

Deep Learning Inference on Heterogeneous Mobile Processors: Potentials and Pitfalls

Sicong Liu
Northwestern Polytechnical
University

Wentao Zhou
Harbin Engineering University

Zimu Zhou
City University of Hong Kong

Bin Guo*
Northwestern Polytechnical
University

Minfan Wang
Chang'an University

Cheng Fang
Northwestern Polytechnical
University

Zheng Lin
Northwestern Polytechnical
University

Zhiwen Yu
Harbin Engineering University
Northwestern Polytechnical
University

ABSTRACT

There is a growing demand to deploy computation-intensive deep learning (DL) models on resource-constrained mobile devices for real-time intelligent applications. Equipped with a variety of processing units such as CPUs, GPUs, and NPUs, the mobile devices hold potential to accelerate DL inference via parallel execution across heterogeneous processors. Various efficient parallel methods have been explored to optimize computation distribution, achieve load balance, and minimize communication cost across processors. Yet their practical effectiveness in the dynamic and diverse real-world mobile environment is less explored. This paper presents a holistic empirical study to assess the capabilities and challenges associated with parallel DL inference on heterogeneous mobile processors. Through carefully designed experiments covering various DL models, mobile software/hardware environments, workload patterns, and resource availability, we identify limitations of existing techniques and highlight opportunities for cross-level optimization.

*Corresponding author: guob@nwpu.edu.cn.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

ACM MobiSys '24, June 03–07, 2024, Tokyo, Japan

© 2024 Association for Computing Machinery.

ACM ISBN 978-x-xxxx-xxxx-x/YY/MM...\$15.00

<https://doi.org/00000.00000>

CCS CONCEPTS

• **Human-centered computing** → Ubiquitous and mobile computing; • **Computing methodologies** → Machine learning.

KEYWORDS

Heterogeneous processors, parallel DL inference

ACM Reference Format:

Sicong Liu, Wentao Zhou, Zimu Zhou, Bin Guo*, Minfan Wang, Cheng Fang, Zheng Lin, and Zhiwen Yu. 2024. Deep Learning Inference on Heterogeneous Mobile Processors: Potentials and Pitfalls. In *Proceedings of In The 22nd ACM International Conference on Mobile Systems, Applications, and Services (ACM MobiSys '24)*. ACM, New York, NY, USA, 6 pages. <https://doi.org/00000.00000>

1 INTRODUCTION

The deployment of deep learning (DL) models has shifted from cloud-centric to mobile devices for on-device intelligence [2, 7, 13, 20]. This transition enables various applications that interact intelligently with users in *real-time*, including biometric authentication on smartphones [20], arm posture tracking on smartwatches [13], 3D object detection on headsets [7], and language translation on home devices [2].

Commercial mobile system-on-chips (SoCs) increasingly integrate a variety of processing units, such as CPUs, GPUs, DSPs, and NPUs, enabling an emerging trend towards accelerating DL inference through *parallel* execution across *heterogeneous* resources [9–11, 25]. It involves mapping DL computations to the processors best suited for their execution, which takes place at two levels (see Fig. 1). *i)* At the *resource-friendly frontend compilation level*, a DL model is compiled into a computational directed acyclic graph (DAG). *ii)* At the *model-adaptive backend compilation level*, the DAG

operators are then mapped and scheduled across heterogeneous computing/memory resources for parallel execution. As mainstream mobile DL frameworks typically restrict computations to a *single* processor at a time [6, 22, 26], existing research developed various strategies to achieve *cross-processor* parallel inference speedup across heterogeneous processors, including *i*) optimizing the distribution of computations by branching and partitioning subgraphs [25], layers [11], or operators [10] among processors; *ii*) balancing loads across processors to minimize asynchronous waiting times [9]; and *iii*) reducing data communication overhead associated with merging results from concurrent executions [28].

Despite prior efforts on the feasibility of parallel inference across processors, their efficacy in *real-world mobile contexts* remains less-explored. We hypothesize that backend compilation-level optimization for maximal utilization of heterogeneous resources, as many existing techniques do [10, 11, 25], may not yield optimal overall performance due to the *diversity* and *dynamics* of the mobile device ecosystem. By evaluating representative techniques across different DL models, mobile software/hardware environments, workload patterns, and runtime resource availability, our empirical study strives to answer the following questions when applying these techniques in practice.

