

智冷未来： 液冷技术的全栈实践与零碳跃迁

海悟集团

彭 琰

科技 · 品质 · 省心

目录

Contents

01

算力革命与协同散热必然性

02

冷板技术攻坚

03

开放架构与合作

04

液冷微模块方案

05

风液融合方案

06

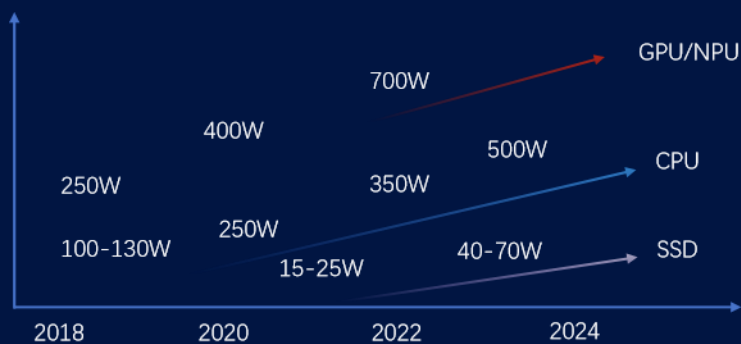
TCO/TCE深度分析

07

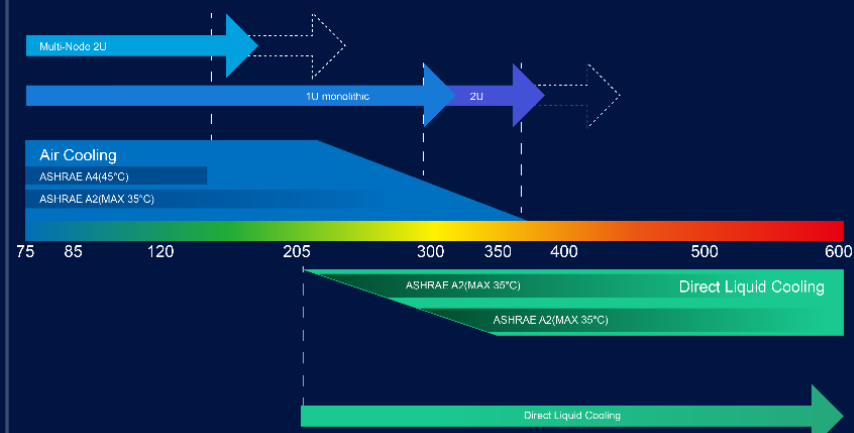
余热循环经济体系

算力革命与协同散热必然性

功率密度持续上升



风冷散热到达极限
CPU: TDP 350W GPU 750W



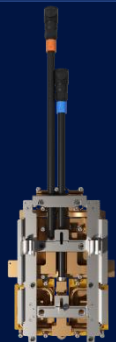
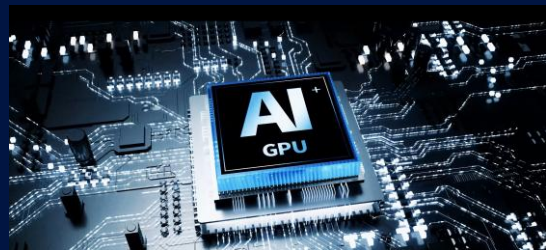
注: ASHRAE推荐

通用计算14kW
AI服务器40/120kW

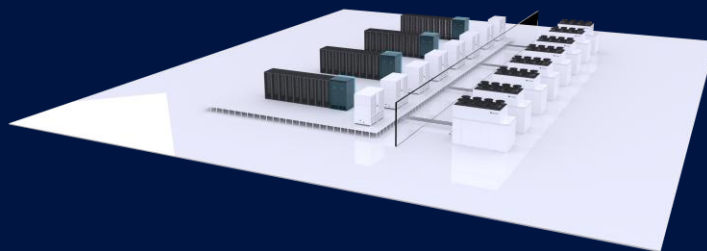
	通用服务器		AI推理服务器		AI训练服务器 (DGX H100)		AI训练服务器 (GB200)	
	数量	功耗 (W)	数量	功耗 (W)	数量	功耗 (W)	数量	功耗 (W)
CPU	2	410	2	700	2	700	2700 (1CPU+2GPU)	
GPU			8	600	8	5600		
芯片组	1	15	1	15	1	15		
内存+硬盘+其它	17	180	25	400	49	800		
TDP合计	605		1715		7115			
实际功耗	700		2000		10200		9800	
U位	2U		4U		8U		2U	
柜	14kW		20kW		40kW		120kW	

