自测考试等各类专业考试中超过了 70%，甚至 90% 的人类考生，展现了超越人类一般专业人士的能力，让人们感受到生成式 AI 带来的震撼。

随着大模型浪潮的席卷，国内各厂商也纷纷跟进，陆续发布自己的大模型。百度发布文心大模型，阿里推出通义千问大模型，腾讯推出混元大模型，华为推出盘古大模型，商汤发布日日新大模型，科大讯飞发布星火认知大模型……一个多月涌现十多个大模型。总体而言，当前国内大模型仍处于起步阶段，能力水平与 ChatGPT 等国外领先水平仍然有不小的差距，有待进一步的优化提升^[3]。

2 生成式 AI 的技术要素及其应用前景

2.1 生成式 AI 的技术架构和产业生态

生成式 AI 技术架构如图 1 所示，生成式 AI 从技术架构看主要包括 4 个层次：基础设施层、算法模型层、人工智能生成内容（artificial intelligence generated content, AIGC）服务层、AIGC 应用层。

（1）基础设施层

基础设施层的大模型训练和推理算力主要基于搭载图形处理器（graphics processing unit, GPU）或智算加速芯片的智算服务器承载，同时大模型训练还需要通过 InfiniBand 网络或基于融合以太

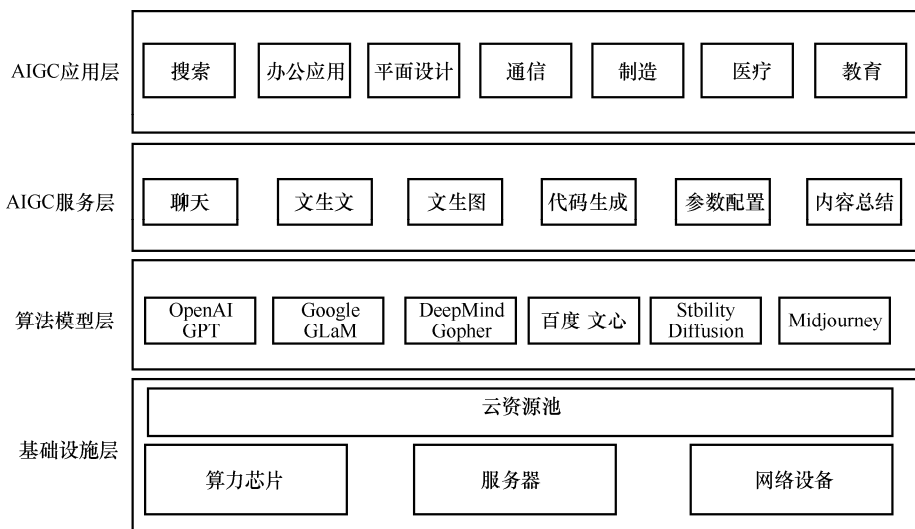


图 1 生成式 AI 技术架构

网的远程直接数据存取（remote direct memory access over converged Ethernet, RoCE）交换网络连接，形成高性能智算集群。在基础算网硬件之上，由云商整合各类硬件资源，封装后以算力服务的方式提供给上层模型训练使用。

（2）算法模型层

算法模型层基于基础设施层提供的算力，算法开发商进行通用大模型或者垂直领域大模型的训练和推理。大模型提供者主要有两类：第一类是技术实力较强，致力于构建通用大模型的科技公司；第二类是致力于构建专业领域大模型的提供者，大多基于开源算法框架，以较少的通用数据结合高质量专业数据，以更低的成本训练出专业领域的大模型。

（3）AIGC 服务层

AIGC 服务层在算法模型的基础上提供特定的生成式 AI 服务，如聊天、文生图、代码生成等，这类服务主要由大模型提供商提供。生成式 AI 服务也包括专业领域的内容生成能力，如生成网络配置指令等，这类服务通常由各领域专业厂商提供。

（4）AIGC 应用层

AIGC 应用层包括基于生成式 AI 技术的创新

应用和集成了生成式 AI 能力的传统应用。随着各类 AIGC 服务能力的完善和对外开放，各行各业的应用开发者都可以调用 AIGC 服务能力开发各类领域应用，应用生态将呈现百花齐放的局面^[4]。

