



团 体 标 准

T/CECA-G 0355—2025

数据中心能效“领跑者”评价要求

Assessment requirements for top-runner energy efficiency of data centers

2025-07-01 发布

2025-07-02 实施

中 国 节 能 协 会 发 布

目 录

前言 II

1 范围 1

2 规范性引用文件 1

3 术语和定义 1

4 评价指标 3

5 主要配套设备要求 3

6 计算和测试方法 4

附录 A（规范性）通算数据中心测试方法 9

附录 B（资料性）归一化系数设置方法 10

附录 C（资料性）模型相关影响因素说明 11

附录 D（资料性）推荐测试模型信息 12

附录 E（资料性）AI 模型超参数和 FLOPs 计算值示例 13

前 言

本文件按照 GB/T 1.1—2020《标准化工作导则 第 1 部分：标准化文件的结构和起草规则》的规定起草。

请注意本文件的某些内容可能涉及专利。本文件的发布机构不承担识别专利的责任。

本文件由中国节能协会提出并归口。

本文件起草单位：中国标准化研究院、中国移动通信集团设计院有限公司、中国信息通信研究院、青海省质量和标准研究院、青海省发展改革委、施耐德电气（中国）有限公司、北京英沣特能源技术有限公司、北京智网数科技术有限公司、北京真视通科技股份有限公司、中电投工程研究检测评定中心有限公司、浪潮电子信息产业股份有限公司、厦门吉快科技有限公司、浪潮计算机科技有限公司、海光信息技术股份有限公司、广东美的暖通设备有限公司、中数（南京）检测认证有限公司、合修能（北京）节能环保科技有限公司、国家气象信息中心、北京阿玛西换热设备制造有限公司、天翼物联科技有限公司、国投中标质量基础设施研究院有限公司、天津华信惠悦科技有限公司、北京康斯特仪表科技股份有限公司、威凯检测技术有限公司、浙江柿子新能源科技有限公司。

本文件主要起草人：彭妍妍、张学斌、陈鸣飞、孙丽玫、李洁、叶晓剑、郭亮、姜宇光、王月、韩玉仲、杨硕、白莉婷、张培、杨朋霖、赵昱、孙立峰、刘海潮、文莉萍、邹元霖、娄小军、董现、孔艳荣、张佳琪、姜凌菲、朱梓厚、马莉莎、卫翔宇、何冠成、游政贤、张新昌、邱惠悻、王洪旭、郑立新、李金波、江兴方、雷海涛、张译文、郭浩、宋小亮、王晶、王丽峰、朱承兴、王政、栗新汉、程辉、李宗华、江高炎、田张新、孙小亮、王珏旻。

本文件为首次发布。

数据中心能效“领跑者”评价要求

1 范围

本文件规定了数据中心能效“领跑者”的能效评价指标、测试和计算方法。

本文件适用于已建、新建或改扩建的数据中心建筑单体或模块单元能源效率和碳排放效率的评价和管理。

2 规范性引用文件

下列文件对于本文件的应用是必不可少的。凡是注日期的引用文件，仅所注日期的版本适用于本文件。凡是不注日期的引用文件，其最新版本（包括所有的修改单）适用于本文件。

GB/T 778.1 饮用冷水水表和热水水表 第1部分：计量要求和技术要求

GB 17167 用能单位能源计量器具配备和管理通则

GB/T 32910.1 数据中心 资源利用 第1部分：术语

GB 40879 数据中心能效限定值及能效等级

GB 43630—2023 塔式和机架式服务器能效限定值及能效等级

GB/T 50063 电力装置电测量仪表装置设计规范

GB 50174 数据中心设计规范

GB 50462 数据中心基础设施施工及验收标准

YD/T 4024 数据中心液冷服务器系统总体技术要求和测试方法

T/CCSA 327—2021 数据中心碳利用效率技术要求和测试规范

T/CCSA 370—2022 数据中心网络能效比测量规范

3 术语和定义

GB 50174、GB 40879、GB/T 32910 界定的以及下列术语和定义适用于本文件。

3.1

数据中心 data center

由信息设备场地（机房），其他基础设施、信息系统软硬件、信息资源（数据）和人员以及相应的规章制度组成的实体。

3.2

智能数据中心 intelligence data center

指通过使用大规模异构算力资源，包括通用算力（CPU）和智能算力（GPU、FPGA、ASIC 等），主要为人工智能应用（如人工智能深度学习模型开发、模型训练和模型推理等场景）提供所需算力、数据和算法的数据中心。