NVIDIA GPU从A100开始提供标准的液冷版本

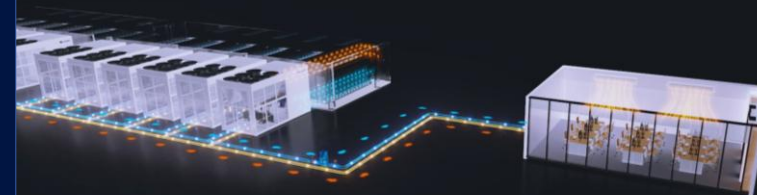
海悟三级协同策略



部件级
液冷覆盖80%高热元件



架构级
动态风液比例调配



系统级
余热能源网络接入

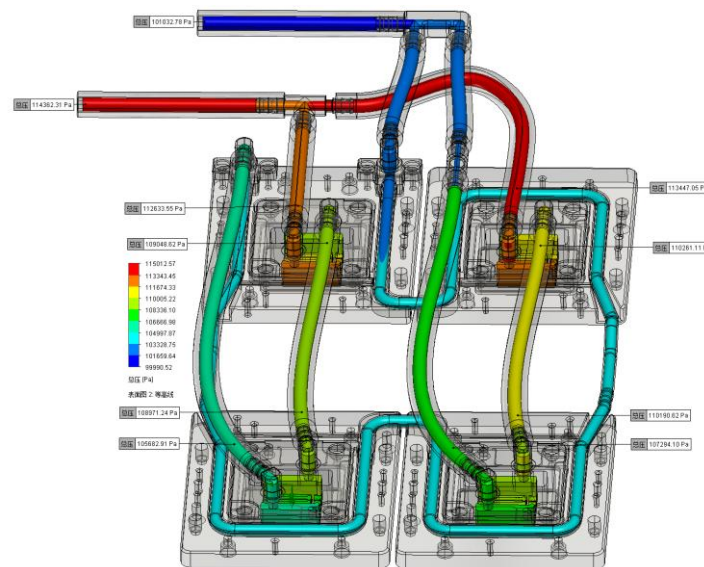
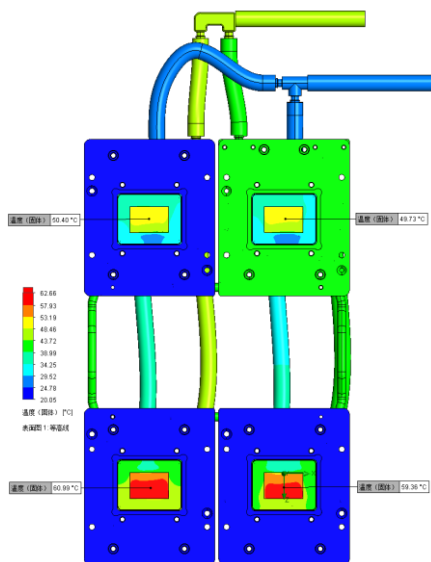
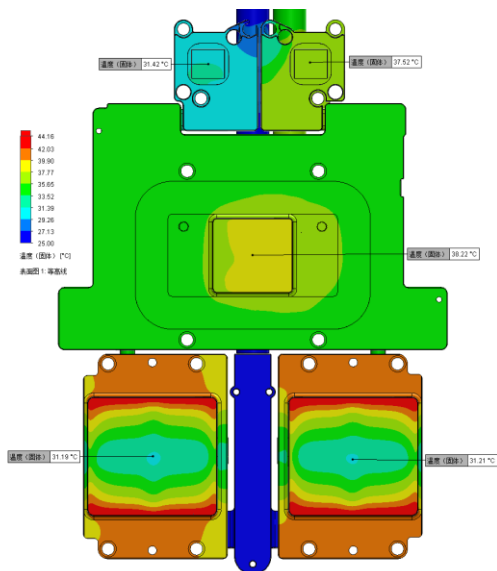
➤ 冷板相关技术攻坚