(I) Are existing strategies for parallel inference across mobile processors adequately effective? For previous solutions focused on *backend compilation-level* optimization, our study implies that their *design space* could be improved both *within* and *beyond* the current scope.

- **Unsupported Operators and Processor Fallbacks:** The operators generated at the backend compilation-level optimization are not supported by all processors, leading to operational fallbacks. For instance, when GPUs encounter unsupported operators, the DL computations are forced to revert to CPUs, leading to suspensions of processes and under-utilization of resources. Our experiments show that unsupported operators *vary* across heterogeneous mobile processor combinations, with significant fallbacks in certain cases *e.g.*, for NN-Stretch in ResNet-50 on the Xiaomi 9, the ratio of unsupported operators during parallel inference on the CPU and GPU reaches approximately 48%.
- **Cross-Level Optimization Opportunities:** Existing solutions agnostic to DL models often exhibit *overparameterization* that can be optimized before DAG conversion. Our experiments demonstrate that incorporating frontend compilation-level optimizations, such as pruning, can improve resource utilization and accelerate DL inference at the backend. Also, optimizing operator fusion/parallelism and memory swap/allocation at the model-adaptive backend compilation level can enhance data reuse and eliminate memory segmentation,

broadening the scope of frontend optimization. This indicates a potential to expand the design space via cross-level optimization, prioritizing the parallel execution of crucial operations while selectively managing less critical ones, optimizing the system scheduling for higher concurrency and resource availability.

(II) Is parallel inference across processors always beneficial? Most strategies maximize the resource utilization across processors for DL task alone. Such *optimization objectives* may not fit in real-world mobile contexts.

- **Granularity of Parallel Scheduling:** On mobile heterogeneous multiprocessor devices, a too *large* scheduling granularity can lead to idle chip cores, undermining the efficiency of utilizing multiple processors. In contrast, an overly *fine* granularity may cause process suspensions and blocks due to increased data transmission delays. Furthermore, our experiment suggest that allocating computations to a single processor could, in some instances, outperform distribution across heterogeneous processors due to scheduling overhead.
- **Presence of Competing Processes:** Achieving maximum resource utilization for DL inference does not necessarily result in optimal overall system performance, particularly in multi-process mobile applications. Dedicating excessive computational resources to DL inference can impair the functionality of other components, *e.g.*, UI responsiveness in smartphones and AR apps. Given the *dynamic nature* of runtime resource availability and competing process demands in mobile applications, pinpointing optimal resource utilization levels is essential for improving the overall performance without compromising other processes.

Our main contributions are as follows.

- We conduct pilot studies tailored to evaluate the performance of parallel DL inference strategies across heterogeneous mobile processors in real-world conditions. They assess representative solutions in setups critical to practical applications, yet under-explored in existing empirical studies.
- We identify the limitations and opportunities in the design space and objectives of existing solutions. These insights would facilitate effective parallel inference across heterogeneous resources amid the inherent diversity and dynamics of the mobile device ecosystem.

2 METHODOLOGY

Optimization Stack. As previously mentioned, parallel DL inference across heterogeneous processors can be optimized at two levels, *i.e.*, frontend and backend compilation. Fig. 1 illustrates the generic workflow to map the DL computation

to different resources. (i) *Frontend compilation* translates DL models into an intermediate representation, i.e., *computational graph*, eliminating redundant operators in the graph. (ii) *Backend compilation* further compiles the *computational graph* into an executable binary file, involving three sequential processes: ① *Graph optimization* merges operators that conform to fusion rules and schedule the computation graph for parallel execution based on four levels: intra-operator, inter-operator, sub-graph, and task. ② *Operator distribution* assigns each operator to processors based on operator types and hardware support. ③ *Memory allocation* allocates memory space for data in the graph based on operator allocation and the available memory resource of processors. The *dynamics* in mobile resource availability, system constraints, workload, and input data can influence the performance of generic compilation optimization strategies.

