

面向视频解析任务的异构 GPU 性能建模与中立评估框架研究

王俊¹, 刘本军²

1. 湖北集防科技有限公司, 湖北 宜昌 443000

2. 湖北三峡职业技术学院电子信息学院, 湖北 宜昌 443000

DOI: 10.61369/TACS.2026050045

摘要 : 视频结构化解析是智能安防系统的核心技术, 但面临多源异构硬件选型难题。针对当前评估中“硬件参数导向”的局限性, 本文提出一种面向视频解析任务的异构 GPU 中立评估框架。通过分析 GDDR 与 HBM 显存架构的技术差异, 构建包含有效解码并发度、推理吞吐量、带宽-计算平衡性、跨框架适配成本、能耗效率的五维评估模型。基于该框架对 NVIDIA Ampere、AMD CDNA2、华为昇腾等架构进行建模与验证, 发现传统显存位宽指标存在架构依赖性偏差。提出的任务-架构匹配度指标 (TAM) 为多品牌 GPU 选型提供了方法论参考。

关键词 : 异构计算; 视频结构化解析; GPU 性能建模; 中立评估框架

Research on Heterogeneous GPU Performance Modeling and Neutral Evaluation Framework for Video Parsing Tasks

Wang Jun¹, Liu Benjun²

1. Hubei Jifang Technology Co., Ltd., Yichang, Hubei 443000

2. School of Electronic Information, Hubei Three Gorges Polytechnic, Yichang, Hubei 443000

Abstract : Video structured parsing is a core technology of intelligent security systems, yet it faces the difficulty of selecting multi-source heterogeneous hardware. Aiming at the limitations of "hardware parameter-oriented" evaluation in current research, this paper proposes a heterogeneous GPU neutral evaluation framework for video parsing tasks. By analyzing the technical differences between GDDR and HBM memory architectures, a five-dimensional evaluation model is constructed, covering effective decoding concurrency, inference throughput, bandwidth-computation balance, cross-framework adaptation cost, and energy efficiency. Based on this framework, modeling and verification are carried out on architectures such as NVIDIA Ampere, AMD CDNA2, and Huawei Ascend. It is found that the traditional video memory bitwidth index suffers from architecture-dependent bias. The proposed task-architecture matching (TAM) metric provides methodological support for multi-brand GPU selection.

Keywords : heterogeneous computing; video structured parsing; GPU performance modeling; neutral evaluation framework

引言

随着智慧城市建设推进, 城市级视频监控系统数据量呈指数级增长。以“雪亮工程”为例, 单节点需处理数千路高清视频流的实时结构化解析, 涉及目标检测、行为分析等多模态 AI 任务^[1]。GPU 因其强大的 SIMT 计算能力成为主流加速方案。

然而, 视频解析场景下的 GPU 选型面临异构架构评估困境。当前市场存在多种技术路线: NVIDIA Ampere 架构中, A100 采用 HBM2e 显存, A40 采用 GDDR6 显存; AMD CDNA2 (MI210) 基于 HBM2e; 华为昇腾系列采用自研达芬奇架构与 HBM2e^[2-5]。这些架构在显存类型、计算单元设计等方面存在本质差异, 传统“显存带宽”“核心频率”等参数难以客观反映实际性能^[6]。

现有研究存在三方面局限: 一是侧重单一品牌 GPU 的算法优化, 缺乏跨平台对比; 二是评估指标过于依赖峰值算力 (TFLOPS), 忽视视频解析任务的 I/O 密集特性与内存墙 (Memory Wall) 约束^[6]; 三是现有基准测试 (如 MLPerf Inference^[7]) 侧重通用 AI 模型, 缺乏针对视频解析全链路 (解码-预处理-推理-后处理) 的专项评估。本文提出一种去架构依赖的中立评估框架, 核心贡献包括: (1) 解构 GDDR 与 HBM 显存体系在视频编解码、模型推理环节的技术差异; (2) 建立视频解析任务的计算-访存模式模型, 识别性能瓶颈; (3) 提出 TAM 指标及 AHP-熵权组合赋权法; (4) 基于公开技术文档验证 HBM 与 GDDR 架构在视频解析场景下的性能差异。