流场优化：正常流(normal flow)转为并行流(parallel flow)可以大幅增加换热系数

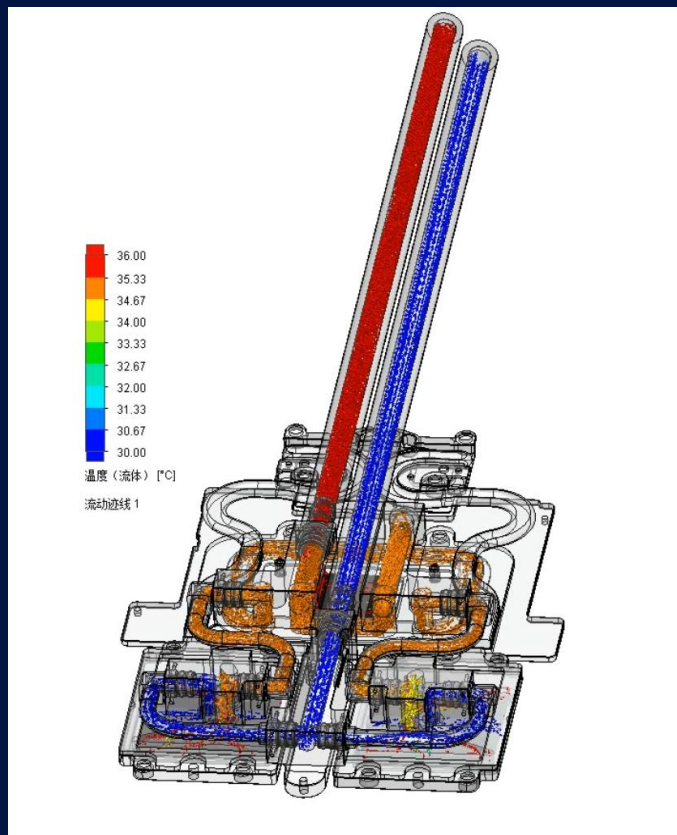
微流道设计：优化通道宽高比 (~7.4) 和鳍片宽高比 (~8)，实现低热阻 (0.07 K/W @ 0.36 L/min)

热性能：可冷却高达300 W/cm²，热阻低至0.02 °C/(W/cm²)，优于传统冷板

其他优化技术：拓扑优化和CFD模拟改善流体分布，减少热点，提升温度均匀性 (如I型歧管，热阻最低)



冷板相关技术攻坚



冷板流道仿真云图



东莞市科学技术局
WWW.DGSTB.DG.GOV.CN

高效高热流密度散热 关键技术研究及产业化

2025年

政府专项：
2025年立项《高效高热流密度散热
关键技术研究及产业化》



2026年两相液冷中试验证

➤ 开放架构与合作

海悟配合开放架构，按照IEC、OCP、H标准



IEC标准



OCP标准



H标准

➤ 开放架构与合作



技术参数

- 机架规格：42U/48kW
- 响应速度：漏液检测 < 50ms
- 交付内容：CDU/PDU/歧管/UQD

商业模式

- ODM合作流程：
- 需求对接 → 联合设计 → 产线共线 → 联合交付
- 香港项目定位：亚太区液冷技术验证节点

➤ 液冷微模块方案

场景：

- 面向金融、政府、企业等中小机房场景

技术：

- 采用预制化模块和全链路温控技术，通过液冷架构实现高密度计算环境

价值：

- 高效制冷高密无忧
- 多重保护、安全可靠
- AI寻优、高效运营
- 快速部署、灵活扩容

亮点：

- 连续两次获得德国红点奖

风液融合方案

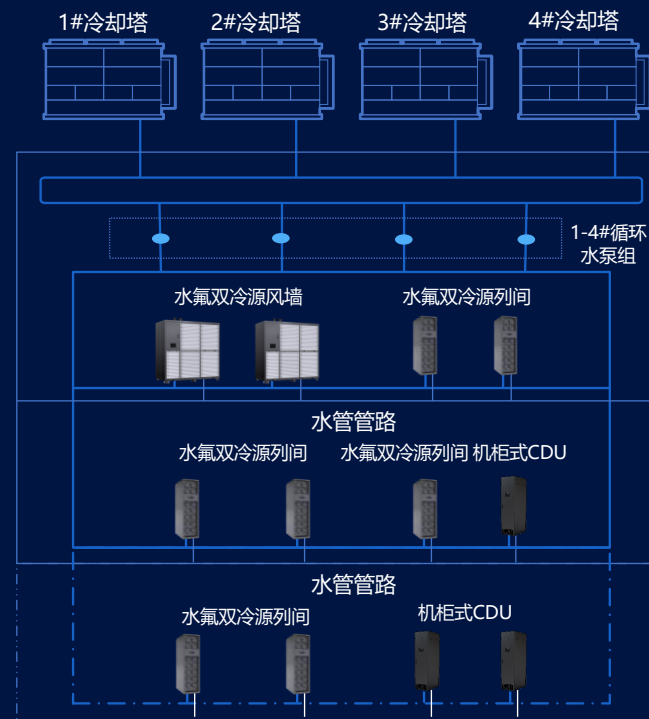


全景应用

全程可靠

全域适配

全速交付



2023年 (初期)

一次侧冷源建设完成,
适配液冷预留接口;
室内100%风冷

2025年 (扩容)