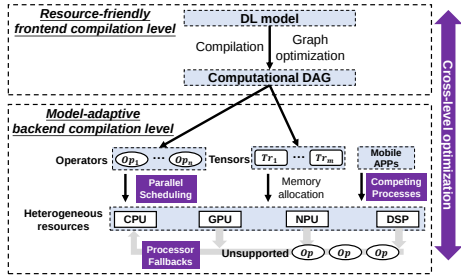


Figure 1: Illustration of parallel DL inference workflow across heterogeneous processors on mobile devices.

Existing Strategies. To our best knowledge, no solution covers the entire optimization stack above. We list representative resource-efficient DL inference methods in two categories based on our generic compilation workflow.

- **Frontend compilation** optimization level reduces the redundancy of DL models to minimize resource demand.
 - **Pruning** [4] removes redundant DL model parameters, channels, and connections.
 - **Low-rank decomposition** [27] adopts SVD to use basis vectors to represent the convolution kernel.
 - **Parameter/activation quantization** [23] represents 32-bit parameters with 8-bit widths.
- **Backend compilation** optimization level improves hardware resource availability and utilization.
 - **Operator fusion** [15] fuses adjacent operators into a new operator to save I/O of intermediate results.
 - **Mace** [26]: *single-processor* inference strategy.
 - **TensorFlow Lite** [6]: *single-processor* strategy.
 - **NN-Stretch** [25]: converts DL models into multi branches to increase computation parallelism.
 - **μ Layer** [11]: divides DL layers into channels and executes diverse channels across processors.
 - **CoDL** [10]: partitions DL model operators into operator chains for operator parallel execution.

Table 1: Specifications of mobile devices for evaluation.

Mobile devices	SoC	Processor conf.	RAM
D_1 : Xiaomi 9	Snapdragon 855	CPU+GPU+DSP	8GB
D_2 : Huawei Nova 7	Kirin 985	CPU+GPU+DSP+NPU	8GB
D_3 : Redmi K30 Pro	Snapdragon 865	CPU+GPU+DSP	6GB
D_4 : Xiaomi 11 Pro	Snapdragon 888	CPU+GPU+DSP	8GB
D_5 : Honor V30 Pro	Kirin 990	CPU+GPU+DSP+NPU	8GB
D_6 : Honor 9	Kirin 960	CPU+GPU+DSP	6GB
D_7 : Huawei Matepad Pro	Snapdragon 870	CPU+GPU+DSP	6GB
D_8 : Pixel 6	Google Tensor	CPU+GPU	8GB

– **Memory allocation** [21] releases memory in time.

Overall Experimental Setups. We aim to assess the limitations and opportunities of the above methods for parallel DL inference across heterogeneous processors, with a special focus on the *diversity* and *dynamics* of mobile devices. We experiment with 8 commercial mobile devices equipped with heterogeneous multiprocessors (see Tab. 1), and simulate various *internal* and *external* system factors affecting the runtime resource availability on these mobile devices (e.g., DL workloads, DVFS strategies, competing processes). In addition to accuracy and latency, we also include multiple intermediate performance metrics such as the reduction in DL resource demands (e.g., model size, memory usage), and the improvement in runtime resource utilization (e.g., utilization rate, cache hit rate, and frame drop rate). As with other empirical studies and benchmarks [5, 14, 24, 30], we mainly test DL models for *visual* related tasks, covering representative models including YOLOv2 [17], VGG-16 [19], PoseNet [31], Fast Style Transfer [1], RetinaFace [3], ResNet-18/50/34 [8] and RegNetX-1.6GF/4GF [16].

