



通算超节点性能基准工具

ClusterBench 研制进展及应用场景解读

陈海

OAI社区Benchmark项目群组长

中国电子技术标准化研究院信息技术研究中心硬件研究室主任

通算超节点性能基准工具

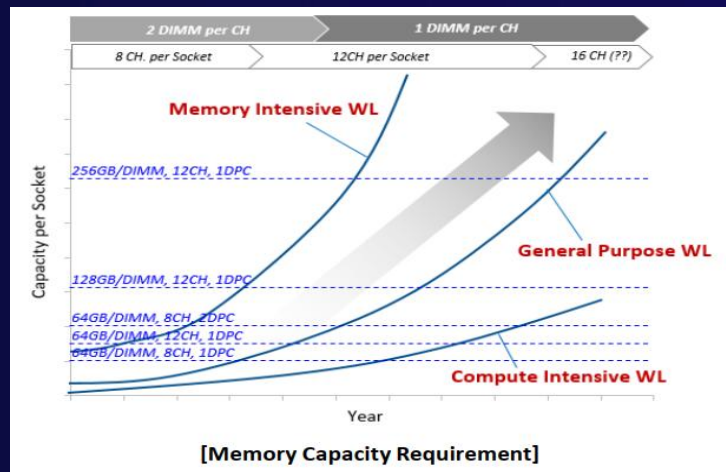
ClusterBench 研制进展及应用场景解读

陈海

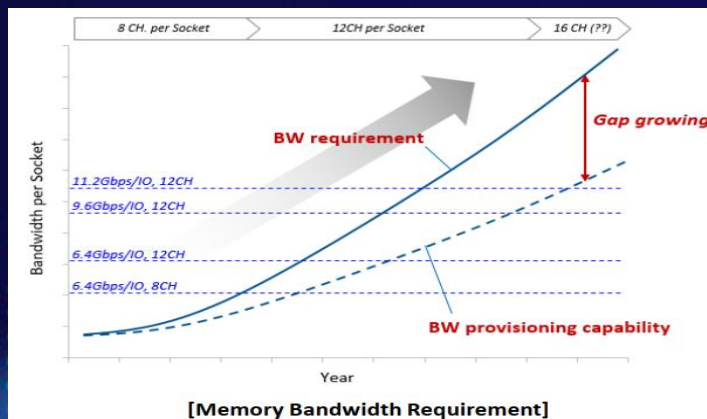
OAI社区Benchmark项目群组长
中国电子技术标准化研究院信息技术研究中心硬件研究室主任