3.3

液冷 liquid cooling

一种采用液体带走发热器件热量的数据中心产品，适用于需提高计算能力、能源效率、部署密度等应

用场景。

注：液冷分为接触式及非接触式液冷两种，接触式液冷是指将冷却液体与发热器件直接接触的一种液冷实现方式，包括浸没式和喷淋式液冷等具体方案。非接触式液冷是指冷却液体与发热器件不直接接触的一种液冷实现方式，包括冷板式等具体方案。

3.4

电子信息设备 **electronic information equipment**

对电子信息进行采集、加工、运算、存储、传输、检索等处理的设备，包括服务器、交换机、存储设备等。

3.5

数据中心总耗电量 **total Electricity consumption of data center**

维持数据中心运行所消耗电能的总和。

注：包括电子信息设备、冷却设备、供配电系统和其他辅助设施的电能消耗。

3.6

数据中心信息设备耗电量 **electricity consumption information devices**

数据中心内各类信息设备所消耗电能的总和。

注：用于数据信息计算、处理、存储、传输、交换的计算机设备、网络设备、服务器设备、存储设备等信息设备直接用电。

3.7

数据中心电能比 **ratio of electricity consumption of data center**

统计期内，数据中心在信息设备实际运行负载下，数据中心总耗电量与信息设备耗电量的比值。表征数据中心基础设施层的电能利用效率（Power Usage Effectiveness, PUE），比值越接近 1 能源利用效率越高。

3.8

数据中心信息设备碳使用效率 **carbon usage effectiveness of data center**

数据中心能源使用产生的碳总排放量与信息设备使用能源比值。表征数据中心基础设施层碳使用效率（Carbon Usage Effectiveness, CUE）。

3.9

面向通用计算的数据中心算力能效 **computing energy efficiency of general data center**

统计期内，数据中心综合算力与数据中心总用电量的比值，表征通算数据中心算力能效（Computing Energy Efficiency, CEE）。比值越高表明能源利用效率越高。

注：面向通用计算的数据中心指以常规服务器、通用存储设备和标准网络交换设备等构建，可支撑各类常见业务数据处理的硬件设施集合。

3.10

面向通用计算的数据中心算力碳效 **computing carbon efficiency of general data center**

统计期内，通用数据中心能源使用所产生的总碳排放量与综合算力的比值，表征通算数据中心算力碳效（Computing Carbon Efficiency, CCE）。

3.11

面向 AI 应用的数据中心算力能效 **computing energy efficiency of data center for AI applications**

统计期内，在面向 AI 应用的数据中心执行 AI 相关任务的实际运行负载下，该数据中心完成的浮点

数计算次数与总能耗的比值，表征面向 AI 应用的数据中心算力能效（AI Energy Efficiency，AEE）。比值越高表明能源利用效率越高。

注：面向 AI 应用的数据中心指以 GPU、TPU 等 AI 加速芯片为核心，搭配高速内存、低延迟高速存储设备以及适配的高速低延迟网络交换设备构建，专为人工智能模型训练、推理等应用提供强大算力支撑的硬件设施集合。

3.12

面向 AI 应用的数据中心算力碳效 **computing carbon efficiency of data center for AI applications**

统计期内，面向 AI 应用的数据中心执行 AI 相关任务的实际运行负载下能源使用所产生的总碳排放量与综合算力的比值，表征面向 AI 应用的数据中心算力碳效（AI Carbon Efficiency，ACE）。

4 评价指标

4.1 面向基础设施的数据中心能效评价指标

面向基础设施的数据中心能效先进水平应满足 GB 40879—2021 中 4.1 规定的能效 1 级，同时满足第 5 节的要求。面向基础设施的数据中心能效节能水平应满足 GB 40879—2021 中 4.1 规定的能效 2 级，同时满足第 5 节的要求。

4.2 面向通用计算的数据中心能效评价指标

采用相关测试软件工具对通算数据中心进行测量，综合算力能效先进水平和节能水平应符合表 1 的规定。

表 1 面向通用计算的数据中心评价指标

单位：标准分每瓦时

指标	能效等级	
	先进水平	节能水平
通算数据中心能效比	≥50	≥20

4.3 面向 AI 应用的数据中心能效评价指标

采用自然语言处理、计算机视觉、商业应用相关模型作为测试工具进行能效测试，实际应用能效先进水平应符合表 2 规定。

表 2 面向 AI 应用场景的实际能效指标先进水平

指标	先进水平/GFLOPs/Wh
面向 AI 应用的数据中心能效比	≥224

5 主要配套设备要求

5.1 冷却系统

制冷机组的能效应符合以下要求：

- 冷水机组应符合 GB 19577 中一级能效要求
- 单元式空气调节机应符合 GB 19576 中一级能效要求
- 溴化锂吸收式冷水机组应符合 GB 29540 中一级能效要求
- 水(地)源热泵机组应符合 GB 30721 中一级能效要求