3 EXPERIMENTAL RESULTS

3.1 Are existing strategies for parallel inference adequately effective?

3.1.1 Performance of different parallel strategies over diverse mobile devices. We compare the inference latency of six strategies: ① Mace framework executes inference on CPU, ② Mace on GPU with the buffer type, ③ Mace on GPU with the image type, ④ μ Layer on CPU+GPU, ⑤ CoDL on CPU+GPU in parallel with buffer type, and ⑥ CoDL on CPU & GPU in parallel with the image type. We test them on five models, i.e., YOLOv2, VGG-16, PoseNet, Fast Style Transfer (FST), and RetinaFace, across three mobile devices, i.e., Snapdragon 855 (D_1), Snapdragon 870 (D_7), and Kirin 985 (D_2). Figure 2 shows the result. *First*, different parallel strategies exhibit varying speedup performance across diverse DL models and mobile devices. On D_1 , CoDL achieves optimal speedup for parallel inference using image data type, while on D_2 , its latency compared to running on a single GPU ranges from 1.4 $\times$ to 1.8 $\times$. This stems from variations in buffer and image data type support in Adreno and Mali architecture GPUs. *Second*, executing inference on a single processor may outperform distribution across heterogeneous processors in certain cases. Parallel inference of FST on CPU+GPU using

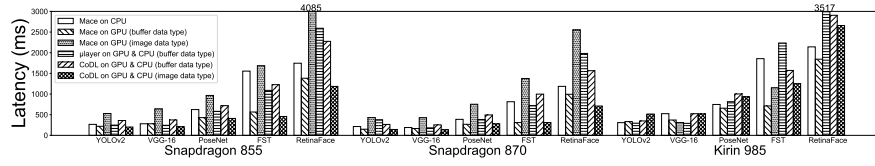


Figure 2: Latency on different mobile devices with diverse parallel inference strategies.

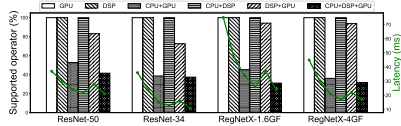


Figure 3: Operator support over diverse models.

Table 2: Performance of cross-level optimization of parallel inference, *i.e.*, CoDL, on Snapdragon 855.

Levels	Methods	Top accuracy (%)	Memory usage (MB)	Latency (ms)	Speedup (%)
Baseline	ResNet-18 [6]	75.23	47.24	213.24	
Resource-friendly	Low-rank decomposition [27]	73.73	15.61	197.53	7.37
frontend compilation	Pruning [4]	71.31	23.91	146.73	31.19
Model-adaptive	CoDL [10]	76.23	47.24	189.01	11.36
backend compilation	Operator fusion [15]	76.23	47.24	136.60	35.91
	CoDL [10]+Low-rank decomposition [27]	73.72	15.60	132.96	37.65
Cross-level combination	CoDL [10]+pruning[4]	71.31	23.89	131.66	38.35
	CoDL [10]+pruning[4]+operator fusion[15]+memory allocation [21]	73.36	11.53	103.23	48.4%

μ Layer slows down by $1.9\times \sim 3.1\times$ compared to on GPU, due to scheduling overhead.

Summary. The performance of various parallel strategies on diverse models varies. Even within the same parallel strategy, performance diverges. There’s *no one-fit-all* parallel strategy, adaptation is essential to enhance resource utilization across diverse mobile platforms and models.

3.1.2 Challenges in unsupported operators and process fallbacks. To test the presence of unsupported operators in existing parallel inference strategies, we use the Snapdragon 855 (D_1) to evaluate six inference schemes across four diverse DL models: *i.e.*, single-processor inference on GPU, DSP using TensorFlow Lite, as well as parallel inference on CPU+GPU, CPU+DSP, CPU+GPU, and CPU+GPU+DSP using NN-stretch. Figure 3 shows the results. *First*, unsupported operators are ubiquitous across diverse parallel schemes and DL models. The presence of unsupported operators triggers process fallback, causing computations to regress from specialized accelerators like GPU, DSP, *etc.* to the CPU. *Second*, the operator fallback leads to resource idleness and process waiting. In ResNet-50, the unsupported operator ratio of parallel inference on CPU and GPU reaches about 48%. We note that cross-processor parallel inference introduces new operators, decreasing the percentage of supported operators.

Summary. At the backend compilation level, considering operator support for cross-processor distribution is crucial. Also, adaptive redistribution can enhance underutilized processor resources in cases of process fallback at runtime.