云数据中心面临挑战：内存带宽与容量成为制约算力持续提升的最大瓶颈，内存利用率低与TCO高问题并存

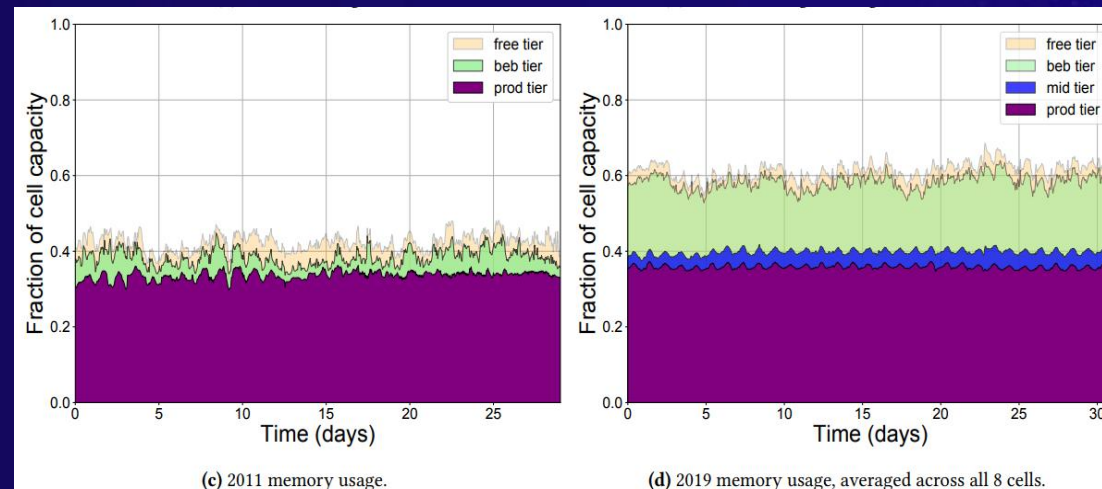
内存容量增长难以满足内存bound负载需求



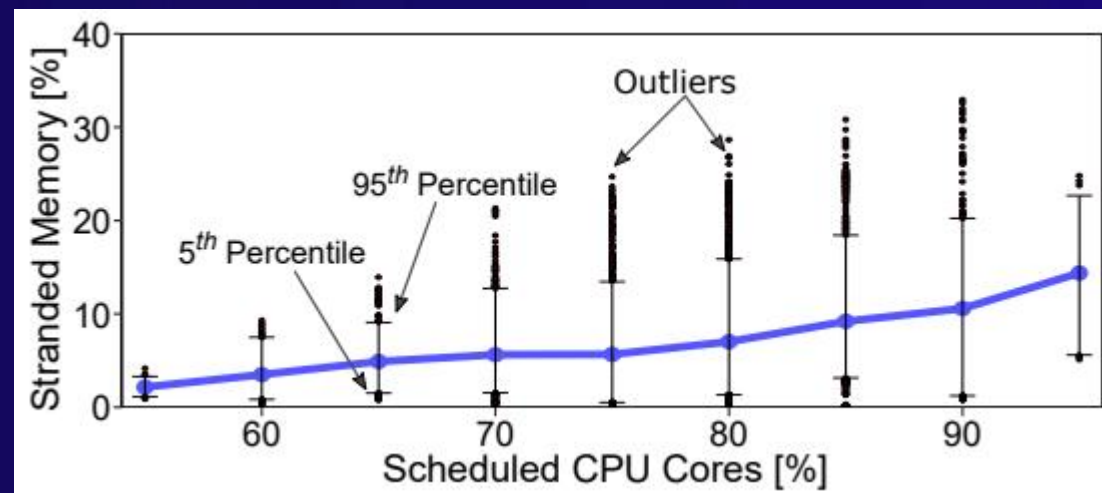
内存带宽与实际需求存在剪刀差



Google: 40%以上的内存未被使用



Microsoft: 高达25%的内存无法分配出去

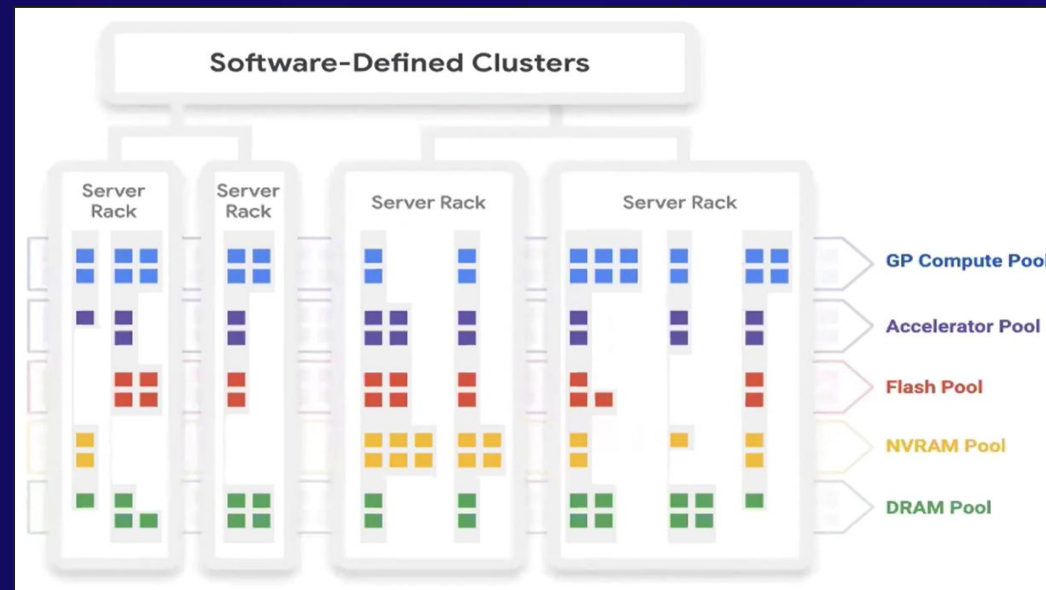
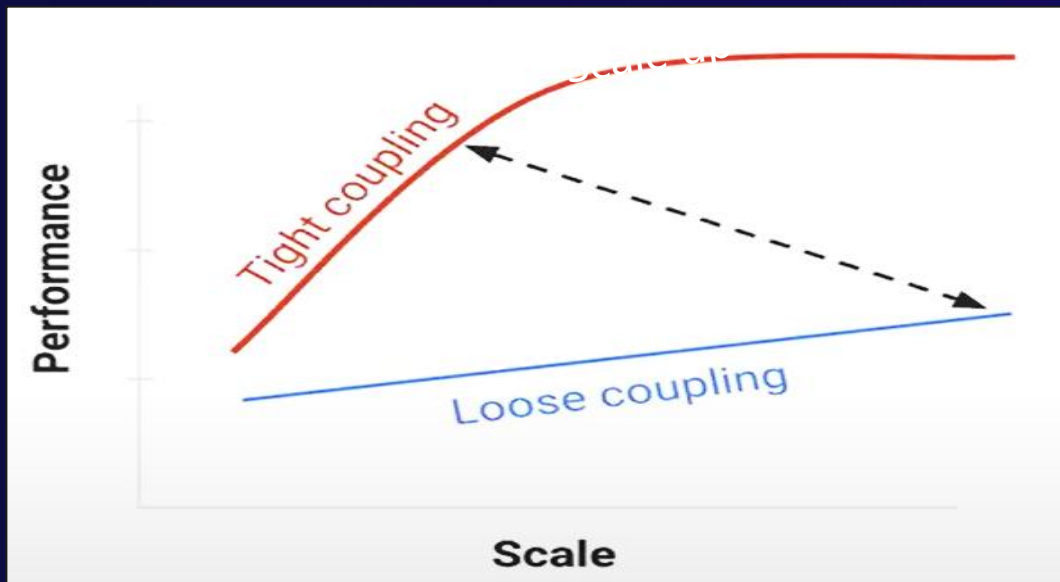


计算网络、存储之间的边界将继续模糊，紧耦合的集群系统逐渐成为产业趋势

互联网等下一代基础设施发展方向核心：分布式紧耦合，DSA，极致资源编排，高效计算系统

□ 性能：紧耦合分布式系统更容易获得Scale up的性能

□ 资源：消除4大组件之间静态配比，更容易满足应用负载的多样性



“Scale out” to “Scale up”：

- 传统松耦合思路导致器件更大规模和更低性能获得
- 高带宽、低时延和远程内存语义形成紧耦合分布式系统，在一定规模下具备更高性能

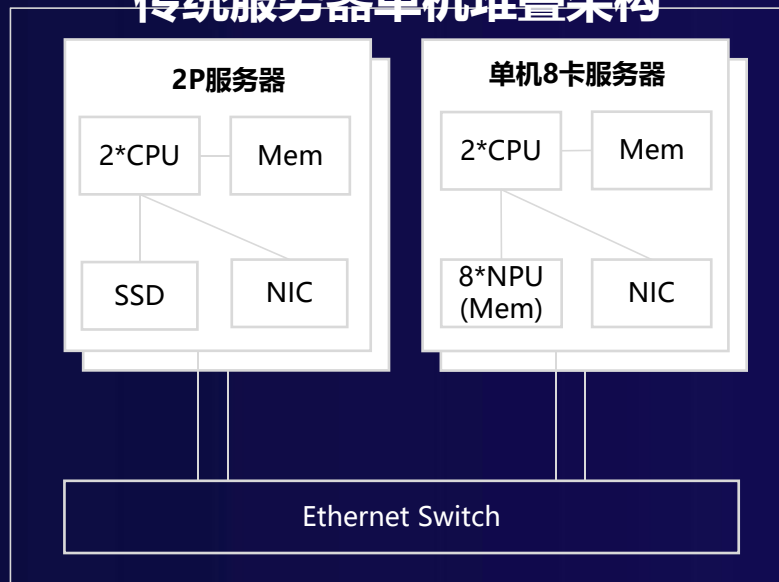
软件定义基础设施：

- 传统分布式以服务器节点为基础，带来负载不均衡、资源碎片化、算力与IO不协调带来的冗余浪费，是影响分布式系统的关键
- 内存成本占比>50%，解决内存池化超分是TOP诉求