——风管送风式空调机组应符合 GB 37479 中一级能效要求
 ——低环境温度空气源热泵（冷水）机组应符合 GB 37480 中一级能效要求
 合规性验证方式：查验中国能效标识官网的备案信息。

5.2 电子信息设备

塔式和机架式服务器的能效应宜符合以下要求：

——塔式和机架式服务器能效应符合 GB 43630 中一级能效要求
 合规性验证方式：查验数据中心提供的相关证明材料或第三方认证证书。

5.3 配电设备

变压器、电机、泵、风机、压缩机等配电设备的能效宜符合以下要求：

——电力变压器能效应符合 GB 20052 中一级能效要求
 ——电动机能效应符合 GB 18613 中一级能效要求
 ——容积式空气压缩机能效应符合 GB 19153 中一级能效要求
 ——通风机能效应符合 GB 19761 中二级及以上能效要求
 ——潜水电泵产品能效应符合 GB 32030 中一级能效要求

注：合规性验证方式：查验中国能效标识官网的备案信息。

6 计算和测试方法

6.1 面向基础设施的数据中心能效碳效计算和测试方法

6.1.1 能效和碳效概述

本文采用电能利用效率（Power Usage Effectiveness, PUE）表征数据中心设施层的综合能源利用效率，通过总能耗电量与 IT 设备能耗量的比值进行计算；采用 CUE 表征数据中心设施层的综合碳排放效率水平，通过总碳排放量和 IT 设备能耗量评估量化，涵盖数据中心建筑单体的电力、冷却、照明等多系统能耗评估。

6.1.2 指标计算方法

a) 能效计算方法

面向基础设施的数据中心能效计算方法按照 GB 40879 的规定进行。

b) 碳效计算方法

面向基础设施的数据中心的碳效计算方法是将数据中心的总碳排放量除以 IT 设备总能源使用量。公式如下。

$$CUE = \frac{CE}{E_{IT}} \quad (1)$$

式中：

CE ：数据中心测量期内数据中心总碳排放量，单位为 kg；

E_{IT} ：IT 设备总能源使用量，单位为千瓦时（kWh），按 GB 40879 统计范围进行测试和计算。

测量期内的数据中心的总碳排放量，优先采用专业机构碳排放核算数据进行计算。在缺少核算数据时，可通过用电量的碳排放进行简化计算。计算公式为：

$$CE = E * \eta \quad (2)$$

式中：

E: 测量期内数据中心单位时间不包含绿色电力的总耗电量。

η : 中国生态环境部各区域最新电网碳排放因子, 单位是吨/MWh, 等效单位为 kg/kWh。

以下章节的碳排放简化计算方式不再赘述。

6.1.3 测试方法

面向基础设施的数据中心能效测试方法按照GB 40879的第6节规定测试方法进行。

本文规定的碳效核算方法均按照标准T/CCSA 327—2021中的第4节核算步骤和方法进行。以下章节对碳效的测试方法不再做单独说明。

6.2 面向通用计算的数据中心能效碳效计算和测试方法

6.2.1 能效和碳效概述

本文采用通用计算(以下简称: 通算)算力能效(Computing Energy Efficiency, CEE)表征通算数据中心IT基础设施(计算服务器、数据存储、网络设备)的综合能效水平; 采用通算算力碳效(Computing Carbon Usage Efficiency, CCUE)表征通算数据中心IT基础设施综合碳排放效率水平。

6.2.2 指标计算方法

a) 能效计算方法

通算数据中心综合算力能效为数据中心综合算力与用电量的比值, 即单位用电量产生的综合算力, 按公式(1)计算。

$$CEE = \frac{C}{E} = \frac{\sum k * U * P * N}{E} \quad (1)$$

式中:

CEE: 通算数据中心综合算力能效值, 单位为标准分每瓦时;

C: 通算数据中心综合算力总和, 代表了实际处理业务的能力;