3.1.3 Cross-level optimization opportunities. To evaluate the potential of cross-level optimization for enhancing parallel inference, we integrate various techniques with CoDL [10]

on ResNet-18, as shown in Table 2. *First*, optimizations at the frontend compilation can also improve resource utilization and accelerate DL inference. Low-rank decomposition reduces latency by 7.37%, leveraging smaller-rank matrices that require less memory access. Channel pruning further reduces latency by 31.19%, eliminating redundant channels. *Second*, backend compilation-level optimization, such as operator fusion and CoDL parallelism, can improve resource availability, broadening the optimization space for frontend compilation while sacrificing less accuracy to meet resource budgets. CoDL on CPU+GPU speeds up parallel inference by 11%, and operator fusion achieves a significant 35% latency reduction. *Fourth*, cross-level optimization optimizes the accuracy-efficiency tradeoff, with the combination of pruning, operator fusion, memory allocation, and CoDL achieving the optimal tradeoff and reducing latency by 48.4%.

3.2 Is parallel inference across processors always beneficial?

3.2.1 Performance of parallel strategies with diverse scheduling granularity. To test the impact of parallel scheduling granularity in real-world mobile contexts, we conducted experiments on Resnet-50 using Snapdragon 855 (D_1) with diverse parallel granularities, *i.e.*, ① sub-graph parallelism using NN-stretch across CPU+GPU+DSP, ② inter-operator parallelism using NN-stretch on CPU+GPU, ③ intra-operator parallelism using CoDL across CPU+GPU, and ④ task parallelism on CPU, GPU, and DSP using Tensorflow Lite. Additionally, we simulate dynamic contexts by introducing varying numbers of competing processes, *i.e.*, 0, 1, and 2. As shown in Figure 4, *first*, coarse-grained parallel granularity, *i.e.*, sub-graphs, exhibits optimal performance, *e.g.*, 19ms, when no competing processes, owing to lower data merging time. *Second*, fine-grained parallel granularity, *i.e.*, intra-operator, has the lowest latency when there are 2 competing processes, attributed to the low computation load per processor. However, it causes process suspensions which increase data transmission delay. With a high-competing workload, CoDL achieves the lowest latency. While in low workload, its latency increases by 48% than NN-stretch. *Third*, with a heavy workload, this impact on coarse-grained parallelism is significant. *Fourth*, allocating computations to a single processor could, in CPU high-workload cases, outperform across heterogeneous processors due to scheduling overhead.

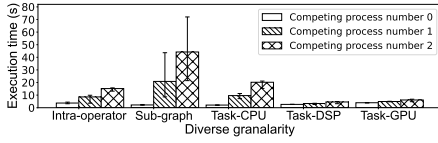


Figure 4: Impact of parallel scheduling granularity.

Table 3: Latency estimation errors on diverse devices.

Model	Methods	Estimated inference latency (ms)						Estimation error (%)			
		σ_0 (offline)	σ_1	σ_2	σ_3	σ_4	$\sigma_1 \uparrow$	$\sigma_2 \uparrow$	$\sigma_3 \uparrow$	$\sigma_4 \uparrow$	
YOLOv2	Mace [26] on GPU	219	262	320	270	372	19	46	23	70	
	CoDL [10] on CPU+GPU	170	217	277	263	315	28	63	55	85	
VGG-16	Mace [26] on GPU	240	272	312	301	464	13	30	25	93	
	CoDL [10] on CPU+GPU	171	220	307	326	387	29	80	91	126	
PoseNet	Mace [26] on GPU	430	442	503	637	768	2	17	48	79	
	CoDL [10] on CPU+GPU	328	335	395	648	725	2	20	98	121	

The scheduling time for a single GPU is $2.5 \times$ larger than the latency reduction brought by CPU-GPU parallelism. *Fifth*, CPU are more sensitive to competing than GPU or DSP. With 3 competing requests, the latency on CPU increases by $10.2 \times$, while for GPU and DSP, only increases by $2.1 \times$ and $1.8 \times$.

Summary. Mobile resource dynamics necessitates adaptive adjustment of parallel scheduling granularity.

