

大模型 AI 在 PaaS 平台上的性能监控与优化技术

程超

云学堂信息科技(江苏)有限公司北京分公司

DOI:10.12238/acair.v2i3.8644

[摘要] 随着人工智能技术的快速发展,大模型AI在各行各业的应用日益广泛。PaaS(Platform as a Service)平台作为云计算服务的重要组成部分,为企业提供了快速构建和部署应用的环境。然而,大模型AI在PaaS平台上的运行对性能监控与优化提出了更高要求,尤其是在GPU(图形处理单元)资源的利用方面。本文旨在探讨大模型AI在PaaS平台上的性能监控与优化技术,分析其技术原理、实施策略及未来发展趋势,为相关从业者提供理论支持和实践参考。

[关键词] 大模型AI; PaaS平台; 性能监控; 优化

中图分类号: X924.3 **文献标识码:** A

Performance Monitoring and Optimization Techniques of Large Model AI on PaaS Platform

Chao Cheng

Yunxuetang Information Technology (Jiangsu) Co., Ltd. Beijing Branch

[Abstract] With the rapid development of artificial intelligence technology, the application of large model AI is increasingly widespread in all walks of life. PaaS (Platform as a Service) platform, as an important component of cloud computing services, provides enterprises with an environment for quickly building and deploying applications. However, the operation of large model AI on the PaaS platform puts forward higher requirements for performance monitoring and optimization, especially in terms of the utilization of GPU (Graphics Processing Unit) resources. This paper aims to explore the performance monitoring and optimization techniques of large model AI on the PaaS platform, analyze its technical principles, implementation strategies and future development trends, and provide theoretical support and practical reference for relevant practitioners.

[Key words] Large Model AI; PaaS Platform; Performance Monitoring; Optimization

引言

PaaS平台通过将计算资源、存储资源和网络资源封装成一个独立的虚拟环境,提供给开发者使用,从而简化了应用开发和部署的复杂度。大模型AI因其强大的数据处理和学习能力,在PaaS平台上得到了广泛应用。然而,大模型AI的复杂性和资源消耗特性,特别是对GPU资源的大量需求,使得其在PaaS平台上的性能监控与优化成为一项重要而复杂的任务。

(1)PaaS平台的发展与应用。随着云计算技术的不断发展,PaaS(Platform as a Service)平台作为一种提供应用开发、部署和管理环境的服务模式,在企业中的应用日益广泛。PaaS平台为开发者提供了便捷的基础设施和服务,使他们能够专注于应用的开发和创新,无需过多关注底层的硬件和软件管理。

(2)性能监控与优化的重要性。性能监控与优化成为确保平台稳定运行、提高用户满意度和降低运营成本的关键环节。有效的性能监控可以及时发现潜在的性能瓶颈,而优化技术则可以针对性地解决这些问题,提升平台的性能和资源利用率。

1 性能监控与优化的实践案例

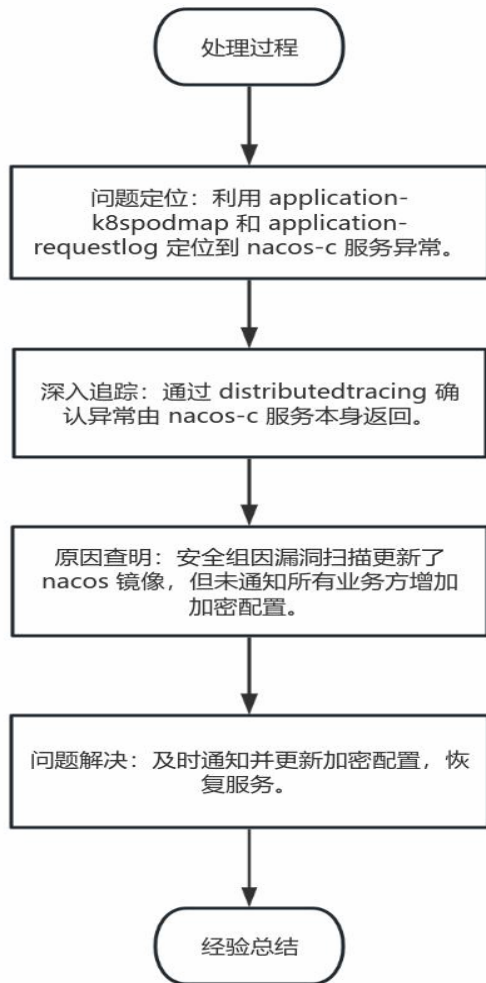
1.1案例一:中国移动磐基PaaS平台的故障定界与性能隐患发现

背景介绍:磐基PaaS平台自身业务模块已微服务化并部署在Kubernetes平台上,其可观测性能力基于eBPF实现,可提供全栈无盲点的调用链追踪等功能。

1.1.1业务页面空白-服务注册中心变更

运维人员发现磐基-巡检页面一直没有数据刷新且页面空白,但F12未报任何错误。通过查看application-k8spodmap,发现很多服务访问nacos-c存在异常(包括巡检功能)。接着在application-requestlog页面看到不少调用nacos服务报403(forbidden)异常。在distributedtracing中对存在异常的调用发起追踪,确认异常是nacos-c服务本身返回。经排查,原来是安全组在漏洞扫描时发现当前nacos版本存在漏洞,负责nacos的研发组更新了镜像,需要使用nacos服务的业务方在代码层面增加加密配置,但未及时通知其他研发组,影响了业务正