E: 测量期内数据中心单位时间的总耗电量, 单位为瓦时(Wh), 按GB 40879统计范围进行测试和计算; 采用基准软件测量计算、存储、网络分别的算力值, 考虑到不同设备算力值的量级不一致问题, 不同算力值通过归一化方法解决量纲差异, 最后通过汇总计算得到数据中心综合算力能效。

算力指标考虑了通用服务器、数据存储以及网络交换三个部分的综合算力。由于数据中心业务性能复杂多变, 且业务数据有保密性要求, 导致业务算力难以准确量化评估, 所以, 本文件的算力评估方法, 采用实际负载水平与设备能力的乘积来计算, 可以最大程度表征业务算力的真实情况。

算力指标的计算公式如下:

$$C = \sum k * U * P * N \\ = \lambda * U_{\text{计算}} * P_{\text{计算}} * N_{\text{计算}} + \beta * U_{\text{存储}} * P_{\text{存储}} * N_{\text{存储}} + \gamma * U_{\text{网络}} * P_{\text{网络}} * N_{\text{网络}} \quad (2)$$

式中:

k: 服务器, 存储, 网络的归一化系数; λ 、 β 、 γ 分别为计算、存储和网络的归一化系数;

U: 实际测量期内计算、存储、网络设备的负载水平;

P: 是服务器、存储、网络设备的算力值。

N: 是服务器、存储、网络设备的数量;

不同设备的算力值差异较大, 可以采用相关设备参数、基准软件等方式获取, 具体可参考附录A。

综合算力能效的测评方法考虑到计算、存储、网络在实际应用中的比例不同, 在计算方法中增加归一化系数(λ 、 β 、 γ)来调节不同产品的权重值, 提高算力能效测评结果的真实性。具体系数参考附录B。

b) 碳效计算方法

通算数据中心综合算力碳效为数据中心总碳排放量与综合算力比值，即单位用碳排放产生的综合算力，按公式(1)计算。

$$CCUE = \frac{C}{CEE} \quad (1)$$

式中：

CCUE：数据中心综合算力碳效值，单位为每标准分的碳排放量；

C：数据中心综合算力总和，代表了实际处理业务的能力；

CE：测量期内数据中心总碳排放量，单位为kg。

6.2.3 测试方法

测试请参照附录A进行。

6.3 面向 AI 应用的数据中心能效碳效计算和测试方法

6.3.1 能效和碳效概述

本文采用面向 AEE 表征面向 AI 应用场景下数据中心的真实能效情况，即数据中心 IT 基础设备（计算服务器、AI 服务器、数据存储、网络设备）之间协同效率的综合水平评估。采用面向 ACE 表征面向 AI 应用场景下的碳排放效率水平，对数据中心实际运行下的碳排放效率进行评估。

6.3.2 指标计算方法

a) 能效计算方法

数据中心运行AI模型时的综合能效为综合性能和能耗之比，即单位能耗下的性能大小。本文采用模型浮点运算次数（Floating-Point Operations, FLOPs）对数据中心实际应用性能进行量化，如公式（1）所示。

$$PerfAEE = \frac{FLOPs}{WE} \quad (1)$$

式中：

AEE：数据中心整体能效大小，单位为浮点数运算次数每瓦时（GFLOPs/Wh），等效单位为 TFLOPs/kWh等；

FLOPs：测试期间，被测服务器对象进行的实际浮点数计算次数，单位为GFLOPs，等效单位为 TFLOPs等；

WE：模型执行期间被测服务器对象的能耗值，单位为瓦时（Wh），等效单位千瓦时（kWh）等。

实测中可针对数据中心应用领域及特点，被测者自行选择执行阶段，即训练或推理，或根据数据中心任务阶段占比情况进行两阶段综合测试。用公式（2）对算力能效进行综合计算：

$$AEE_{cluster} = \omega_{train} * PerfAEE_{train} + \omega_{inference} * PerfAEE_{inference} \quad (2)$$