通用计算场景面对芯片性价比挑战，需通过计算架构创新提升系统性价比

从传统服务器单机堆叠到紧耦合架构演进，需要新的量化评估基准牵引系统设计优化

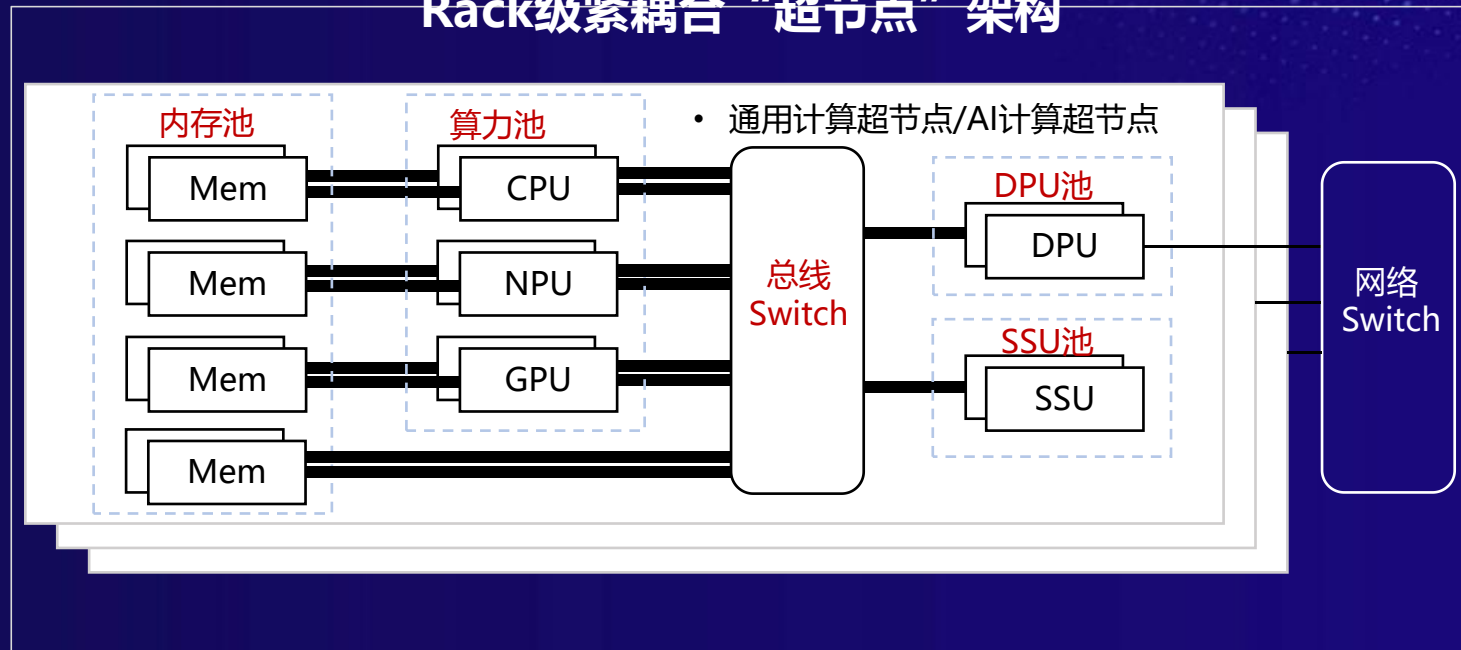
传统服务器单机堆叠架构



传统性能基准关键特征：

- 1、**测试对象**：部件、单机及单场景为主，计算、存储、网络分开测试
- 2、**测试模式**：100%满载压力测试为主
- 3、**测试负载**：核心算子、proxy 应用为主
- 4、**测试指标**：规格算力为主

Rack级紧耦合“超节点”架构



新架构性能基准关键挑战：

- 1、**测试对象**：集群系统级为主，如何横向对比？
- 2、**测试模式**：资源配置动态伸缩如何模拟？
- 3、**测试负载**：计算、内存、IO等密集型负载，动态压力负载输入如何设计？
- 4、**测试指标**：较传统分布式系统，性价比、资源利用率指标如何统计？

计算系统架构演进带来评测变化：从传统分层解耦模式服务器单机算力对比，演进到可体现解决方案全栈能力的集群有效算力对比

DCPerf: Meta构建Benchmark来支撑核心业务CPU采购，解决传统基准无法准确表示复杂DC负载痛点

- 针对现有性能测试套（如SPEC CPU、CloudSuite）无法模拟数据中心复杂业务、系统级特征表现相似但微架构级特征表现不相似的2大挑战，提出与数据中心应用具备相似软件栈、系统层级及微架构层级特征的DCPerf性能测试套，在4代多核服务器（核数36到176）上验证，误差仅为3.3%

关键技术一：生产级高保真负载模拟

解决问题1：传统基准仅模拟功能性，忽略软件栈、CPU扩展性以及“数据中心税”，无法正确模拟生产级负载；

DCPerf Automation Framework

Representative Applications

TaoBench

FeedSim

DjangoBench

Mediawiki

SparkBench

VideoTranscode

Hooks

Perf

CPU Util

MemStat

NetStat

CPUFreq

Power

uArch

Topdown

Microbenchmarks for Datacenter Taxes

Folly

FBThrift

zstd

.....