作者简介:

王俊 (1974—), 男, 湖北集防科技有限公司, 研究方向为视频监控、视频大数据分析。

刘本军 (1976—), 男, 湖北三峡职业技术学院, 研究方向为计算机网络、视频监控。

一、异构 GPU 架构技术差异分析

（一）显存子系统架构对比：GDDR vs HBM

显存带宽是视频解析任务的性能瓶颈之一，当前主流 GPU 采用两种显存架构：（1）GDDR 体系：以 NVIDIA A40 为代表，采用 384bit 位宽，通过高频（16 Gbps）提升带宽，总带宽达 696 GB/s，优势在于成本可控、容量扩展灵活（48 GB GDDR6），但高负载下易遭遇热节流^[2]。（2）HBM 体系：以 NVIDIA A100、AMD MI210、华为昇腾 910B 为代表，采用 3D 堆叠技术实现 4 096 bit 位宽，频率相对较低（1.6~2.0 Gbps），但总带宽可达 1.6~3.2 TB/s，通过硅中介层缩短传输距离，显著降低访存延迟与功耗^[3-5]。

关键差异在于：视频解析任务涉及大规模特征图的频繁读写。在低算术强度（AI<10）的预处理阶段，HBM 架构的高并行通道数具有优势；而在高算术强度（AI>100）的推理阶段，GDDR 架构通过高频率可维持竞争力^[6]。

（二）计算单元与视频解码能力差异

不同厂商针对 AI 推理与视频解码采用不同硬件策略，关键差异如表 1 所示。

表 1 异构 GPU 架构视频解析能力对比				
对比维度	NVIDIA A100 80GB	NVIDIA A40 48GB	AMD MI210 64GB	华为昇腾 910B 64GB
架构	Ampere (GA100)	Ampere (GA102)	CDNA2 (gfx90a)	达芬奇
显存配置	HBM2e 80 GB 2.0 TB/s	GDDR6 48 GB 696 GB/s	HBM2e 64 GB 1.6 TB/s	HBM2e 64 GB 3.2 TB/s
峰值算力 (INT8)	624 TOPS	~150 TOPS	181 TOPS	640 TOPS
视频解码单元	5×NVDEC (第 5 代)[9]	2×NVDEC	2×VCN 2.6.0	32×DVPP
官方解码能力	157 路并发 [9]	未直接标称	未直接标称	未直接标称
数据来源	^{[8][9]} 官方白皮书	^[4] 官方数据表	^[5] 官方规格页	^[7] 华为白皮书

视频解析任务呈现阶段性特征：前端解码（高 I/O）、中段推理（高计算）、后处理 NMS（高分支）。由表 1 可见：（1）A100 与 A40 采用同代 Ampere 架构，但显存类型截然不同（HBM2e vs GDDR6），可隔离显存体系差异；（2）A100 配备 5×NVDEC，1080p30 HEVC 并发解码能力达 157 路^[9]，A40 配备 2×NVDEC，各厂商在视频处理单元数量上差异显著；（3）NVIDIA Tensor Core、AMD Matrix Core、华为 Cube Core 的异构设计使单一峰值算力指标无法反映实际计算需求。因此，需建立面向视频解析任务特征的分阶段评估模型。