3.2.2 Performance of parallel strategies with diverse processor utilization rates. Maximizing resource utilization for DL inference may not always lead to optimal system performance, especially in multi-process applications. We employ a user to simultaneously run DL inference on a smartphone while using a video playback app, e.g., YouTube, showcasing the performance tradeoff between DL inference and other system operations. Using Auto Clicker software, we establish a fixed screen scrolling path and speed, continuously scrolling the YouTube app to maintain a consistent UI rendering frame rate and workload. In this setup, we first obtained the frame drop rate and GPU utilization when users using video playback app without any DL model inference tasks over a 2-minute period. We then evaluate six single-/cross-processor inference strategies over a 20-minute period with VGG-16 on device D_1 . Figure 5 shows the results. *First*, excessive GPU utilization during DL inference led to significant frame drop rates in the YouTube video app refresh. Utilizing Mace on GPU (with buffer type) resulted in frame drops exceeding 90%. *Second*, compared to execution on a single GPU, parallel execution results in a lower frame drop rate. This is because parallel execution offloads a portion of the computational burden onto the CPU, reducing the load on the GPU, which decreases the UI frame drop rate.

Summary. In mobile systems, identifying optimal processor utilization levels for parallel inference is demanded to balance DL inference performance with other system tasks.

3.2.3 Offline vs. runtime latency profiler. Existing latency profilers are mainly *offline*, e.g., linear regression based on the number of computations (i.e., FLOPs) [12], complex black-box machine learning [29], and platform-aware methods (i.e., concurrency) [10]. However, the non-stationary dynamics in mobile environments widen the gap between actual

Table 4: Impact of data reuse in parallel inference.

	Frame-level reuse	Layer-level reuse	Latency per frame (ms)	Improvement	Visual quality (VMAF)	Degradation
Setups	0	0	99.8	-	49	-
	1	1	93.2	6.6% ↓	47	4.1% ↓
	2	2	86.2	13.6% ↓	42	14.3% ↓

and estimated latency. To observe the impact of resource dynamics on latency, we simulate five resource conditions $\sigma_0 \sim \sigma_4$ commonly seen in mobile systems[18]: σ_0 (offline): run DL inference solely on GPU with 30 °C, which affects dynamic voltage and frequency scaling (DVFS); σ_1 : increase the mobile GPU temperature to 60 °C; σ_2 : increase the mobile GPU temperature to 70 °C; σ_3 : let competing processes occupy the cache to tune the cache hit rate to be 20%; σ_4 : suspend half of the operators not supported by GPU to CPU for execution. As shown in Table 3, *first*, the speedup performance under dynamic resource condition $\sigma_1 \sim \sigma_4$ significantly vary from σ_0 which has abundant resources (i.e., offline settings). *Second*, both GPU inference and CPU-GPU parallel schemes experience increased latency as the temperature rises in σ_2 , due to the need for DVFS to prevent overheating. *Third*, parallel inference benefits from a higher cache hit rate, leading to lower latency than single-processor. Parallel inference distributes computation across processors, mitigating the impact of GPU computation slowdown caused by high-temperature throttling. In σ_3 , increased cache competing results in frequent cache misses and memory access, increasing latency. *Fourth*, suspending operators from GPU to slower CPUs increases computation latency, and CPU-GPU data transfer incurs higher transmission latency. In σ_4 , offloading GPU-unsupported operators to CPU results in a latency increase by 93% for Mace and 126% for CoDL.

Summary. Runtime latency profiling is desired to integrate non-stationary mobile resource dynamics.

3.2.4 Opportunities in integrating mobile DL inference regularity into backend data reuse. Mobile DL inference on continual data streams, e.g., video, follows a *predictable lifecycle*. It's beneficial to determine when previous intermediate results will be accessed for the final time, and when it's safe to release resources. We test a model on device D_1 . In Table 4, the reuse in parallel inference led to a 6.6% reduction in latency with a slight decrease, i.e., 4.1%, in visual quality (VMAF).

4 CONCLUSION

Parallel inference on heterogeneous mobile processors can boost the performance of computing-intensive DL models. However, mainstream mobile DL frameworks typically utilize one processor, limiting speedup. Despite prior efforts, the effectiveness of parallel inference in real-world mobile scenarios with diversity and dynamics is less-explored. We conduct a cross-level benchmark to assess potential and pitfalls of parallel DL inference on heterogeneous mobile processors. More insights and heuristics are needed on this topic.