Workload	Web	Ranking	Data Caching	Big Data	Media Processing
Benchmarks in DCPerf	MediaWiki DjangoBench	FeedSim	TaoBench	SparkBench	VideoTranscode
Performance metrics	Peak RPS	RPS under latency SLO	Peak RPS and cache hit rate	Throughput	Throughput
Req. proc. time	Seconds	Seconds	Milliseconds	Minutes	Minutes
Peak CPU util.	90-100%	50-70%	80%	60-80%	95-100%
Thread-to-core ratio	N(100)	N(10)	N(10)	N(1)	N(1)
Per-server RPS	N(1K)	N(100)	N(1M)	N(10)	N(10)
RPC fanout	N(100)	N(10)	N(10)	N(10)	0
Instructions per request	N(1B)	N(10B)	N(1K)	N(10B)	N(1M)

软件架构；DCPerf与工作负载的映射关系，包括各基准的性能指标（如RPS、吞吐量）、CPU 利用率及线程模型，直观体现多样化建模思路。

关键技术二：微观指标对齐

解决问题2：传统基准仅关注宏观性能（如吞吐量），忽略底层硬件特征（如指令周期、缓存命中率），导致代际验证预测性能表现存在偏差。

IPC Per Physical Core

Benchmark	Prod & DCPerf	SPEC2017
Cache (prod)	1.2	1.1
TaoBench	1.8	1.8
Ranking (prod)	1.0	1.4
FeedSim	1.2	1.4
IG Web (prod)	2.6	2.6
DjangoBench	2.0	1.6
FB Web (prod)	0.6	0.8
Mediawiki	1.5	3.3
SparkBench	2.1	1.9
Spark (prod)	2.5	1.8
500.perlbench	3.3	2.5
502.gcc	1.5	1.8
505.mcf	1.5	1.8
520.omnibench	1.5	1.8
523.xalan	1.5	1.8
525.x264	1.5	1.8
531.deepjv	1.5	1.8
541.leela	1.5	1.8
548.exchange2	1.5	1.8
557.xz	1.5	1.8

Percentage of Pipeline Stalls (%)

Category	Prod	DCPerf	SPEC2017
Frontend stalls	36	34	20
Backend stalls	16	13	24
Retiring	39	45	47

L1-Cache Miss (MPKI)

Benchmark	Prod & DCPerf	SPEC2017
Cache (prod)	56	54
TaoBench	17	14
Ranking (prod)	55	46
FeedSim	39	31
IG Web (prod)	7	12
DjangoBench	3	9
FB Web (prod)	2	4
Mediawiki	4	4
SparkBench	1	1
Spark (prod)	1	1
500.perlbench	1	1
502.gcc	1	1
505.mcf	1	1
520.omnibench	1	1
523.xalan	1	1
525.x264	1	1
531.deepjv	1	1
541.leela	1	1
548.exchange2	1	1
557.xz	1	1

Memory Bandwidth (GB/s)

Benchmark	Prod & DCPerf	SPEC2017
Cache (prod)	29	17
TaoBench	31	30
Ranking (prod)	19	21
FeedSim	36	29
IG Web (prod)	36	33
DjangoBench	36	33
FB Web (prod)	36	33
Mediawiki	36	33
SparkBench	36	33
Spark (prod)	36	33
500.perlbench	16	43
502.gcc	43	68
505.mcf	50	50
520.omnibench	18	18
523.xalan	5	5
525.x264	8	8
531.deepjv	3	3
541.leela	3	3
548.exchange2	3	3
557.xz	21	21

对比 DCPerf、生产负载及 SPEC 的 IPC、内存带宽、L1 I-Cache miss 率

通算超节点ClusterBench定位：覆盖通用计算集群关键要素，支撑通算超节点及多节点等性能评估

- **ClusterBench** 是一款面向整机柜/计算集群系统的性能评测工具，主打通算超节点、池化系统、多节点/整机柜产品性能评估测试，辅助解决计算集群性能评估系统性、准确性难题
- **用户**：云数据中心、互联网、运营商、芯片厂商、服务器厂家等
- **使用场景**：售前测试、业务部署配置优化、采购规模评估、系统性能瓶颈分析

测试对象

评估整个集群软硬件综合业务性能：

- **服务器**：CPU、内存、硬盘、网卡、DPU等
- **交换机**
- **OS**：编译器、调度器、运行时、加速库等

测试节点：1~N节点

测试对象：CXL-based系统/UB-based系统/以太扩展的系统（arm+x86+……）

业务场景



基础能力

多维度评测模型

整体体现为：

整机柜/集群综合性能（即有效算力）

应用性能：各APP场景业务性能

可细分为：

扩展性：应用扩展性、集群扩展性

服务质量Qos/可靠性：90%响应时间

计算能力：整型、浮点（双精度、单精度）

访存能力：带宽、时延

IO能力：网络带宽/延时、存储吞吐

能效：待定

通算超节点性能基准ClusterBench 整体技术方案

ClusterBench技术点总览



关键技术

□ 综合指标分层设计（满足从基础到业务性能观测评估需求）：

1. **L0-L2综合指标体系**：确定指标设计原则，基于关键技术指标、客户需求指标构建指标备选池，各指标分别设计统计方案；
2. **性价比指标**：性能由各负载的吞吐量标准化后几何平均得出，成本由关键资源(CPU+内存)规格基于单位资源得出，性价比=性能分/资源分；

□ 负载套件（负载真实）：

1. **业务代表性**：基于TOP算力场景/互联网客户场景/未来趋势场景构建场景库，抽取典型用例，提取用例核心业务流，通过多个用例特征进行核密度分析，选取TOPN密度作为备选用例；
2. **特征多样性**：对备选用例的特征经过数据预处理、降维后，进行相似度分析，若两用例特征相近，则通过特征空间投影进一步分析各维度的相似度和分布密度，确保桶散点图中，特征点均匀分布，总体特征向量正交；