式中：

ω_{train} ：被测对象进行训练的能效权重，通常取0、0.5、1；

$\omega_{inference}$ ：被测对象进行推理的能效权重，通常取0、0.5、1，使得训练和推理权重相加和为1；

AEE_{train} ：被测对象训练阶段数据中心能效大小；

$AEE_{inference}$ ：被测对象推理阶段数据中心能效大小。

$Perf_{train}$ $Perf_{inference}$ 进一步， $AEE_{train}Perf_{train}$ 的详细定义如下，

$AEE_{inference} Perf_{inference}$ 定义类似，不赘述。

$$PreAEE_{train} = \sum_i^n \omega_i \frac{FLOPs_i}{EW_i} \quad (3)$$

式中：

ω_i n ：为参与测试的模型任务个数；

ω_i ：为第 i 个测试模型任务的权重系数；

$FLOPs_i$ ：为第 i 个测试模型任务执行中的 FLOPs，通常指精度为 FP16 下的浮点数计算个数，若精度不一致，则进行换算；

EW_i ：为第 i 个测试模型任务执行期间所消耗的能量，单位为瓦时（Wh），千瓦时（kWh）。

b) 碳效计算方法

面向AI应用场景的算力碳效为数据中心总碳排放量与AI算力比值，即单位算力产生的碳排放量，按如下公式计算。

$$C_{effACE} = \frac{F_{total}FLOPs}{CEC} = \frac{\sum_i^n FLOPs_i}{E * \eta} \quad (1)$$

式中：

C_{effACE} ：一次浮点数计算产生的碳排放量；

$F_{total}FLOPs$ ：定义同（1）能效计算方法；

CE ：被测对象在测试期间的总碳排放量， gCO_2 单位为 kg。

E 和 η 定义同 6.1.2 节中碳效计算方法的式（2）定义。

6.3.3 测试方法

6.3.3.1 测试要求

测试环境要求及测量设备按照 GB 40879-2021 中第 6.1 节规定执行。

本文规定测试的开始时间：对于训练任务，启动模型执行指令作为开始时间，此时数据在服务器存储设备中，未加载至内存；对于推理任务为请求方发起模型请求作为开始时间。

本文规定测试的结束时间：对于训练任务以模型完成最后一轮次的训练并且保存完毕最终的模型参数作为结束时间；对于推理任务为请求方收到完整的返回结果作为结束时间。

6.3.3.2 能耗测量

测试期间功耗单位为瓦（W）。

被测对象为物理机，测量位置参考 GB 40879—2021 中电能消耗测量点示意图中的 8 点；被测对象为虚拟机，则应支持输出服务器的即时功率。

服务器可支持直接获取任务执行期间的总能耗量，或总能耗量采用 AI 服务器单机执行期间平均功率进行估算。具体在模型执行期间，周期性采样测量服务器的负载功率，求均值。模型训练中采样周期至少

横跨一批次样本的完整计算过程，在推理中采样周期分布在整個推理期间，且采样密度均不小于 10%。

6.3.3.3 性能测量

基于模型的组成结构、应用领域、参数情况、代表性、开源情况以及可复现度等多因素推选出本标准中使用的评估模型工具，基于相关模型对数据中心算力能效进行统一的评估。其中用于训练或推理测试的模型和数据集工具信息见附录 D，具体模型推荐结构和 FLOPs 参考取值及计算方式示例见附录 E。

面向 AI 应用数据中心训练或推理测试过程应符合以下规定：

1) 训练测试步骤

1.准备服务器基础环境：安装硬件环境的驱动、固件、**docker** 等基础软件配置。若硬件已有配套软件，该步骤不考虑。

2.测试准备：

- a) 准备模型和数据集：选定模型和数据集，下载数据集，编写代码或下载开源代码；
- b) 如需要，进行数据集相关格式转换和预处理等；
- c) 构建模型运行环境：安装模型依赖环境，或拉取镜像创建启动容器。

3.运行测试：

a) 启动测试：基于模型训练入口启动执行模型，期间记录过程日志、监控 loss 下降情况、开始结束时间、能耗采样信息等，测试期间模型参数设置按照附录 C 中要求进行；

b) 结果验证：对测试结果进行检查，是否达到预期目标。若不达标，则进行调整重新按流程测试。

4.结果回收：返回测试信息，基于模型和超参数、能耗等信息计算能效值。

5.输出结果：至少进行 5 次重复测量，求多次测量平均值作为最终的测试结果。

2) 推理测试步骤

1.准备服务器基础环境：安装硬件环境的驱动、固件、**docker** 等基础软件配置。若硬件已有配套软件，该步骤不考虑。

2.测试准备：

- a) 准备模型和数据集：选定模型和数据集，下载数据集，编写代码或下载开源代码；
- b) 如需要，进行数据集相关格式转换和预处理等；
- c) 构建模型运行环境：安装模型依赖环境，或拉取镜像创建启动容器。