2.2 ChatGPT 大模型的技术发展

ChatGPT 大模型的技术迭代之路如图 2 所示。

ChatGPT 为代表的生成式 AI 的技术迭代分为两个主要阶段。第一阶段激发模型潜力，侧重参数规模增大，该阶段通过对大量训练数据的学习，获得各种知识，并以模型参数的形式存储起来。通过不断增加模型参数，充分挖掘模型潜力，当参数量达到一定级别时，会呈现能力涌现的状况，爆发出大量新的能力，但是参数规模达到多大会出现能力涌现，产业界尚无明确结论，大多数观点认为在数十亿到千亿参数规模之间。第二阶段即人类反馈强化学习阶段，侧重人类意图对齐。该阶段又包含 3 个步骤：（1）监督微调，收集有代表性的问题，标注员标注问题对应的合理答案，通过少量经过标注的问答对模型进行微调；（2）设计一个模仿人类打分的奖励模型，对语言模型生成的回答打分，让符合人类意图的答案获得高分；（3）强化学习训练阶段，让语言模型和奖励模型进行交互训练，从而使得语言模型生成的内容不断趋近人类意图^[5-6]。

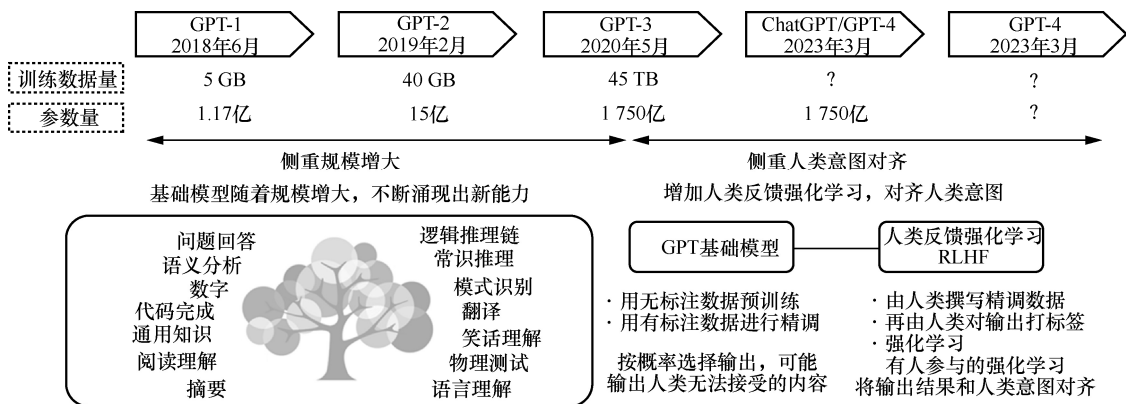


图2 ChatGPT 大模型的技术迭代之路

2.3 生成式 AI 的应用前景

ChatGPT 作为一个拥有接近人类表达水平的自然语言生成系统，凭借出色的知识整合和语言描述能力，将进一步打破人和机器的边界，在各个领域都具有巨大的发展潜力，有可能重塑众多产业生态。从其生成的内容形式来看，可以大致分为文本生成、代码生成、图像生成、音频生成、视频生成和其他内容六大类。

在文本生成方面，主要应用领域集中在商业办公、文字写作、医疗教育、智能搜索及个人应用方面。在商业办公领域，微软宣布将 ChatGPT 模型植入 Office 全家桶，并正式发布 Microsoft 365 Copilot。在代码生成方面，ChatGPT 可以帮助开发人员快速生成程序代码，编写速度更快、注释更清晰、理解更方便且代码效率更高，将带来研发效率的显著提升。在图像生成方面，可代替人类设计师创作一些图像作品，如广告海报、电影海报等，还可以根据输入的文字提示生成复杂的插图和场景，大大提高电影、游戏等娱乐领域的生产效率。