增设液冷系统,
一次侧无需改造

2028年 (未来)

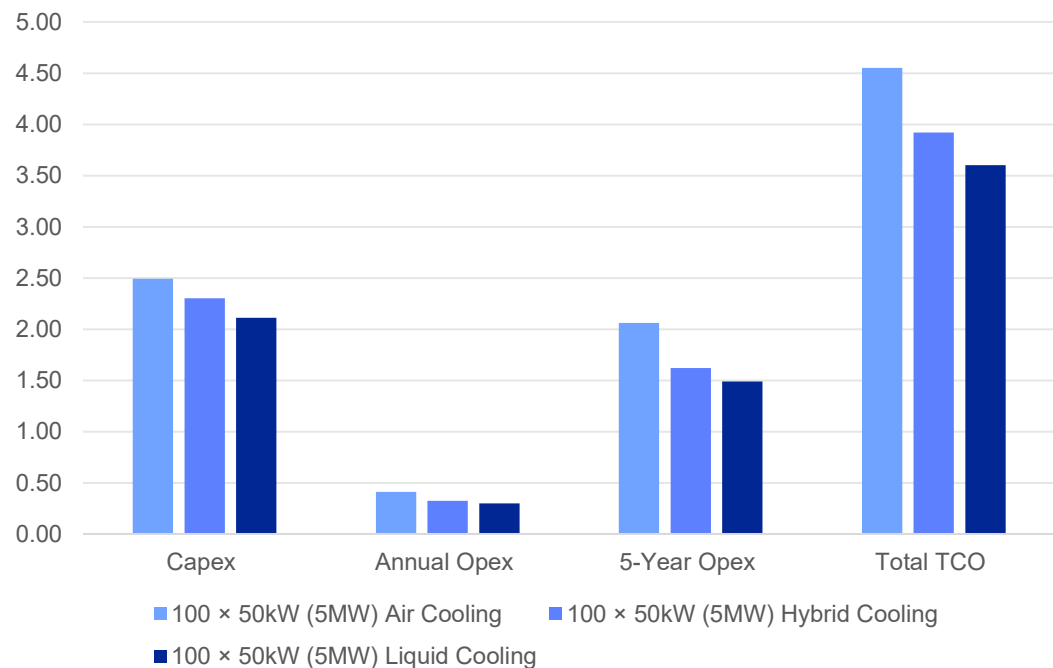
兼容高密度液冷;
按需弹性调节风液比

动态调控逻辑

- 实时监测: 200+个温度/压力传感器
- 调节策略: $\pm 10\%$ 风液出力比例
- 切换阈值: GPU温度 $> 70^{\circ}\text{C}$ 时自动增强液冷

TCO分析

50kW/Rack TCO Analysis (亿元)



资本支出 (Capex) 计算:

风冷的成本为每瓦特 7.02 美元 (每瓦特 49.84 人民币)

液冷的成本为 5.68 美元加上 1.3 除以紧凑系数 (C) (等值人民币), 其中紧凑系数 C 为 (密度 / 10 千瓦)

混合式冷却的资本支出是风冷和液冷资本支出的平均值

运营成本 (Opex) 计算:

总功率 = 信息技术 (IT) 设备功率 × 电源使用效率 (PUE)

年度运营成本 = 总功率 × 8760 小时 × 0.71 元 / 千瓦时 (0.1 美元 / 千瓦时)

维护成本每年约为 355 万至 1775 万元人民币, 根据密度按比例调整

总拥有成本 (TCO) 计算:

Capex + 5年Opex.

算力碳效分析



《数据中心算力碳效白皮书》（ODCC）明确定义：

服务器碳效模型定义

服务器碳效 CEPS：表示单位算力所对应的碳排放量

公式：CEPS = C / S

C：碳排放总量

S：服务器算力性能

碳排放量计算 (C)

计算公式：C = n × E使用

n：碳排放因子，0.6808 kgCO₂/kWh

E使用：全生命周期（5年）能耗

E使用 = Pserver × 24 × 365 × 5

Pserver：通过 BenchSEE 工具 测得服务器在11种负载下的平均功耗（W/T）

算力性能计算 (S)