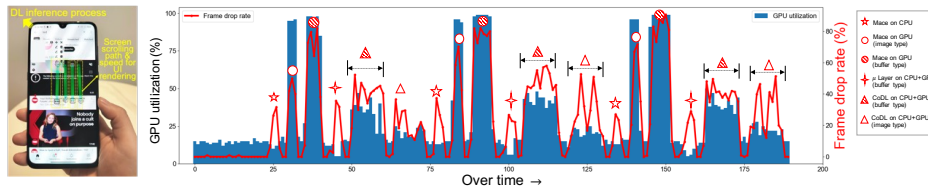


Figure 5: GPU utilization and mobile user interaction responsiveness (e.g., frame drop rate) at runtime.

ACKNOWLEDGEMENTS

This work was supported by the National Science Fund for Distinguished Young Scholars (62025205), the National Natural Science Foundation of China (No. 62032020, 62102317), and CityU APRC grant (No. 9610633).

REFERENCES

- [1] Jie An, Haoyi Xiong, Jiebo Luo, Jun Huan, and Jinwen Ma. 2019. Fast universal style transfer for artistic and photorealistic rendering. *arXiv preprint arXiv:1907.03118* (2019).
- [2] Luca Arrotta, Gabriele Civitarese, and Claudio Bettini. 2022. Dexar: Deep explainable sensor-based activity recognition in smart-home environments. *Proceedings of IMWUT* 6, 1 (2022), 1–30.
- [3] Jiankang Deng, Jia Guo, Evangelos Ververas, Irene Kotsia, and Stefanos Zafeiriou. 2020. Retinaface: Single-shot multi-level face localisation in the wild. In *Proceedings of CVPR*. 5203–5212.
- [4] Hongxiang Fan, Stylianos I Venieris, Alexandros Kouris, and Nicholas Lane. 2023. Sparse-DySta: Sparsity-Aware Dynamic and Static Scheduling for Sparse Multi-DNN Workloads. In *Proceedings of MICRO*. 353–366.
- [5] Haogang Feng, Gaoze Mu, Shida Zhong, Peichang Zhang, and Tao Yuan. 2022. Benchmark analysis of yolo performance on edge intelligence devices. *Cryptography* 6, 2 (2022), 16.
- [6] Google. 2017. TensorFlow Lite. <https://www.tensorflow.org/lite/performance/measurement..>
- [7] Yongjie Guan, Xueyu Hou, Nan Wu, Bo Han, and Tao Han. 2022. DeepMix: mobility-aware, lightweight, and hybrid 3D object detection for headsets. In *Proceedings of MobiSys*. 28–41.
- [8] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. 2016. Deep residual learning for image recognition. In *Proceedings of CVPR*. 770–778.
- [9] Joo Seong Jeong, Jingyu Lee, Donghyun Kim, Changmin Jeon, Changjin Jeong, Youngki Lee, and Byung-Gon Chun. 2022. Band: coordinated multi-dnn inference on heterogeneous mobile processors. In *Proceedings of MobiSys*. 235–247.
- [10] Fucheng Jia, Deyu Zhang, Ting Cao, Shiqi Jiang, Yunxin Liu, Ju Ren, and Yaoxue Zhang. 2022. Codl: efficient cpu-gpu co-execution for deep learning inference on mobile devices. In *Proceedings of MobiSys*. 209–221.
- [11] Youngsok Kim, Joonsung Kim, Dongju Chae, Daehyun Kim, and Jangwoo Kim. 2019. μ layer: Low latency on-device inference using cooperative single-layer acceleration and processor-friendly quantization. In *Proceedings of EuroSys*. 1–15.
- [12] Hanxiao Liu, Karen Simonyan, and Yiming Yang. 2018. Darts: Differentiable architecture search. *arXiv preprint arXiv:1806.09055* (2018).
- [13] Miaomiao Liu, Sikai Yang, Wyssanie Chomsin, and Wan Du. 2022. Real-time tracking of smartwatch orientation and location by multitask learning. In *Proceedings of SenSys*. 120–133.
- [14] Chunjie Luo, Xiwen He, Jianfeng Zhan, Lei Wang, Wanling Gao, and Jiahui Dai. 2020. Comparison and benchmarking of ai models and frameworks on mobile devices. *arXiv preprint arXiv:2005.05085* (2020).