生成式 AI 被产业界广泛认为是人工智能与内容领域深度融合的代表，是变革未来生活和工作模式的底层核心技术。尤其在数字技术与传统产业融合的背景下，生成式 AI 将与大数据、物联网、云计算、数字孪生、XR 等数字技术一同构筑数字技术底座，驱动产业数字化变革升级^[7-8]。

3 生成式 AI 对通信领域的影响分析

3.1 生成式 AI 对算力的需求和影响

生成式 AI 大模型的训练和推理对算力的需求不同。大模型训练是计算密集型处理，需要高性能 AI 集群。而推理是多事务并发处理，数据中心或者边缘 DC 部署的 GPU 服务器即可满足要求。

以 GPT-3 为例，完成一个模型需要进行多轮训练，一轮训练所需算力为 3 640 PFLOPS×day，需要高性能 AI 集群（性能优化数据中心（performance optimized datacenter, POD））保障算力。一个训练作业在一个 POD 内完成，一般不跨 POD。

以英伟达 A100 POD 为例，其包含 140 台服务器，1 120 块 A100 GPU，算力可达到 336 PFLOPS。其中，GPU 有二级互联架构：在服务器内部，8 块 GPU 间使用高达 600 GB/s 的 NVLink 总线互联；服务器之间使用 InfiniBand 网络，连接各服务器的 8 个 200 Gbit/s 网口，可为服务器提供 1.6 Tbit/s 的互联带宽。使用 A100 POD 可以在 20 天内完成一轮 GPT-3 级别模型的训练。

高性能 AI 集群在大幅提升算力的同时，也使得 AI 芯片及服务器功率有了较大提升，单台服务器功率达到 5 kW 以上，液冷服务器及液冷数据中心将逐渐成为刚需。

生成式 AI 的推理一般不需要专用高性能 AI 集群。对于单个模型实例的推理需求，单台 GPU

服务器,甚至单块高性能 GPU 卡即可完成。但随着用户并发数上升以及图片、视频内容的增加,推理算力需求将大幅度增长,需要多台普通 GPU 服务器才能完成。

相比而言,中小模型自 2016 年起已经广泛应用于各种场景,中小模型参数量小于 10 亿,通常在千万到亿数量级,随着 GPU 处理能力的增加,中小模型训练已经不需要多块 GPU 并行处理,而推理算力要求更低,所以中小模型使用支持高速串行计算机扩展总线(peripheral component interconnect express, PCIe)接口的 GPU 服务器即可满足需求。

3.2 生成式 AI 对数据中心网络的影响

生成式 AI 大模型对数据中心网络的影响主要体现在模型训练所使用的高性能集群,大模型训练需要在传统数据中心网络中为高性能集群的每个 POD 单独构建一个高速数据交换平面,用于不同服务器上的 GPU 间远程直接数据存取(remote direct memory access, RDMA)数据传输。高速数据交换平面需要有如下能力。

(1) 高带宽

大模型训练需要将算法、数据拆分到数百或者数千块 GPU 卡上,因此 GPU 卡间需要较高的互联带宽。当前服务器内部 GPU 之间的互联总线带宽已经达到数太比特每秒,如 A100 NVLink 总线带宽达到 4.8 Tbit/s(600 GB/s),而单个服务器对外仅能提供 200 Gbit/s $\times$ 8 合计 1.6 Tbit/s 的接入能力,因此,服务器之间的网络带宽已成为制约 AI 集群性能的瓶颈。对于服务器之间的组网,国外大多采用 InfiniBand 网络,由于 InfiniBand 网络技术封闭,国内更倾向采用 RoCE 无损以太网。当前无论是 InfiniBand 网络还是 RoCE,都已经开始引入 100 Gbit/s/200 Gbit/s 接入端口和 400 Gbit/s 汇聚端口,并开始向 800 Gbit/s 端口能力演进。

(2) 低时延

大模型的训练过程中跨 GPU 数据交换频繁发

生,除高带宽外,低时延对于大模型的训练也非常重要。InfiniBand 网络时延最低可以低于 1 μ s,采用 RoCE 技术的无损以太网时延目前在 5~10 μ s 水平,需要进一步优化。

