

بسمه تعالی

درخواست طرح پیشنهادی (RFP)

پژوهش، طراحی، نمونه سازی، تست و تدوین دانش فنی تراشه

شتاب دهنده CNN مبتنی بر MAC Array

با فناوری ۴۰ نانومتر Mixed Signal

تیر ماه ۱۴۰۵

۱- مقدمه

در سال‌های اخیر با رشد بی‌سابقه نیاز به پردازش هوش مصنوعی، پردازش ابری، شبکه‌های پرسرعت و سامانه‌های ذخیره سازی داده، طراحی شتاب‌دهنده سخت‌افزاری به یکی از ارکان اصلی معماری پردازش‌های پیشرفته تبدیل شده است. در حالی که مدل‌های زبانی بزرگ (LLMها) بیشتر به معماری transformer متکی هستند، شبکه‌های CNN (Convolutional Neural Network) ستون اصلی پردازش تصویر، بینایی ماشین و تشخیص ویدئو (مانند الگوریتم‌های YOLO یا ResNet) به شمار می‌روند. شتاب‌دهنده‌هایی که برای اجرای الگوریتم‌های CNN طراحی می‌شود، به دلیل ساختار تکراری کانولوشن‌ها و نیاز به بازاستفاده مداوم از داده‌ها، از جریان‌های داده اختصاصی (مانند Weight-Stationary یا Row-Stationary) برای حذف گلوگاه‌های حافظه استفاده می‌کنند. یک شتاب‌دهنده تخصصی CNN دارای بخش‌های زیر است:

- بافرهای خطی (Line Buffers): رجیسترهای سخت‌افزاری که نوارهای افقی کوچکی از تصویر را نگه می‌دارند تا هنگام حرکت پنجره کانولوشن روی تصویر، نیازی به بارگذاری مجدد کل فریم نباشد؛
- واحدهای Pooling سخت‌افزاری: بلوک‌های منطقی که عملیات Max-Pooling یا Average-Pooling را فوراً و بدون هدر دادن چرخه‌های کلاک پردازنده اصلی انجام می‌دهند؛
- مکانیزم حذف صفرها (Sparsity Engines): شبکه‌های CNN پس از لایه ReLU حاوی صفرهای زیادی هستند. این کنترلرها با تشخیص صفر، عملیات ریاضی را کاملاً نادیده می‌گیرند که این کار سرعت را دو برابر کرده و مصرف توان را نصف می‌کند.

از جمله مهم‌ترین شتاب‌دهنده‌های طراحی شده برای CNNها می‌توان به موارد زیر اشاره کرد:

- ۱- میکروکنترلرهای کم‌مصرف جا سازی شده برای استنتاج لبه (Edge IoT): این تراشه‌ها برای دوربین‌های هو شمند "همیشه روشن" و تجهیزات پوشیدنی طراحی شده‌اند تا مدل‌های بینایی ماشین را با یک باتری کوچک اجرا کنند:

- تراشه‌های MAX7800 / MAX7802 (ساخت Analog Devices): این میکروکنترلرها دارای یک شتاب‌دهنده سخت‌افزاری اختصاصی CNN با ۶۴ پردازنده موازی هستند. این لایه سخت‌افزاری، حلقه‌های درهم‌تنیده پردازشی (Nested Loops) را به صورت موازی اجرا کرده و با نگه داشتن تمام وزن‌ها روی خود تراشه، مصرف توان را به چند میلی‌وات کاهش می‌دهد.

۲- پردازنده‌های جاسازی شده بینایی در لبه (VPUs / NPUs): این کمک پردازنده‌ها (Co-processors) به کامپیوترهای لبه یا سیستم‌های صنعتی متصل می‌شوند تا فید چند دوربین را به صورت هم‌زمان و بی‌درنگ پردازش کنند:

- Hailo-۱۰ / Hailo-۸: این شتاب‌دهنده ساختاری منحصربه‌فرد دارد که در آن حافظه و عناصر مسیریابی مستقیماً در کنار عناصر پردازشی توزیع شده‌اند. این چیدمان، ساختار لایه‌به‌لایه‌ی شبکه‌های CNN را شبیه‌سازی کرده و نرخ فریم بسیار بالایی را با کمترین توان مصرفی ارائه می‌دهد؛
- Intel Movidius / Keem Bay: یک واحد پردازش بینایی (VPU) اختصاصی است که به خطوط استریم سخت‌افزاری مجهز شده تا عملیات کانولوشن دوبعدی را با سرعت بالا روی تصاویر اعمال کند.

۳- مقیاس دیتاسنتر و ابری:

- Google TPU (نسخه‌های اولیه): اگرچه نسخه‌های مدرن TPU گوگل اکنون برای مدل‌های زبانی (LLM) بزرگ شده‌اند، اما معماری نسخه‌های ۱ تا ۳ آن‌ها کاملاً بر پایه آرایه‌های سیستمیک (Systolic Array) برای شتاب‌دهی به جبر خطی و کانولوشن‌های CNN طراحی شده بود. در این ساختار، داده‌ها در یک شبکه دوبعدی از بلوک‌های ضرب-جمع (MAC) بدون نیاز به نوشتن و خواندن مداوم در حافظه خارجی به حرکت درمی‌آیند.