二、视频解析任务特征与性能瓶颈建模

（一）视频结构化解析流水线分析

典型视频解析任务可分解为五个阶段，各阶段对硬件资源需求差异显著，如表 2 所示。

表 2 视频结构化解析流水线五阶段分析				
阶段	计算特征	主要瓶颈	关键资源需求	算术强度 (AI)
视频解码	硬解码单元负载，GPU Shader 利用率低	解码器并发路数限制，显存带宽	单路 1080p@30 fps 约 1.5 GB/s 原始数据	低 (AI<1)
图像预处理	像素级操作 (Resize/Normalize)	显存带宽（读写双通）	低算术强度任务，HBM 带宽利用率比 GDDR 高 15%~20%	极低 (AI ≈ 2~5)
深度学习推理	卷积层 (GEMM)，全连接层混合	计算单元利用率，权重加载带宽	ResNet-50 推理，INT8 精度优先	中等 (AI ≈ 50)
后处理	NMS/ 坐标映射，时序滤波	分支预测效率，原子操作延迟	CPU-GPU 协同，内存随机访问	低 (AI<10)
结构化输出	JSON/XML 编码，网络 I/O	PCIe 带宽，CPU-GPU 协同效率	网络吞吐，序列化开销	极低 (AI<1)

（二）内存墙约束下的有效算力模型

视频解析任务的实际性能受“内存墙”效应制约。依据 Roofline 性能模型^[6]，有效算力利用率 η 定义为：

$$\eta = 1 / (1 + (B_{required} / B_{actual}) \cdot (1/AI))$$

其中，B_{required} 为算法所需带宽，B_{actual} 为硬件提供带宽，AI 为算术强度（每字节访存对应的浮点运算数）。

对于 ResNet-50 前向推理，AI ≈ 50 FLOPs/Byte。若 GPU 提供 1 TB/s 带宽、峰值算力 10 TFLOPS，则 η = 83.3%。但对于预处理阶段（AI ≈ 2）， η 降至 16.7%，此时性能完全由显存带宽决定。关键洞察：HBM 架构的高带宽在 AI < 10 的低算术强度任务（解码、预处理）中具有决定性优势，而 GDDR 架构仅在 AI > 100 的高计算密度任务中才能发挥峰值性能。

三、异构 GPU 中立评估框架构建

（一）传统硬件参数的架构依赖性缺陷

现行选型常依赖以下参数，但存在架构偏向性：（1）显存位宽：HBM 天然具有高得多位宽（4 096 bit vs 384 bit），但技术原理不同，直接对比无意义；（2）基础频率：GDDR6 可达 16~21 Gbps，HBM2e 仅 1.6~2.0 Gbps，但后者通过位宽补偿实现更高总带宽；（3）峰值算力（TFLOPS）：仅反映理论能力，忽视 I/O 瓶颈；（4）显存类型：限定特定类型即排除其他架构。这些参数未考虑视频解析任务的阶段性特征，也未考虑不同架构在内存墙

约束下的实际效率差异^[6]。

（二）五维中立评估指标体系

本文提出基于任务实际负载的五维指标，消除架构差异带来的评估偏差：

指标1：有效解码并发度（EDC） $EDC = \min(Ndecoder, Vmemory / (Fframe \times Tbuffer))$

其中 Ndecoder 为总解码并发能力，Vmemory 为可用显存，Fframe 为单帧数据量，Tbuffer 为时序缓冲帧数。

指标2：推理吞吐量标准化值（NIT）。采用 ResNet-50 为基准，测量 INT8 精度下实际 FPS 并归一化：

$NIT = Throughputactual / Throughputpeak \times 100\%$

指标3：带宽 - 计算平衡系数（BCB）。评估架构设计是否匹配视频解析任务的访存 - 计算比：

$BCB = (Bactual \times Altask) / Fpeak$

理想值接近 1，过高表示计算单元闲置（带宽瓶颈），过低表示带宽过剩（计算瓶颈）。

指标4：跨框架适配成本（AC）。量化从 PyTorch/ TensorFlow 迁移至目标平台的工程复杂度，包括算子支持率、模型转换成功率等。

指标5：能耗效率（EE）。定义为单位功耗下的有效解析路数：

$EE = EDC / Pcard \text{ (路 / 瓦)}$

（三）任务 - 架构匹配度（TAM）模型

采用 AHP- 熵权组合赋权法确定权重。步骤1：AHP 确定主观权重，构建判断矩阵（表3），得 WAHP = [0.263, 0.263, 0.158, 0.105, 0.105]T。一致性检验：λ max=5.02，CR=0.004<0.1，通过检验。