3.运行测试：

a) 启动测试：基于模型推理入口启动执行模型，期间记录过程日志、吞吐率、开始结束时间、能耗采样信息等，测试期间模型参数设置按照附录 C 中要求进行；

b) 结果验证：对测试结果进行检查，是否达到预期目标。若不达标，则进行调整重新按流程测试。

4.结果回收：返回测试信息，基于模型和超参数、能耗等信息计算能效值。

5.输出结果：至少进行 5 次重复测量，求多次测量平均值作为最终的测试结果。

附录 A (规范性) 通算数据中心测试方法

A.1 电功耗/电能测量

通算数据中心设备总耗电量 E 应按照 GB 40879—2021《数据中心能效限定值及能效等级》规定的测量条件和测量方式进行测量，得到时间段 t 单位时间对应的 E 。

A.2 服务器能效 SEE 测量

通用计算服务器的SEE，采用相关测试软件进行测量（如BenchSEE或Spec）。可通过现场搭建环境测试，或由厂家提供正式测试报告，或厂家官网提供对外规格。

A.3 数据存储能效 DEE 测量

DEE 参考采用能效基准工具进行测量（如 BenchDEE），可以通过现场搭建环境测试，或者由厂家提供测试报告，或者厂家官网提供的对外规格。

A.4 网络能效 NEE 测量

NEE根据T/CCSA 370—2022《数据中心网络能效比测量规范》测量EER值，NEE是EER值的倒数，可通过现场搭建环境测试，或由厂家提供测试报告，或厂家官网提供的对外规格。由于网络占比较低，此项为可选。

A.5 数据中心信息设备负载水平和综合算力测量

数据中心信息设备负载水平 U 有两种计算方法。方法一按照设备类型分组计算，比如数据中心有三种服务器，每组服务器数量为 A ， B ， C 。分别计算 U_A ， U_B ， U_C ，代入公式计算。方法二所有设备负载水平取平均值，代入公式计算。

数据中心信息设备能效计算方法应按照单台设备的能效乘以设备数量来计算总算力。比如数据中心有三种服务器，每组服务器数量为 A ， B ， C 。服务器对应能效为 x ， y ， z 。总能效为 $A * x + B * y + C * z$ 。

通用服务器的 $U_{计算}$ 是实际占用 CPU 资源的百分比，在测试周期内持续采集 CPU 的占用率数据，取算数平均值代入公式中计算。 $U_{存储}$ 是数据存储的实际占用 CPU 资源的百分比，从存储管理软件中可以读取。在测试周期内持续采集存储产品占用 CPU 资源的百分比，取算数平均值代入公式中计算。 $U_{网络}$ 是网络设备的实际出口交换机上下行总带宽与可用总带宽的比值。通过开源软件或商用软件监测。监测数据的间隔建议为 60 秒，最长不超过 15 分钟，计算运行综合算力能效的监测周期宜采用 24 小时统计值，最短不少于 12 小时。年度综合算力能效应采用 8760 小时统计平均值。

附录 B
(资料性)
归一化系数设置方法

B.1 数据中心综合算力能效归一化系数

考虑到计算、存储、网络量纲不一致，需要进行归一化处理，归一化与权重统一用下表表示。数据中心算力能效归一化系数依据目前市场的通用情况统计取值，具体如下表：计算设备的系数约为0.85，数据存储设备的系数约为0.1，网络设备的系数约为0.05，使用可以按照实际情况进行调整。