□ 波峰波谷的压测模式（模拟客户业务潮汐特征）：

1. **负载压测引擎**：使用半开环压测模型，通过压测线程数控制到达率，各workload通过参考机器平均时延确定压测线程数；
2. **波峰波谷测试流**：三个阶段各执行30min，波谷控制到达率为~30%，波峰控制在~90%，burst有多种模式允许配置，默认采用基于Tile的burst策略；

□ 高易用高可靠的测试框架，对接云管平台（易用性）：

1. **一键式安装部署和运行**：抽象被测系统适配层，对不同的SUT各实现子步骤，框架通过统一流程执行 依赖安装->节点部署->服务启动；
2. **对接云管平台**：用户提供云管平台服务，框架对接实现控制，对K8S和OpenStack抽象统一控制逻辑实现负载层无感知；

测试框架：支持一键自动化部署、测试，支持虚机、容器部署负载、部署时间、测试时间、兼容不同的集群部署

实现方案

支持集群基准测试：

1. **节点控制**：PAgent负责创建管理物理机，VAgent负责管理虚拟机和容器，和节点的关系为一对一，Controller和Agent通过grpc通信和数据文件传输；
2. **运行监控**：Agent定时(默认5s)采集资源使用数据，通过心跳上报Controller，Controller汇总形成报表供查看；

支持虚机、容器部署负载

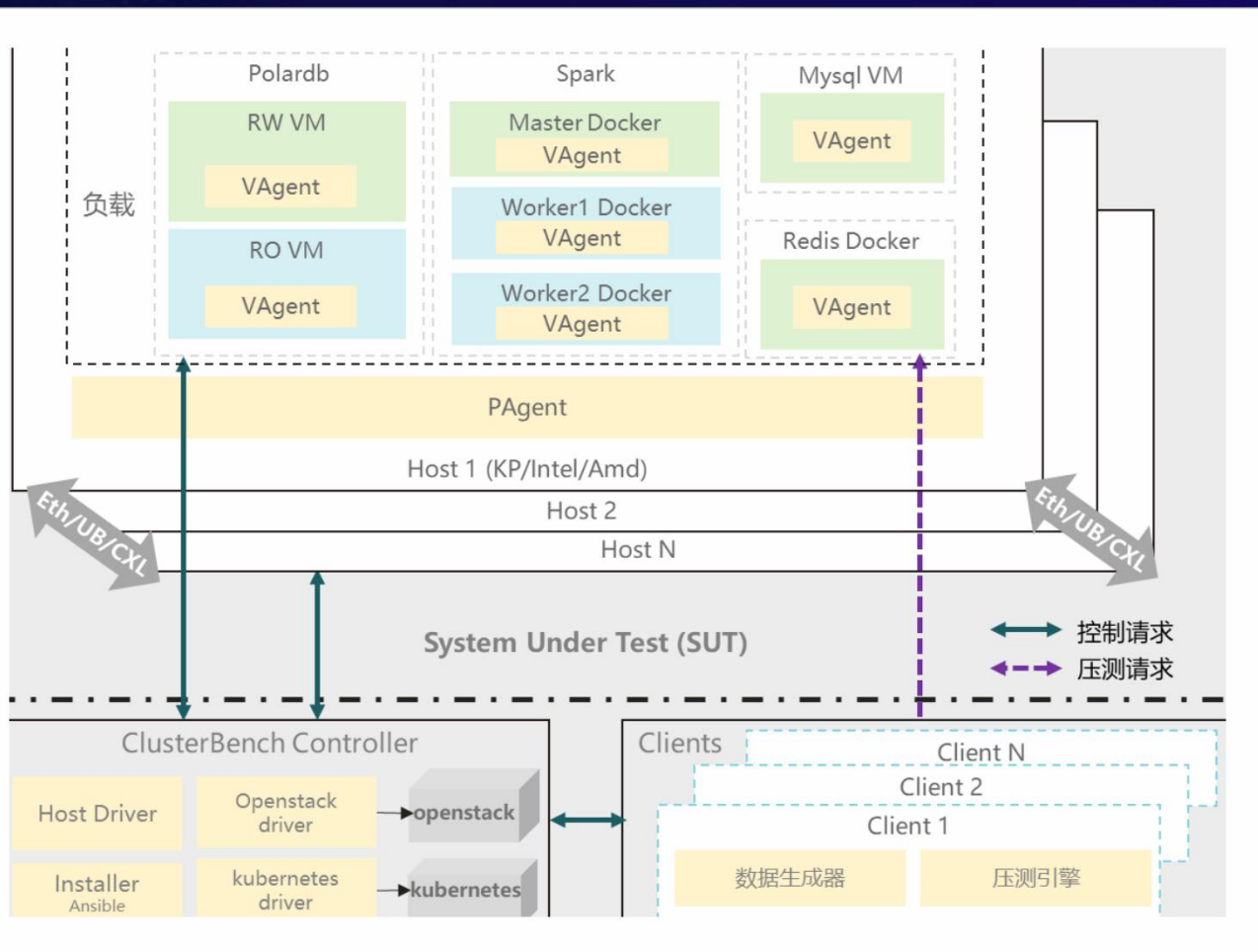
1. **云管平台控制**：用户提供云管平台，框架对接平台实现创建虚拟机容器、超分、调度、资源监控等控制，对K8S和OpenStack抽象统一控制逻辑实现负载层无感知；
2. **超分控制**：允许用户配置超分比，默认为CPU 1.5，内存1.2，基于超分比调用云管接口限制节点资源分配；