(3) 零丢包

RDMA 传输对丢包的容忍度极低,千分之一的丢包率将导致传输效率急剧下降,2%的丢包率将导致 RDMA 吞吐率下降为 0。因此,以太网承载 RDMA 协议时,丢包率要尽可能小,最好能做到零丢包。

除上述 3 个能力要求外,采用无损以太网构建高性能数据交换平面时,设备和组网模式对网络规模、性能也会造成一定影响。采用全盒式两层 CLOS 组网可以支持千块 GPU 卡全互联,而使用框式(Chassis)交换机单层 CLOS 组网即可实现千块 GPU 卡全互联,使用 Chassis 加盒式交换机两层组网可实现超大规模万块级 GPU 卡全互联,不过两层组网相较单层组网时延有少许增加。

除上述成熟的无损以太网设备组网模式外,还可以使用分布式解耦机框(distributed disaggregated chassis, DDC)新型组网模式。DDC 可提供端到端确定性流控,基于信元交换的超低时延,其时延和无损性能可比肩 InfiniBand 网络,但目前技术和产品尚不成熟,还需要进一步发展完善。

对于中小模型的训练以及大、中、小各类模型的推理,不需要 GPU 间高速数据交换,当前数据中心网络技术和组网均可以满足,无须特殊的网络设计。

3.3 生成式 AI 对通信大网的影响

在 toH/toC 领域,用户更多依赖生成式 AI 进行内容制作并在网络中传输,导致用户对网络的流量需求、带宽要求、网络质量进一步提升;在 toB 领域,AIGC 导致办公和生产环节对网络依赖度进一步提升,增加工作时间段网络流量,对网络的可靠性和服务质量也有更高要求。两类场景不仅会推



动互联网总流量提升,也将影响目前的互联网流量潮汐效应,使得流量周期内峰谷间的变化更加平缓。

本文以如下参考模型估算:用户提交问题及 ChatGPT 返回答案,平均一次交互假定 1 000 字,根据抓包结果单向交互大概为 3 MB。以每天每用户提交 10 次提问请求,一天业务量在 24 h 均匀分布进行估算,2023 年 2 月,ChatGPT 平均每天的访问用户数约为 3 500 万,用户交互的南北向流量约为 $(35 \times 10^6 \times 3 \times 10^6 \times 10 \times 8) / (24 \times 60 \times 60) = 97.2 \text{ Gbit/s}$ 。即使使用强度再放大 100 倍达到 10 Tbit/s 级别,带宽在全球互联网 1 000 Tbit/s 体量中也仅占 1%,文本型问答产生的南北向网络流量有限,对互联网带宽几乎无影响。

大模型对带宽增长的推动效果受大模型生成内容的类型和应用规模的影响,随着大模型多模态能力的成熟,信息的交互媒体形式从文本向语音、图片、视频、全息影像的高阶形态升级,带宽增长将带来数量级的变化。近年来,互联网骨干网带宽年增长率已下降至 20% 水平,当前应跟踪 ChatGPT 类大模型应用产生的流量变化情况,进行建模评估并及时应对。

当前大模型训练需要采用集中式部署,在一个数据中心内部的 POD 中完成训练,而推理节点可按需采用集中式或分布式部署。这种训练/推理分离模式下,训练 POD 与推理节点间只需要传递训练完毕的推理模型(文件大小为 GB 量级),这与目前的互联网应用服务的分布式部署机制相同,对东西向流量无明显影响。

4 运营商在生成式 AI 发展浪潮中的定位和应对

4.1 运营商的定位探讨

运营商在生成式 AI 产业链的定位及能力要求如图 3 所示,结合 AIGC 技术架构可以将大模型 AIGC 产业链中的参与者大致分为 3 类角色:

基础设施服务商,提供训练和推理所需的算力、网络服务;算法模型提供商,完成大模型的训练、部署、迭代,以应用程序接口(application program interface, API)或其他方式提供服务;AIGC 应用服务商,调用前者提供的模型能力,开发面向特定领域的应用,如代码自动生成应用、图片自动生成应用等。对电信运营商而言,选择哪一类或几类角色,对能力要求不同,未来的收益与风险也不同,需要结合自身特长和未来战略定位考虑。