表B.1 数据中心综合算力能效归一化系数表

λ	β	γ
0.85	0.1	0.05

被测者也可以根据实际测试数据来调整服务器、存储、网络，三个算力的系数。如果采用其它性能测试软件，归一化值需要重新调整。

附录 C
(资料性)
模型相关影响因素说明

C.1 模型实际运行参数

测试模型执行中对计算量和计算效率有影响的相关信息整理，以及相关要求。

表 C.1 模型实际运行相关参数说明

维度	描述	举例	要求
运行阶段	模型运行的不同状态	训练、推理	-
评估方式	模型是否可更改	开放、封闭	对标封闭，可修改不做限定的内容
框架	框架的性能与优化、图优化与自动微分	TensorFlow、PyTorch、MXNet	不做限定
并行策略	硬件并行优化	多GPU训练、TPU Pod	不做限定
	并行策略	数据并行、模型并行、混合并行、混合精度训练	限定
计算优化策略	优化算法和内存布局，在软件层面提升性能	flash attention 2	限定
超参数	训练前设定	网络结构超参数：层数、每层单元数、激活函数、输入输出维度、注意力头数； 训练过程超参数：批次大小、训练轮数	不做限定，推荐使用本标准推荐参数
精度	数据存储的精度	FP32、FP16、INT8	不做限定，根据实际评估时进行换算
训练相关信息	优化器	Adam、SGD、Lamb、RMSProp。影响训练的稳定性 and 速度	限定
	学习率调度	学习率调度策略帮助模型在训练过程中更快地收敛	限定
	正则化技术	Dropout、L2 正则化、早停（early stopping）等正则化方法可以防止模型过拟合，保证模型的泛化能力	限定
	是否采用重计算	为了节约内存，在反向传播的时候重新进行正向计算	限定
	归一化：批归一化、层归一化	通过规范化每一层的输入，减小权重初始化和学习率设置的敏感度，加速训练过程	限定，增加了少量的额外 FLOPs 计算量，推荐与模型设定保持一致
训练数据的质量与数量	数据预处理与增强	对数据集进行固定，此处不做修改	限定，必须使用推荐通用数据集和数据量大小

附录 D
(资料性)
推荐测试模型信息

D.1 推荐测试模型

用于性能测试的模型领域、目标值和数据集等信息，包括训练和推理两个阶段，分别见表 D.1 和 D.2。

表 D.1 用于训练测试的模型信息

领域	子方向	模型	目标	数据集
计算机视觉 (CV)	图像分类	resnet-50	ACC: 0.759	Imagenet
	目标检测	Retinanet	34.0% mAP	OpenImages
自然语言处理 (NLP)	自然语言编码	bert-large	ACC: 0.72	Wikipedia 2020/01/01
	大模型	llama2-70b (LLama2 70B-LoRA finetuning)	0.925 cross entropy loss	SCROLLS GovReport
		llama2-7b	交叉熵损失函数 5.0 左右	openwebtext, 如无特殊说明, 训练前 1 亿个 token
商业应用 (CA)	推荐系统	DLRM-dcnv2	AUC: 80.31%	Multihot Criteo Terabyte

表 D.2 用于推理测试的模型信息

领域	子方向	模型	推理目标	数据集
计算机视觉 (CV)	图像分类	resnet-50	ACC: 76.0%	imagenet2012 validation
	目标检测	Retinanet	mAP: 37.5%	OpenImages mlperf validation
自然语言处理 (NLP)	自然语言编码	bert-large	F1: 90.0%	SQuAD v1.1 validation set
	大模型	llama2-70b (LLama2 70B-LoRA finetuning)	Rouge1: 45.0% Rouge2: 22.0% RougeN: 28.0%	1M-GPT4-Augmented.parquet
		llama2-7b	ACC: 46.0%	MMLU
商业应用 (CA)	推荐系统	DLRM-dcnv2	ACC: 81.0%	Multihot Criteo

附录 E
(资料性)

AI 模型超参数和 FLOPs 计算值示例

E.1 测试模型的超参数和对应FLOPs值参考

以下整理了用于数据中心性能测试的模型超参数和 FLOPs 信息，如表 E.1 所示。其中罗列的 FLOPs 为批次大小为 1 且仅进行一次正向传播的模型理论计算量，通常模型的一次迭代的计算量主要包含正向传播、反向传播和参数更新，因当前模型中参数更新涉及计算量占比较小，此处忽略不计。且通常认为反向传播的计算量是正向传播的 2 倍。

表 E. 1 本文中用于测试的模型超参数和 FLOPs 参考

领域	子方向	模型	超参数信息			一次正向传播的计算量 (FLOPs/样本)
计算机视觉 (CV)	图像分类	resnet-50	参数	参数说明	经典取值	3.8×10^9
			$I_h \times I_w$	输入尺寸	$224 \times 224 \times 3$ (RGB 图像)	
			$K_h \times K_w$	卷积核大小	例如: 7×7 , $K_h = K_w$	
			stride	步幅	例如: 2	
			C_{in}	输入通道数	第一层的输入通道数为 3, 之后的为上一层的输出通道数	
			C_{out}	输出通道数	例如: 64	
			$H_{out} \times W_{out}$	输出尺寸	$H_{out} = W_{out} = (I_h - K_h) / \text{stride} + 1$, 例如: $(224 - 7) / 2 + 1$	
			batch_size	批次大小	实际设定而定	
			epoch	训练轮次	仅对训练而言	
	目标检测	Retinanet	参数	参数说明	经典取值	2.53×10^{10}
			batch_size	训练批大小	32	
			H_out 和 W_out	输入图像的尺寸	Image Size = 800 x 800	
			C_in 和 C_out	分别是输入和输出通道数	视代码而定	