一键式安装部署和运行：

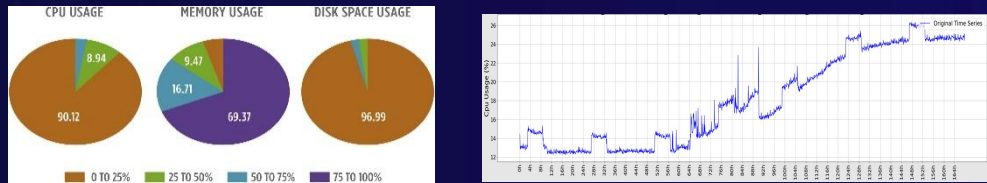
1. **兼容性**：抽象SUT适配层，对不同的SUT各实现子步骤，框架通过统一流程执行 依赖安装->节点部署->服务启动；
2. **自动化部署**：通过Ansible实现部署流程配置化，实现多节点并发部署，将部署时间缩短在30min内，并提供详细的错误日志；

灵活的测试方法：

1. **类sysbench模式**：提供测试数据生成，压测模型。测试用户实际的业务应用；
2. **类spec virt模式**：提供镜像，测试数据，压测模型，预置负载，模拟业务调度场景

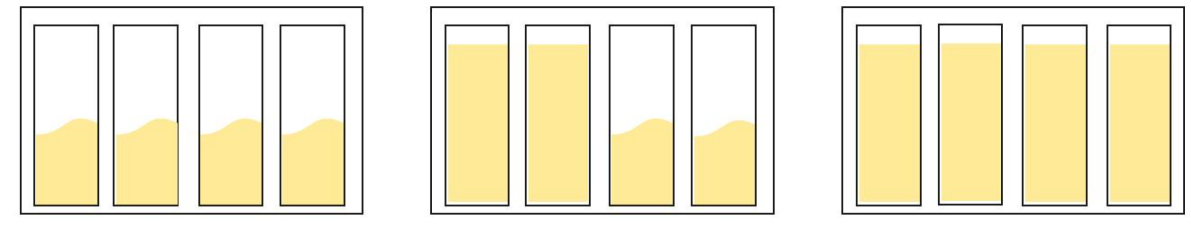


测试模式：三阶段模拟业务波峰波谷不同时段负载压力变化，更贴合DC实际业务特征



参考公有云ESC弹性计算服务，实际业务中资源使用率较低<30%，存在负载潮汐

压测流程



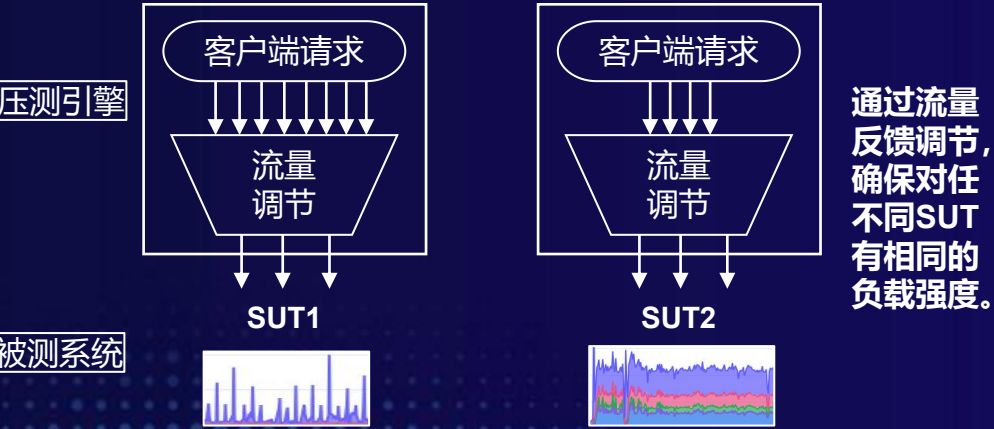
波谷阶段：
以Tile为单位启动实例，压力强度控制在10%-30%（可配置）

Burst阶段：
模拟负载潮汐：部分Tile压力强度骤升为高负载，其他保持轻负载

波峰阶段：
模拟负载高峰期：全部Tile均为满额负载，达到性能瓶颈点~90%

压测模型设计

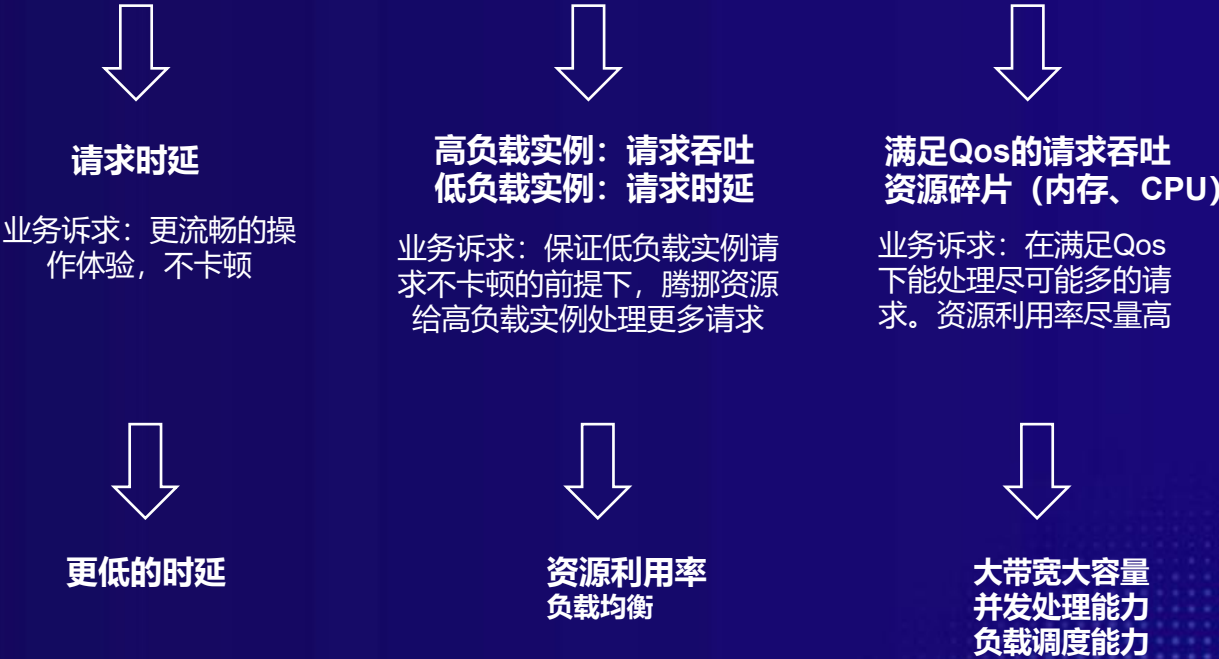
- **负载强度定义：**单位时间内被测系统的请求到达率。
- **压测模型：**同一负载水平下，对所有SUT保持负载强度一致。