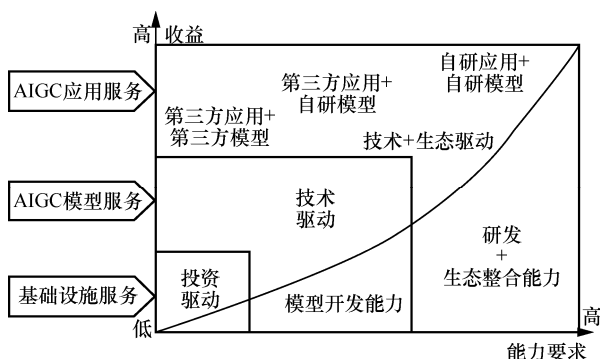


图 3 运营商在生成式 AI 产业链的定位及能力要求

笔者认为提供基础设施服务是基本定位,与电信运营商目前的云网发展战略一致,也是运营商擅长的服务模式。当前三大运营商均已通过专业公司对外提供云网基础服务,如天翼云、移动云、联通云等,面向大模型的训练、推理需求,通过增强智算算力,即可对外提供智算服务。

从 AIGC 服务安全性考虑,未来服务政府或信息敏感单位的大模型需要高安全、高可靠能力,电信运营商应积极布局。算法模型服务属于技术驱动,对于运营商的研发能力要求较高,需要引入高水平的 AI 科学家,积累相关技术能力。此类服务技术门槛高,未来收益也将比基础算网服务高。

对于最上层的 AIGC 应用服务,将是非常丰富的生态发展模式。运营商可以基于自主训练的大模型,首先针对通信服务自身构建 AIGC 应用,再结合自身优势市场,如智慧家庭领域,自主开发或与第三方企业共同开发 AIGC 应用,服务于

家庭用户健康、教育、娱乐等需求,或面向政企用户提供 AI 解决方案。

总体而言,运营商可以结合自身的能力、特点选择不同的角色定位。提供基础设施服务是起步模式,随着运营商在生成式 AI 方面技术能力和生态整合能力的提升,可逐步提供基础设施+模型+应用的综合性产品服务^[9]。

4.2 运营商的智算及网络规划建议

大模型训练所需算力比较特殊,由于模型参数量巨大,训练过程中需要在数百上千个 GPU 加速卡之间以 RDMA 方式高速传送中间数据,所以大模型训练需要高性能集群架构的算力。而大模型推理对算力的要求不高,一般单台 GPU 服务器就可以满足单个大模型推理。未来随着大模型推理进一步优化,单次推理的算力需求有望降低到一块 GPU 卡,因此大模型推理过程并不需要高性能 AI 算力集群。

(1) 训练侧智算算力规划建议

从 Gartner 技术成熟度曲线看,当前生成式 AI 正处于炒作曲线的高峰,资本市场关注度高,大量互联网或 AI 初创企业跃跃欲试,呈现出对大模型训练算力旺盛的需求。以百度、阿里、腾讯为代表的主流厂商已经自主构建智算算力,基本不会租用运营商训练算力。随着生成式 AI 技术达到炒作高峰后会经历泡沫破灭的阵痛期,初创公司的算力需求有可能会大幅减小,因此需要理性看待当前生成式 AI 所带来的价值和挑战^[10]。所以运营商建设训练集群需要综合考虑如下因素。

- 整体投资回报率以及业务发展的可持续性。
- IT 类产品生命迭代周期较快,折旧周期短。
- 当前由于需求激增,英伟达 GPU 的交付周期普遍较长,同时溢价较高。

建议运营商集团层面统筹考虑智算训练中心的布局规划,统一建设 1~2 个训练集群,首先服务于运营商内部场景的模型训练需求,如网络的运营运维、视频 AI 等,其次考虑对外面

向科研机构、高校、初创公司等提供训练算力服务。对于省级公司,可以关注省内政企行业客户的业务需求,提供中小模型的训练算力服务。部分热点省份,如果客户大模型训练需求较多且投资回报明确,可以考虑按需建设省内训练集群。