			K _h , K _w	K _h 和 K _w 是卷积核的大小	视代码而定	
			N _{class}	类别数	视代码而定	
			L	卷积层数	50	
自然语言处理 (NLP)	自然语言编码	bert-large	参数	参数说明	备注	3.1×10^{11}
			N	Transformer编码器层数	Bert: 12层; Bert-large: 24层	
			d_{model}	隐藏层维度	Bert: 768; Bert-large: 1024	
			d_{ff}	前馈层维度	Bert: 3072; Bert-large: 4096	
			h	注意力头	Bert: 12; Bert-large: 16	
			L	序列长度	Bert: 通常可取128, 256, 512等, 此处取512; Bert-large: 512	
			V	词汇表大小	此处暂设V=32000	
	大模型	llama2-70b (LLama2 70B-LoRA finetuning)	参数	参数说明	备注	70b 的参数 7.5×10^{14} lora 微调参数 1.5×10^{12}
			batch_size (b)	训练批大小	Llama2-7b: - Llama2-70b: 4M (代码中推荐)	
			sequence_length (s)	上下文长度	Llama2-7b: 2048 Llama2-70b: 4096	
			hidden_size (h)	隐藏层大小	Llama2-7b: 4096 Llama2-70b: 8192	
		llama2-7b	i	全连接网络中的中间隐藏层大小	通常为 4*h Llama2-7b: 11008 Llama2-70b: 4*h	2.9×10^{13}
			n	注意力头数量	Llama2-7b: 32 Llama2-70b: 64	
			d	每个头的维度	Llama2-7b: - Llama2-70b: -	
			L	模型层数	Llama2-7b: 32 Llama2-70b: 80	
		DLRM-dcnv2	参数	参数说明	备注	$5^{13} \times 10^{10}$
			embedding_dim	Embedding 维度	128	
			dense_arclayer_size	全连接层架构	[512, 256, 128]	
商业应用 (CA)	推荐系统	DLRM-dcnv2				$5^{13} \times 10^{10}$

			zes		
			over_arch _layer_siz es	特征交互层	[1024, 1024, 512, 256, 1]
			dcn_num_ layers	特征交互层的层数	3
			dcn_low_r ank_dim	每个交互层中的低 秩矩阵维度	512
			num_emb eddings_p er_feature	每个特征嵌入维度	40000000,...12973,108,3 6
			batch_size	训练批大小	65536

E.2 模型基本结构的 FLOPs 计算方式

因考虑到模型的组成基本结构有相似之处，为方便理解，此处仅对基本模型结构的 FLOPs 计算方式进行简要计算说明。实际操作中，可基于已有 FLOPs 计算工具（例如：ptflops、thop）进行快速计算。

卷积层的 FLOPs

$$FLOPs_{conv} = 2 * H_{out} * W_{out} * C_{in} * K_H * K_W * C_{out} * batch_size$$

其中，系数 2 的含义为一次矩阵运算会进行两次浮点数运算，一次乘法和一次加法运算。 H_{out} ， W_{out} 分别为输出尺寸，且 $H_{out} = W_{out} = (I_h - K_H) / stride + 1$ ， I_h 为输入尺寸，stride 为步幅， C_{in} 和 C_{out} 分别为输入输出通道数， K_H 和 K_W 为卷积核大小，batch_size 为批次大小。

全连接层的 FLOPs

$$FLOPs_{fc} = 2 * N * M$$

其中，N 为输入的特征数（上一层的输出大小，通常为卷积层输出展平后的结果），M 为输出的特征数（当前层的神经元数量）。

Transformer 结构的 FLOPs

1) 输入嵌入层

$$FLOPs_{embedding} = 2 * L * d_{model} * V$$

其中，L 为序列长度， d_{model} 为隐藏层维度，V 为词汇表大小。

2) 自注意力层

$$FLOPs_{QKV,total} = 2 * L * d_{model}^2 * h$$

$$FLOPs_{attention} = 2 * L^2 * d_{model}$$

$$FLOPs_{self_attention} = FLOPs_{QKV,total} + FLOPs_{attention}$$

其中，h 为注意力头的个数。 $FLOPs_{self_attention}$ 为一个注意力结构的 FLOPs 值大小。