表3 指标重要性判断矩阵					
指标	EDC	NIT	BCB	AC	EE
EDC	1	1	2	3	3
NIT	1	1	2	3	3
BCB	1/2	1/2	1	2	2
AC	1/3	1/3	1/2	1	1
EE	1/3	1/3	1/2	1	1

步骤2：熵权法确定客观权重，得 Wentropy = [0.238, 0.241, 0.185, 0.172, 0.164]T。步骤3：采用乘法合成法（α=0.5），得最终权重 W = [0.250, 0.252, 0.171, 0.134, 0.133]T。

TAM 计算公式：

$TAM = \sum i=15 wi \times (Mi - Mmin) / (Mmax - Mmin)$

四、异构 GPU 性能评估验证分析

（一）评估对象与数据来源

选取四款代表性 GPU 进行对比，涵盖 HBM 与 GDDR 两大显存体系。硬件规格均来自厂商官方技术白皮书^{[3-5][7][9]}。

表4 评估对象硬件规格

参数	NVIDIA	NVIDIA A40	AMD	华为昇腾
	A100 80GB (HBM)	48GB (GDDR)	MI210 64GB (HBM)	910B 64GB (HBM)
架构	Ampere (GA100)	Ampere (GA102)	CDNA2 (gfx90a)	达芬奇
显存类型 / 容量	HBM2e / 80 GB	GDDR6 / 48 GB	HBM2e / 64 GB	HBM2e / 64 GB
显存带宽	2.0 TB/s	696 GB/s	1.6 TB/s	3.2 TB/s
峰值算力 (INT8)	624 TOPS	~150 TOPS	181 TOPS	640 TOPS
视频解码单元	5 × NVDEC (第5代)[9]	2 × NVDEC[4]	2 × VCN 2.6.0[5]	32 × DVPP[7]
TDP	300 W	300 W	300 W	310 W

注：A40 与 A100 形成同架构下的显存类型对照。A100 解码能力 157 路为官方实测数据^[9]。

（二）单指标性能对比分析

基于 MLPerf Inference v2.1 公开结果^{[6][8]}进行二次换算，各 GPU 在 ResNet-50 模型上的端到端性能如表5所示。在低算术强度（AI<5）的预处理阶段，HBM 架构带宽利用率显著高于 GDDR 架构，差距达 15%~20%；随着算术强度增加，GDDR 架构通过高频率补偿，差距逐渐缩小。

表5 显存带宽利用率对比（%）

算子类型	算术强度 (AI)	A100 (HBM)	A40 (GDDR)	MI210 (HBM)	昇腾910B (HBM)
图像缩放	2.1	45%	32%	62%	58%
归一化	3.5	58%	41%	71%	68%
卷积层1	25.6	82%	78%	85%	88%
卷积层5	78.3	91%	89%	89%	92%

基于硬件规格与公式推导的解码并发度估算如表6所示。A100 的 EDC 受显存容量限制（142 路），A40 受解码单元数限制（72 路）。

表6 视频解码并发性能估算

GPU 型号	总解码并发能力	显存瓶颈路数	实际限制因素	计算依据
A100 (HBM)	157 路 ^[9]	142 路	显存容量	Cdecoder=157(官方数据)Vmemory/(Fframe × Tbuffer)=142
A40 (GDDR)	72 路	85 路	解码单元数	Cdecoder=72 显存未成为瓶颈
MI210 (HBM)	120 路	98 路	解码单元数	Cdecoder=120
昇腾910B (HBM)	128 路	64 路	显存容量	Cdecoder=128 Vmemory/(Fframe × Tbuffer)=64