常运行。通过eBPF监控告警,发现了问题根因并及时进行了处理。处理过程如图1:



(图1)

1.1.2业务潜在性能问题-jedis未升级

在日常巡检中,利用application-k8spodmap面板发现e-k-a*服务访问redis服务存在约30%的异常。结合application-requestlog,确定是hello3命令的请求,但服务端返回“-errunknowncommand`hello”的错误。推测是redis服务端版本低于客户端版本,不支持hello命令(hello命令在redis version 6.0.0以后才支持)。对异常调用发起distributedtracing也证实是redis服务端返回的error。经研发确认,弹性计算的部分业务中使用的jedis客户端版本v3.6.1高于redis服务端v5.0.14版本,导致hello命令报错。进一步调查发现,jedis在v3版本上刚支持redis v6存在bug,会不停发送hello命令,升级到v4后,若发现redis服务端是低版本,则会改为发送auth命令。为规避此报错,与存在报错的业务方沟通后,将jedis统一升级到v4.3.2版本以方便统一管理。

1.1.3业务高峰期不可用-服务503异常

业务方反馈磐基平台的监控中心偶现访问异常(503 service unavailable),运维团队并未主动挂过503。由于业务

服务调用关系复杂,分析多日志也未找到原因,测试环境也未复现。最终利用基于eBPF的可观测性能力,快速定位到问题服务。通过application-k8spodmap发现100...访问m--c服务存在25%的服务端异常情况,在application-requestlog面板分析得知调用/web*-p*-c*//h接口存在503异常。利用distributedtracing面板对异常调用发起追踪,发现是m--c服务返回了503异常。经分析,由于早上上班后大家会集中查看监控数据,带来高峰数据。而整个运营运维业务一直在快速迭代发布,生产部署的微服务pod的副本数相同,通过高峰数据和eBPF的red指标及调用链追踪功能,快速找到了服务瓶颈点。与业务运维同事沟通后,给m--c在最大副本数基础上增加了1倍(增加的副本数根据线上的red指标最终调整得到最优)。

1.2腾讯云音视频PaaS平台与NVIDIA的合作案例

背景:腾讯云音视频PaaS平台专注于技术产品,构建了行业中极速高清智能转码、超低时延快直播的音视频解决方案。其中腾讯云mps媒体处理服务为海内外客户提供全场景极速高清转码、音视频增强、专有云codecsdk等服务。

在该案例中,画质增强是其提供的主要视频类云服务之一,通过3D去噪、色彩增强、超分辨率、插帧等处理技术提高画面清晰度。服务工作管线基于ffmpegfilter构建,其中应用了去噪、色彩增强以及超分辨率等AI模型,这些模型的工作管线除了模型推理,还包含了一系列前后处理操作,如yuv2rgb、copymakeborder、normalize、hwc2chw、chw2hwc、float2uint8、rgb2yuv等。此前,这些前后处理操作使用自定义的CUDA算子在GPU上处理或使用OpenCV在CPU上处理,计算效率并不理想,这使得前后处理在整个工作管线中占据了较大开销,也可能增加CPU负载,影响其执行效率。

为了加速画质增强全链路工作管线,腾讯云音视频PaaS平台与NVIDIA团队合作,利用cv-cuda™加速视频增强AI工作管线中的前后处理模块,结合NVIDIATensorRT™,将视频增强AI全流程置于GPU上进行加速。集成cv-cuda后相比加速前,前后处理部分效率提升16%-38%,端到端效率提升6%-10%。

具体来说,对于1080p输入图像,使用cv-cuda后三种模型工作管线的处理速度有了明显提升。例如,去噪模型的前后处理时间从原本的约14.5ms减少到约12ms;色彩增强模型从约6.5ms减少到约4ms;超分辨率模型从约42ms减少到约30ms。

2 大模型AI在PaaS平台上的性能监控技术

在大模型AI的应用中,GPU(图形处理单元)和CPU(中央处理单元)都发挥着重要作用,但它们具有不同的特点和优势。

CPU是计算机的核心处理器,擅长处理复杂的逻辑运算和顺序控制任务。它具有强大的通用性,可以执行各种类型的计算任务,但在并行处理大量数据方面相对较弱。

GPU则专门为并行处理大量数据而设计,具有大量的核心和高带宽的内存。在大模型AI的训练和推理中,特别是涉及到大量矩阵运算和并行计算的任务时,GPU能够提供比CPU更高的性能和效率。

2.1性能监控指标

大模型AI在PaaS平台上的性能监控指标主要包括以下几个方面(如图2所示)：

CPU利用率	反映大模型 AI 在计算过程中的 CPU 占用情况
内存利用率	反映大模型 AI 在运行过程中对内存的占用情况
网络带宽利用率	反映大模型 AI 在数据传输过程中的网络带宽占用情况
响应时间	反映大模型 AI 对请求的响应速度
吞吐量	反映大模型 AI 在单位时间内处理请求的能力

