

Yash Agarwal

yash-agarwal.org | yashagar@cs.cmu.edu | linkedin.com/in/yashagarwalcs | 412-519-4930

EDUCATION

- Carnegie Mellon University, School of Computer Science** **Pittsburgh, PA**
Masters in Intelligent Information **Systems** - GPA - 3.9/4.0 2025 - 2026
 - Coursework - *Parallel Computation, Computer Systems, DL Systems, LLM Systems, Deep Learning, Multimodal ML*
 - Advisor - Prof. **Tianqi Chen** (NVIDIA, creator of TVM, XGBoost)
- Indian Institute of Technology (IIT), Kharagpur** **Kharagpur, India**
Bachelors in **Computer Science** and Engineering - GPA 8.65/10 2019 - 2023
 - **All India Rank 250** out of 1.2 million applicants for Joint Entrance Examination (JEE Advanced)
 - Publication - YesBut: Evaluating Vision-Language Models **EMNLP 2024**

WORK EXPERIENCE

- Modular Inc** **San Francisco, CA**
AI Performance Engineering Intern – MAX Inference Engine Summer 2026
 - Shipped **DSpark speculative decoding** for **Gemma 4 12B/31B** reaching **519 tok/s**, block-parallel drafter on target **hidden states**, **NVFP4** targets, **FP8 KV cache**, **CUDA graph capture** - **4x** faster decode; beat **vLLM TPOT** by 28%
 - Implemented **BLASST dynamic block-sparse attention** in the warp-specialized **FlashAttention-4** kernel on **Blackwell B200**, skipping near-zero KV blocks mid-online **softmax** - attention up to **42% faster**, **TTFT -20%** at preserved accuracy
 - Fixed a serial softmax stall in the **FA4** kernel with **Nsight Compute** and per-warp **clock64** gantt timelines; shipped **cross-stage-P TMem** scheduling with **QK-first MMA** ordering, freeing score buffers early - **3.4%** faster kernel
 - Accelerated **pinned host-memory** allocation **10x** (50mins → 5mins) for tiered **KV-cache offload** via lock-free **NUMA-pinned** parallel page faulting at 99 GB/s pipelined into GPU registration, **18% faster than CUDA's VMM API**
 - Re-architected serving to batch prompt chunks into decode steps' idle FLOPs with speculation live, replacing exclusive **prefill** steps that froze every user 13–17 ms - prompt cost **53 → 16 μs/token**, per-token latency up to **12.9%** lower
 - Built **model-free speculative decoding** from a **suffix tree** over generated tokens, tree ops injected as **CUDA host callbacks** inside the decode graph to preserve CPU-GPU **async scheduling** - **TPOT -21%** on code, **-11%** on **Spec-Bench**
- Squarepoint Capital** **Bangalore, India**
Software Developer **C++** 2023 - 2025
 - Architected firm's first **FX Options** trading stack, built entire system in **C++20** and Python to pioneer new asset class
 - Debugged critical production crashes using **gdb** core dump analysis and implemented **thread-safety** in memory access
 - Redesigned **multithreaded** Python testing framework to single-threaded architecture eliminating **40%** test flakes

PROJECTS @ CARNEGIE MELLON

- Needle** - Optimizing CUDA GPU Kernels and Autodiff | DL Systems Course Project **Fall 2025**
 - Optimized Matmul **CUDA kernel** by tiling, vectorization and GMEM coalescing, achieving 85% performance of **cuBLAS**
 - Implemented **Flash Attention**, including kernel fusion, tiling, online softmax, achieving 77% performance of PyTorch
 - Architected automatic differentiation Tensor, complete with mathematical ops and nn.Modules for Conv, RNN, MHA
 - Developed NDArray class supporting reshaping, broadcasting, summation, Matmul using **C++ and CUDA** backend
- Malloc, cache, shell** | Computer Systems Course Projects **Spring 2026**
 - Implemented **malloc** using **mmap** and segregated free-lists, optimizing using **Better Fit**, **miniblocks**, removing footers
 - Implemented a **cache** simulator, optimized matrix transpose by tiling and eliminating **evictions**, 4th in a class of 200
 - Parallelized a File System with **RW locks**, enabling **concurrent** reads and writes to shared inodes, blocks, directories
- WhisQer** - Quantization trained Conformer | Deep Learning Course Project **Fall 2025**
 - Architected **quantized** Conformer ASR with 2-bit and 1-bit precision variants, with learnable scaling parameters, compressing a 108MB FP32 model to 45.6MB 2-bit by quantizing encoder FFN and attention projections
 - Implemented teacher student co-training combining CTC and KL divergence loss and straight through estimator

SKILLS

Languages - C++ 23, Python, Rust, Bash | **ML** - PyTorch, JAX, Numpy, CUDA, Triton, Mojo