（三）TAM 综合评估结果

基于各指标归一化值与组合权重，计算 TAM 得分（满分 1.0）：

表7 异构 GPU TAM 综合评估结果

评估维度	权重	A100 (HBM)	A40 (GDDR)	MI210 (HBM)	昇腾 910B (HBM)
EDC	0.25	0.85	0.42	0.72	0.68
NIT	0.25	0.92	0.74	0.78	0.81
BCB	0.17	0.88	0.65	0.95	0.93
AC	0.13	0.95	0.95	0.65	0.60
EE	0.13	0.75	0.68	0.82	0.88
TAM 综合得分		0.86	0.69	0.78	0.77

发现与讨论：（1）HBM 架构（A100、MI210、昇腾 910B）的 TAM 得分显著高于 GDDR 架构（A40），差距达 17%~25%，主要源于低算术强度任务中 HBM 的高带宽利用率优势及 GDDR 的显存带宽瓶颈（696 GB/s vs 2.0 TB/s）。（2）架构互补性：A100 在解码并发与软件生态占优；MI210 在带宽 - 计算平衡性上表现最佳（BCB=0.95）；昇腾 910B 能效比突出（EE=0.88）但适配成本较高；A40 虽在视频解析场景落后，但成本优势明显，适合轻量级推理任务。（3）若仅比较峰值算力，A100（624

TOPS）约为 A40（150 TOPS）的 4.2 倍，但 TAM 得分仅高出 25%，说明 GDDR 架构在特定场景仍具性价比。（4）本评估基于公开技术文档与理论建模，实际部署性能可能受驱动版本、散热条件等因素影响，建议作为选型初筛工具，最终决策需结合实测验证。

五、结论

本文针对大规模视频解析场景下的异构 GPU 选型难题，构建了面向任务特征的中立评估框架。主要结论如下：

（1）传统基于显存位宽、基础频率的评估指标存在架构依赖性，无法客观比较 GDDR 与 HBM 差异。提出的 EDC、NIT、BCB、AC、EE 五维任务导向指标消除了这种偏差，在低算术强度的解码与预处理阶段，HBM 架构带宽利用率显著优于 GDDR 架构（差距 15%~20%）。

（2）采用 AHP-熵权组合赋权法确定 TAM 指标权重，CR=0.004<0.1 通过一致性检验。通过 GDDR 代表（A40）与 HBM 代表（A100、MI210、昇腾 910B）的对比验证，证明了框架在区分显存体系差异方面的有效性。

（3）TAM 评估揭示了不同架构在视频解析场景下的真实性价比。HBM 架构综合得分显著高于 GDDR 架构，但 GDDR 架构在成本敏感型轻量级任务中仍具应用价值。

参考文献

[1] 王坤峰, 荀超, 王飞跃. 平行视觉: 基于 ACP 的智能视觉计算方法 [J]. 自动化学报, 2016, 42(10): 1490-1500.

[2] Abdelkhalik H, Arafa Y, Santhi N, et al. Demystifying the Nvidia Ampere Architecture through Microbenchmarking and Instruction-level Analysis[C]. IEEE HPEC, 2022. arXiv:2208.11174.

[3] NVIDIA Corporation. NVIDIA A100 Tensor Core GPU Datasheet[EB/OL]. 2021.

[4] NVIDIA Corporation. NVIDIA A40 GPU Datasheet[EB/OL]. 2021.

[5] AMD Corporation. AMD Instinct MI210 Accelerator Specifications[EB/OL]. 2021.

[6] Williams S, Waterman A, Patterson D. Roofline: An insightful visual performance model for multicore architectures[J]. Communications of the ACM, 2009, 52(4): 65-76.

[7] 华为技术有限公司. 昇腾 910 AI 处理器技术白皮书 [R]. 2021.

[8] MLCommons. MLPerf Inference v2.1 Results[EB/OL]. <https://mlcommons.org/benchmarks/inference-datacenter/>, 2022.

[9] NVIDIA Corporation. NVIDIA A100 Tensor Core GPU Architecture[EB/OL]. 2020.