理论算力：CPU理论浮点算力 × 时间（5年）× 50%利用率

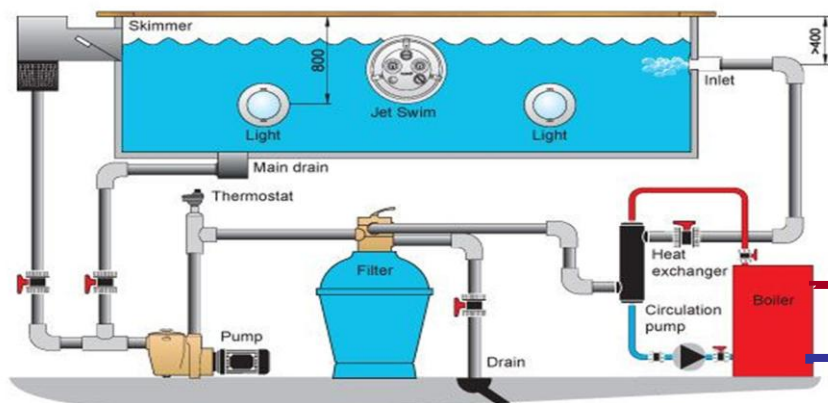
实际算力（Sm）：使用 **SPECrate 2017** 测试

Sm = 50% × SPECrate Integer + 50% × SPECrate Floating

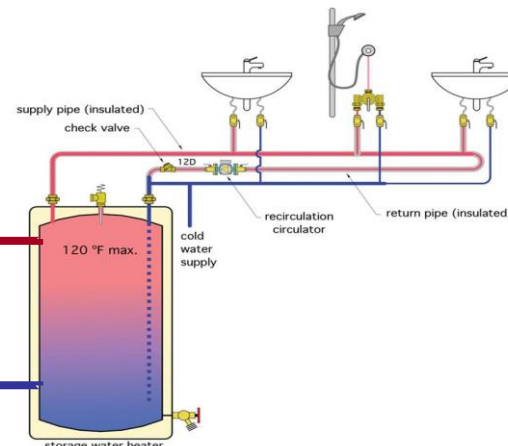
- PUE（Power Usage Effectiveness）：衡量数据中心总能耗与IT设备能耗的比值，反映能源利用效率，但未直接关联碳排放
- TCE（Total Carbon per Compute Unit）：以单位算力的碳排放量为指标，将算力输出与碳排放直接挂钩，更贴合“双碳”目标需求
- 国内项目招标普遍要求披露“单位算力碳排放”（kgCO₂/TFLOPS）。
- NVIDIA在2024 GTC大会发布DGX GB200系统时，同步披露“0.28 kgCO₂/TFLOPS”的碳效数据

➤ 余热循环经济体系

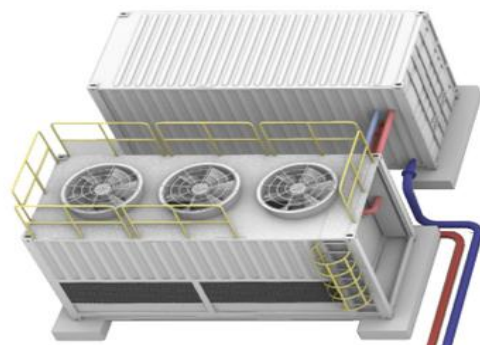
四季恒温游泳池



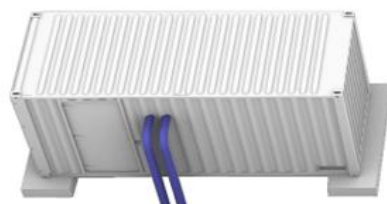
生活热水供应



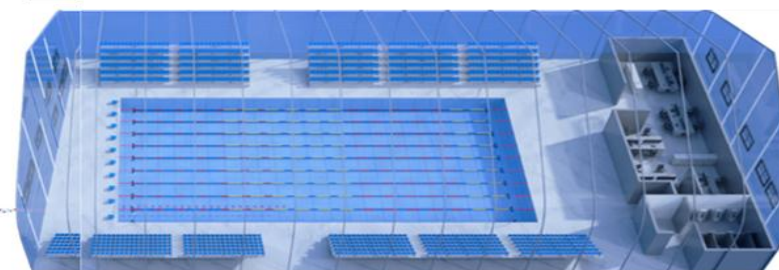
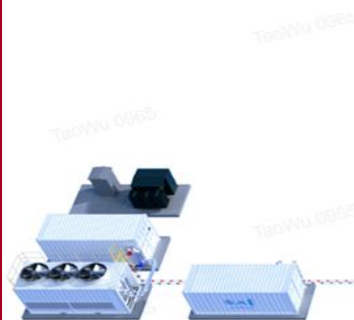
家庭供暖



高性能算力系统
和高效率热回收系统



算力供热系统



感谢您的聆听

THANK YOU FOR LISTENING



电话: 13141287663
邮箱: pengyan@haiwu.com

WWW.HAIWU.COM