(2) 推理侧智算算力规划建议

AIGC 推理侧的算力没有特殊架构需求,普通的 GPU 服务器即可满足,其总算力需求将随用户数量增加逐渐增长。同时大模型推理算力同样可以满足中小模型决策式 AI 的训练要求。综合前述两类需求,推理侧算力可在现有的云资源池中增加 GPU 算力池,具体部署位置可以考虑按省份规划。中小模型训练基本无时延约束,部署省份主要考虑省内用户发展便利。而大模型推理初期时延要求不高,带宽占用有限,可考虑省内集中规划,未来随着 AIGC 应用和用户数量的增加逐渐下沉。

(3) 智算网络规划建议

大模型训练集群内需要高速网络连接所有的服务器,要考虑建设配套的 RoCE 或 InfiniBand 网络。但集群间并无大量数据高速传送的需求,所以集群间只要是普通数据中心网络即可满足需求,无特殊规划需求。

推理侧流量模型与普通互联网业务相比未见明显差异,未来可跟踪 AIGC 应用中视频内容占比的变化,适时做好广域网络的规划,未来 3 年内 AIGC 不会对广域网络造成重大冲击。

4.3 运营商的大模型及应用发展建议

OpenAI 的 GPT 模型演进为通用大模型训练指明了两个重要的技术路线:其一是预训练加人类反馈强化学习的训练模式,其二是增加参数规模产生能力涌现,不断激发模型的潜力。运营商自主训练也将遵循上述两个重要的技术路线。

预训练算法已经有大量公开论文以及部分



开源代码,甚至有部分训练数据集作为基础。但完成可达到 GPT-3.5 水准的大模型预训练,仍然需要解决算法优化和数据集优化等问题,特别是可充分发挥 GPU 集群能力的算法框架优化,这一过程需要一批高水平的 AI 科学家作为团队核心。基于预训练模型进行精调时,可以选择发展通用模型方向,也可以选择发展垂直行业方向。

相比通用大模型,垂直行业大模型只聚焦某个领域,技术门槛相对低一些,所以建议运营商初期可基于部分开源大模型进行预训练,并结合自身通信领域的数据优势,优先在通信领域大模型入手,首先满足自身网络智能化、智能客服等业务需求,通过构建通信领域大模型,验证技术可行性,积累大模型技术能力。再逐步对外提供垂直领域的模型服务,特别是面向政府或特定行业,运营商可借助安全可靠能力优势,为政企敏感行业提供垂直模型服务。长期来看,通用大模型仍然是终极发展目标,建议在垂直大模型取得进展之后,考虑向通用大模型训练发力^[11-13]。

5 结束语

ChatGPT 的热潮席卷全球,开启了人工智能新一轮增长,为千行百业的发展带来了新的想象空间,也给通信行业带来新的机遇和挑战。对于电信运营商而言,一方面需要进一步加强与智算相关的基础设施建设,积极跟进 AI 领域前沿技术发展及基础研究。基于垂直行业及关键场景拓展相关业务布局,积极探索 AI+多模态融合的新场景,挖掘新兴 AI 技术的落地应用。另一方面,也需要高度关注生成式 AI 在数据安全、个人隐私、知识产权等方面带来的风险,积极研究相关监管防范策略,引导产业良性发展,从而在新一轮的人工智能技术革命浪潮中把握先机^[14]。

参考文献:

- [1] 罗戈,张新鹏. 聚焦 ChatGPT: 发展、影响与问题[J]. 自然杂志, 2023, 45(2): 106-108.
LUO G, ZHANG X P. Focusing on ChatGPT: development, impact, and issues[J]. Nature Magazine, 2023, 45(2): 106-108.
- [2] 贵重,李云翔,江为强. 人工智能的发展与挑战: 以 ChatGPT 为例[J]. 电信工程技术与标准化, 2023, 36(3): 24-28.
GUI Z, LI Y X, JIANG W Q. Development and challenge of artificial intelligence—taking ChatGPT as an example[J]. Telecom Engineering Technics and Standardization, 2023, 36(3): 24-28.
- [3] 胡滨雨. 从 ChatGPT 爆火看人工智能大势[J]. 中国电信业, 2023(3): 36-41.
HU B Y. Looking at the trend of artificial intelligence from the explosion of ChatGPT[J]. China Telecommunications Trade, 2023(3): 36-41.
- [4] 史占中,郑世民,蒋越. ChatGPT 与 AIGC 产业链[J]. 上海管理科学, 2023, 45(2): 12-14.
SHI Z Z, ZHENG S M, JIANG Y. ChatGPT and AIGC industry chain[J]. Shanghai Management Science, 2023, 45(2): 12-14.
- [5] 赵志泉,王东波. 数字化时代下 ChatGPT 的开端、发展和影响[J]. 科技情报研究, 2023(2): 37-47.
ZHAO Z X, WANG D B. The beginning, development and impact of ChatGPT in the digital age[J]. Scientific Information Research, 2023(2): 37-47.
- [6] 钱力,刘熠,张智雄,等. ChatGPT 的技术基础分析[J]. 数据分析与知识发现, 2023, 7(3): 6-15.
QIAN L, LIU Y, ZHANG Z X, et al. Analysis of the technical basis of ChatGPT[J]. Data Analysis and Knowledge Discovery, 2023, 7(3): 6-15.
- [7] 刘禹良,李鸿亮,白翔,等. 浅析 ChatGPT: 历史沿革、应用现状及前景展望[J]. 中国图象图形学报, 2023, 28(4): 893-902.
LIU Y L, LI H L, BAI X, et al. Analysis of ChatGPT: historical evolution, application status, and prospects[J]. Chinese Journal of Image and Graphics, 2023, 28(4): 893-902.
- [8] WANG F Y, YANG J, WANG X X, et al. Chat with ChatGPT on industry 5.0: learning and decision-making for intelligent industries[J]. IEEE/CAA Journal of Automatica Sinica, 2023, 10(4): 831-834.
- [9] 胡滨雨,郭敏杰. 新一轮 AI 爆发下电信运营商的挑战和机遇[J]. 中国电信业, 2023(4): 23-27.
HU B Y, GUO M J. Challenges and opportunities for telecom

- operators under the new round of AI explosion[J]. China Telecommunications Trade, 2023(4): 23-27.
- [10] 刘聪. 由 ChatGPT 浪潮引发深入思考与落地展望[J]. 数字经济, 2023(3): 50-52.
- LIU C. Deep thinking and implementation prospects triggered by the ChatGPT wave[J]. Digital Economy, 2023(3): 50-52.
- [11] 兰祝刚, 任然. ChatGPT 对运营商影响及发展建议[J]. 中国电信业, 2023(4): 28-31.
- LAN Z G, REN R. Impact of ChatGPT on operators and development suggestions[J]. China Telecommunications Trade, 2023(4): 28-31.
- [12] 廖军, 谈鹏驹, 张冬月, 等. 人工智能在电信网络中的应用研究[J]. 信息通信技术与政策, 2020(9): 31-34.
- LIAO J, TAN P J, ZHANG D Y, et al. Research on the application of AI in telecom networks[J]. Telecommunications Network Technology, 2020(9): 31-34.
- [13] 陈俊明, 张洁, 左罗. 运营商 AI 能力建设及演进探讨[J]. 邮电设计技术, 2021(12): 83-88.
- CHEN J M, ZHANG J, ZUO L. Discussion on construction and

evolution of operators' AI capability[J]. Designing Techniques of Posts and Telecommunications, 2021(12): 83-88.

- [14] VAN D E A M, BOLLEN J, ZUIDEMA W, et al. ChatGPT: five priorities for research[J]. Nature, 2023, 614(7947): 224-226

[作者简介]



张嗣宏 (1982-), 男, 中兴通讯股份有限公司中国区战略规划总工程师, 主要研究方向为算力网络、云网融合、人工智能、大数据、SDN/NFV、5G 网络及业务发展等。



张健 (1979-), 男, 中兴通讯股份有限公司中国区战略规划总监, 主要研究方向为人工智能、算力网络、云网融合、大数据等。