- [15] Wei Niu, Jiexiong Guan, Yanzhi Wang, Gagan Agrawal, and Bin Ren. 2021. DNNFusion: accelerating deep neural networks execution with advanced operator fusion. In *Proceedings of TH-CPL*. 883–898.
- [16] Ilija Radosavovic, Raj Prateek Kosaraju, Ross Girshick, Kaiming He, and Piotr Dollár. 2020. Designing network design spaces. In *Proceedings of CVPR*. 10428–10436.
- [17] Joseph Redmon and Ali Farhadi. 2017. YOLO9000: better, faster, stronger. In *Proceedings of CVPR*. 7263–7271.
- [18] Haihong She, Yigui Luo, Zhaohong Xiang, Weiming Liang, and Yin Xie. 2023. Accurate Latency Prediction of Deep Learning Model Inference Under Dynamic Runtime Resource. In *Proceedings of NeurIPS*. Springer, 495–510.
- [19] Karen Simonyan and Andrew Zisserman. 2014. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556* (2014).
- [20] Yunpeng Song and Zhongmin Cai. 2022. Integrating handcrafted features with deep representations for smartphone authentication. *Proceedings of IMWUT* 6, 1 (2022), 1–27.
- [21] Arne Symons, Linyan Mei, and Marian Verhelst. 2021. LOMA: Fast auto-scheduling on DNN accelerators through loop-order-based memory allocation. In *Proceedings of AICAS*. IEEE, 1–4.
- [22] Tencent. 2018. ncnn. <https://github.com/Tencent/ncnn>.
- [23] Xudong Wang, Li Lina Zhang, Jiahang Xu, Quanlu Zhang, Yujing Wang, Yuqing Yang, Ningxin Zheng, Ting Cao, and Mao Yang. 2023. SpaceEvo: Hardware-Friendly Search Space Design for Efficient INT8 Inference. In *Proceedings of CVPR*. 5819–5828.
- [24] Yuxin Wang, Qiang Wang, Shaohuai Shi, Xin He, Zhenheng Tang, Kaiyong Zhao, and Xiaowen Chu. 2020. Benchmarking the performance and energy efficiency of AI accelerators for AI training. In *Proceedings of CCGRID*. IEEE, 744–751.
- [25] Jianyu Wei, Ting Cao, Shijie Cao, Shiqi Jiang, Shaowei Fu, Mao Yang, Yanyong Zhang, and Yunxin Liu. 2023. NN-Stretch: Automatic Neural Network Branching for Parallel Inference on Heterogeneous Multi-Processors. In *Proceedings of MobiSys*. 70–83.
- [26] XiaoMi. 2017. XiaoMi/mace. <https://github.com/XiaoMi/mace>.
- [27] Miao Yin, Yang Sui, Siyu Liao, and Bo Yuan. 2021. Towards efficient tensor decomposition-based dnn model compression with optimization framework. In *Proceedings of CVPR*. 10674–10683.
- [28] Chenyang Zhang, Feng Zhang, Kuangyu Chen, Mingjun Chen, Bingsheng He, and Xiaoyong Du. 2023. EdgeNN: Efficient Neural Network Inference for CPU-GPU Integrated Edge Devices. In *Proceedings of ICDE*. IEEE, 1193–1207.
- [29] Li Lina Zhang, Shihao Han, Jianyu Wei, Ningxin Zheng, Ting Cao, Yuqing Yang, and Yunxin Liu. 2021. NN-meter: Towards accurate latency prediction of deep-learning model inference on diverse edge devices. In *Proceedings of MobiSys*. 81–93.
- [30] Qiyang Zhang, Xiang Li, Xiangying Che, Xiao Ma, Ao Zhou, Mengwei Xu, Shangguang Wang, Yun Ma, and Xuanzhe Liu. 2022. A comprehensive benchmark of deep learning libraries on mobile devices. In *Proceedings of WWW*. 3298–3307.
- [31] Christian Zimmermann and Thomas Brox. 2017. Learning to estimate 3d hand pose from single rgb images. In *Proceedings CVPR*. 4903–4911.