- 确定不同负载强度所需的压测线程数

性能评估指标

牵引集群设计方向

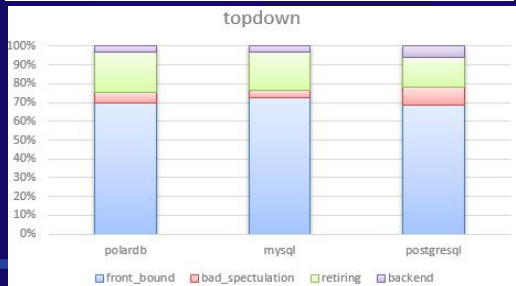
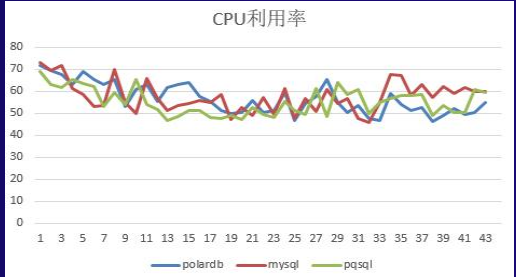
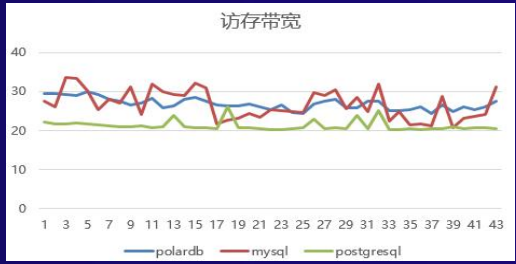


负载选型：围绕互联网+国计民生梳理6大解决方案典型负载，提取负载典型业务模型

领域	典型负载		关注指标
大数据	Spark TPCDS		SQL时延
	Hbase YCSB		吞吐
	Flink nexmark/HiBench		吞吐
数据库	MySQL/PG/PolarDB TPCC		吞吐、时延
	clickhouse/doris/openGauss TPCH		吞吐、时延
	Redis等内存数据库		吞吐
	milvus/VectorDB/openGauss等向量数据库		吞吐、时延
虚拟化（含虚机/容器等）	Kafka		吞吐
	Nginx		IOPS
	Mysql sysbench		吞吐、时延
	Redis		吞吐
	Go(云原生)		吞吐
	DeathStarBench/sofaload		吞吐
搜推广-AI化workload	召回ScaNN		满足QOS的吞吐
	粗排		满足QOS的吞吐
	精排(dlrm+dcn)		满足QOS的吞吐
泛AI场景	预处理	训练场景Pytorch框架下发	时长
		推理场景vLLM框架下发	时长
	多模态数据处理	x264解码	时长
		x265解码	时长
		图像解码	时长
		大模型后处理采样	时长
		图像前处理融合算子 (图像resize/crop/RGB/yuv转换等)	时长
	RAG	Embedding生成和查询	时长
		全文检索BM25	时长
多媒体 视频编解码	x264编码		执行时间*实例数
	x265编码		执行时间*实例数

提取典型负载业务模型（以PolarDB为例）

业务流	数据集	请求	组网
OLTP	单个实例生成10张表， 每表10W行记录	point_select range_select 主从同时下发	主从架构 BP>数据量 线程数>核数
OLAP	dsdgen固定生成	tpch主下发	主从架构