(图2)

2.2性能监控技术

大模型AI在PaaS平台上的性能监控技术主要包括以下几种：

2.2.1基于指标监控

通过预设的数据采集策略和智能代理程序,实时采集大模型AI的各项性能指标,并进行存储和展示。常见的指标存储方法包括时间序列数据库、时间序列缓存和日志存储等。指标展示则通过控制面板、大屏幕展示和实时数据展示等方式进行。

2.2.2基于日志监控

通过监控大模型AI的日志信息,发现潜在的性能问题。日志采集可以通过操作系统内核接口、应用组件提供的接口或第三方Agent进行。日志存储则可以选择文件日志存储、数据库日志存储或云日志存储等方式。日志展示则通过文本或表格的形式进行。

3 大模型AI在PaaS平台上的性能优化技术

3.1资源优化

资源优化在大模型AI于PaaS平台的性能提升中起着至关重要的作用,尤其是对GPU资源的优化。它涉及到对计算资源(包括CPU和GPU)、存储资源和网络资源的精心配置与高效调度,从而为大模型AI的运行创造理想的环境,大幅提升其运行效率。

弹性伸缩是其中一项关键策略。它能够依据大模型AI的实时负载状况,智能化地调整计算资源的分配。当面对负载的骤然增加,比如大量数据的涌入需要进行复杂的分析处理时,系统能够迅速感知并自动增加CPU和内存等资源的分配,确保处理过程的流畅和高效。反之,当负载减轻,资源需求降低时,又能及时释放多余的资源,这种动态的调整不仅保障了性能,还极大地节省了成本。

3.2算法优化

算法优化是提高大模型AI性能的根本途径。通过优化算法结构、参数设置和训练过程,可以显著提升大模型AI的准确性和效率。

模型剪枝:例如,在深度神经网络中,可以基于神经元的激活值或连接权重的大小来判断其重要性。对于那些激活值较小或权重较低的神经元或连接,可以进行剪除。在剪枝过程中,可以采用逐层剪枝的策略,先从对模型性能影响较小的层开始。

量化压缩:可以采用均匀量化的方法,将模型中的浮点数参数按照固定的步长转换为整数。或者使用非均匀量化,根据参数的分布特点进行更精细的量化。为了减小量化对模型性能的影响,可以在量化前对模型进行预训练,让模型适应一定程度的精度损失。

知识蒸馏:在这个过程中,大型的教师模型通常是一个复杂但性能优越的模型。在训练学生模型时,除了使用常规的标签数据,还引入教师模型的输出作为额外的监督信号。

3.3并发与并行优化

在任务调度方面,可采用基于优先级的任务调度算法。先对大模型AI任务进行分类,按照紧急程度、资源需求等因素赋予不同的优先级。例如,对于实时性要求高的任务赋予高优先级,优先分配资源。并行计算是提升大模型AI处理速度的重要手段。在多核处理器环境下,可以采用数据并行的方法,将数据分割成多个子集,分配给不同的核心进行处理。例如,在图像识别模型的训练中,将大量图像数据分别分配给多个核心同时进行特征提取和模型训练。对于分布式计算资源,采用模型并行策略,将大模型拆分成多个部分,分布在不同的计算节点上同时计算。

在负载均衡方面,采用基于硬件的负载均衡器,如F5负载均衡设备,或者基于软件的负载均衡方案,如Nginx等。通过设置合适的均衡算法,如轮询、加权轮询、最少连接等,将大模型AI的访问请求均匀分配到多个服务实例上。实时监测各个服务实例的负载情况,包括CPU使用率、内存占用、网络带宽、GPU使用率等指标。当某个实例负载过高时,自动调整分配策略,将更多请求分配到负载较低的实例上,从而平衡各个服务实例的负载压力,提高系统的稳定性和可扩展性。

4 结束语

总之,大模型AI在PaaS平台上的性能监控与优化技术需要不断发展和创新。通过持续的研究和实践,能够不断提升大模型AI的性能,为用户提供更高效、更可靠的服务。随着技术的进步,我们有望看到更智能、更精准的监控和优化手段的出现,进一步推动大模型AI在PaaS平台上的广泛应用和深度融合。

[参考文献]

[1]叶菁.大模型融入云平台AI+云竞赛打响起跑枪[N].通信信息报,2024-06-12(002).

[2]孔彬.基于AI大模型的智慧运维平台设想[J].广播与电视技术,2024,51(03):166-169.

[3]邹贺铨.大模型融入云平台,信息化走向数智化[J].重庆邮电大学学报(自然科学版),2024,36(01):1-8.

[4]郭治豪,廖素明,覃广荣,等.基于PaaS的流程引擎与党建平台适配性研究[J].现代信息科技,2024,8(02):82-85+91.

作者简介:

程超(1980--),男,汉族,河北张家口人,研究生,研究方向:平台架构,AI。