验证负载典型性：
修改非典型性参数，对比修改前后CPU、访存等关键特征的相似度，共对比3组修改用例，平均相似度0.8接近1，说明负载典型性高。

测试指标：三层指标设计反应真实特征，满足不同场景&用户诉求

数据库：

- 1、关注数据库业务性能提升（TPCC指标）；
- 2、数据库RTO时间；

虚拟化：

- 1、提高资源利用率情况（CPU、内存带宽）；
- 2、RPC通信时延的降低；
- 3、大规格虚机&资源利用率；
- 4、减少虚拟机密度及内存碎片问题；

搜推广业务：

- 1、降低RPC通信时延；
- 2、提升内存带宽，读写速率；

金融交易场景：

- 1、降低网络通信时延；
- 2、提高数据同步速率；

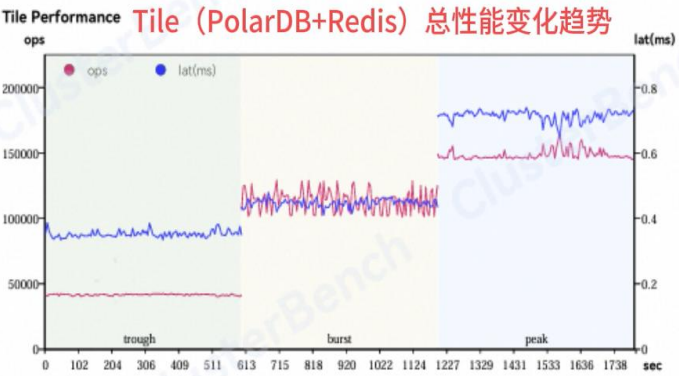
数据库一体机场景：

- 1、高带宽、低时延通信；
- 2、消减IO瓶颈；

层级	主要指标	获取方式	样例
L2 集群系统业务综合指标（峰值性能）	业务吞吐	运行负载，统计单位时间内，集群处理的业务请求量。	ops/rps/tps=1468
	业务时延	从client端发送业务请求到收到返回结果所需要的时间。	280ms
	虚拟机密度	集群部署的虚机/容器个数： $n * x * \sum_{k=1}^x s_k / N$	12 VM / Host
	资源分配率	统计虚拟化分配的CPU、内存与物理CPU、内存的比作为分配率。当分配率大于1时，分配率=超分比。	CPU: 150% 内存: 110%
	资源利用率	测试周期内，满足Qos约束下，集群资源可最大利用的比率（以CPU和内存的平均资源利用率进行代表）。	50%
L1 集群系统关键技术特性指标（中间件性能）	虚拟机迁移时延	虚机从原节点迁移到目标节点的过程中，统计对实例上运行业务的中断时间，耗时越少，迁移性能越好。	100ms
	Rpc通信时延	基于bRPC框架，使用多容器组成RPC通信流，测量容器之间单跳和多跳的时延。	0.21ms
	跨节点内存访问	用Imbench统计跨节点远端内存(CXL/UB/RDMA)的访问时延。	250ns
	特定脏页率下迁移效率	通过microbench构造不同脏页率，内存处于不同脏页率情况下，统计虚拟机迁移时长。	2min
L0 集群基础规格能力指标（基础能力）	集群总算力	基于计算密集型负载运行得到，结果进行标准化。	205
	集群内存访问总带宽	用Imbench统计本地内存的总带宽。	4.5GB/s
	节点互联时延	通过iperf得到集群系统中各种计算节点之间进行数据传输所需经过的延迟。	1000ns
	Non/Cacheable可访问物理内存容量	集群系统中单个处理器可直接访问的Non/Cacheable物理内存容量	500GB

ClusterBench已完成UB/CXL等系统支持，并在OAI主论坛完成发布，后续将进入RC阶段

Open AI Infra



代码仓

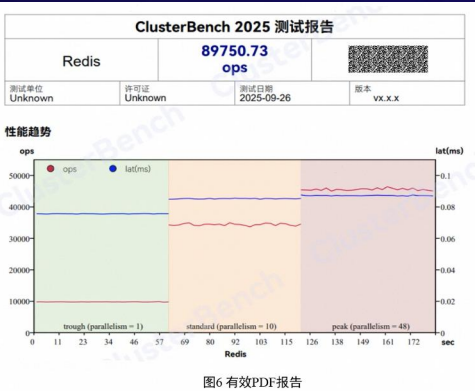
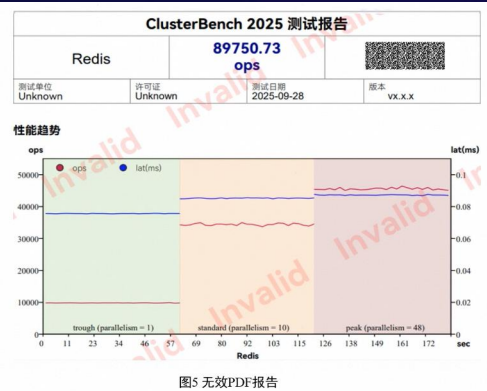
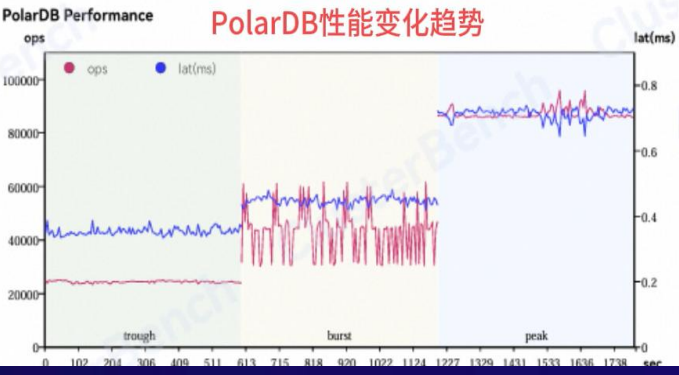
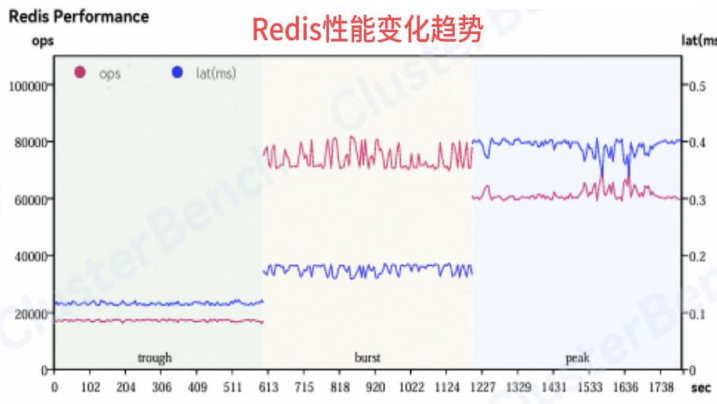


图5 无效PDF报告

图6 有效PDF报告

Open AI Infra社区大会主论坛+ Bench论坛

通算超节点应用场景 workshop

计算性能工程年度会议(高校+产业)

2026.01

2026.02

2026.03

2026.04

2026.06

2026.07-09

2026.10

2026.12

- 工具DEMO演示
- 启动不同平台的兼容性测试 (ARM/X86/CXL等)
- 测试大纲意见
- 征集缺失负载
- 源代码可获取
- 明确测试框架+测试模式
- +测试指标, 形成测试大纲
- 超节点测评, Beta版本发布
- 超节点系统评价指标体系研究报告 (三层指标评价体系, 通算+智算)
- 迭代优化
- Release版本发布
- 超节点测试结果
- 正式发布

邀请产业界联合构建 自主、开放、面向未来的计算基础设施性能基准



计算产品性能基准工作组
Computing Product Performance Benchmark Workgroup

CPPB

WG

AISBench人工智能系统性能评测基准委员会

Artificial Intelligence System Performance Evaluation Benchmark Committee



Global Computing Consortium

全球计算联盟

通用计算：
ClusterBench/CPU Bench
/BenchSEE.....

群聊：集群标准与工具



AI计算：
AISBench.....