۴- بلوک‌های سخت‌افزاری قابل برنامه‌نویسی برای FPGA:

- AMD Xilinx DPU (Deep Learning Processor Unit): یک هسته IP قابل پیکربندی است که به طور ویژه برای CNN بهینه‌سازی شده و عملیات زمان‌بندی، ضرب ماتریسی موازی و لایه‌های Pooling را مستقیماً در منطق برنامه‌نویسی FPGA انجام می‌دهد؛
- Eyeriss (دانشگاه MIT): یک معماری پایه‌ای برای تراشه‌های مدرن CNN است که جریان داده Row-Stationary را معرفی کرد. این روش بازاستفاده از پیکسل‌ها و وزن‌ها را درون رجیسترهای محلی به حداکثر می‌رساند. کدهای متن‌باز VHDL/Verilog این پروژه برای کارهای تحقیقاتی در دسترس است.

۲- هدف فراخوان

هدف از این فراخوان، شناسایی و انتخاب تیم مناسب برای پژوهش، طراحی، نمونه سازی، تست و تدوین دانش فنی یک شتاب دهنده هوش مصنوعی سبک (Lightweight AI Accelerator) برای inference با تمرکز بر اجرای کارآمد شبکه های عصبی کانولوشنی (CNN) برای کاربردهای: بینایی ماشین (Object Detection, Classification)، رباتیک (tracking, navigation)، سیستم های نظارتی (surveillance AI)، IoT هوشمند، پردازش تصویر صنعتی و edge AI camera modules می باشد. الزامات عملکردی برای مدل های هدف که این پردازنده باید پشتیبانی کند شامل موارد زیر می باشد:

- MobileNetV1 INT8: حداقل ۳۰ FPS برای ورودی ۲۲۴×۲۲۴؛
- Tiny-YOLO: حداقل ۱۰ FPS برای ورودی ۴۱۶×۴۱۶؛
- ResNet-18 (quantized): latency کمتر از ۵۰ ms
- MAC array utilization $\geq 60\%$ under defined benchmark workloads (ResNet-18) measured over full inference execution including memory stalls.

شرایط آزمون:

Batch size = ۱

Input data from off-chip DRAM

تمامی اندازه گیری ها باید شامل data movement overhead باشند.

برای توسعه های آتی و نسخه های بعدی این پردازنده، لازم است که از معماری برپایه ریز تراشه (chiplet) استفاده شود تا اولاً بتوان از یک بلوک / ریز تراشه در چندین تراشه استفاده کرد و بدین ترتیب تیراژ ساخت را افزایش داد تا هزینه ساخت قابل توجیه تر باشد. همچنین امکان به روز رسانی یا توسعه بخشی از تراشه بدون طراحی و ساخت مجدد کل آن میسر باشد. علاوه بر این کارایی تراشه های مبتنی بر chiplet معمولاً بالاتر و بازده تولید (yield) آنها نیز بیشتر است و بازاریابی و فروش آنها راحت تر می باشد. برای نیل به این هدف در نسخه های آتی، ضروری است که این تراشه در این نسخه، هر چند بصورت chiplet، طراحی و ساخته نمی شود، لیکن بصورت chiplet ready طراحی شود تا تبدیل آن به chiplet در نسخه های آتی سهل تر باشد. الزامات معماری chiplet ready ذیل در فصول آتی به استحضار می رسد. طراحی هر گونه لایه فیزیکی خارجی (PHY)، از جمله DDR PHY یا PCIe PHY یا ... در دامنه این طرح قرار نمی گیرند و تبادل داده با حافظه خارجی برعهده CPU/SoC میزبان خواهد بود.

۳- مدل اجرای طرح

طرح باید در فرآیند mixed signal نود فناوری ۴۰ نانومتر اجرا شود. در همین راستا، فایل‌ها و اسناد لازم شامل PDK فناوری مذکور، توسط مرکز ملی پیشران طراحی تراشه در اختیار مجری قرار خواهد گرفت و پشتیبانی‌های فنی و نرم‌افزاری از تیم مجری به عمل خواهد آمد. پس از اتمام طراحی، فایل‌های GDSII مربوطه توسط مرکز ملی پیشران طراحی تراشه و با حمایت برنامه ملی میکروالکترونیک، برای ساخت (Tape Out) به کارخانه تولید کننده تراشه (Foundry) ارسال خواهد شد و مجری نیازی به تعامل شخصی با Foundry نخواهد داشت. پس از ساخت تراشه، عملیات باندینگ توسط مرکز ملی پیشران طراحی تراشه و در نهایت تست و ارزیابی آن مشترکاً توسط مجری و مرکز ملی پیشران طراحی تراشه انجام خواهد گردید.

۴- محدوده و دامنه طرح‌های پیشنهادی

معماری سیستم علاوه بر این تراشه شامل یک CPU / SoC host مستقل است که وظیفه مدیریت حافظه خارجی، بارگذاری مدل‌ها، زمانبندی تسک‌ها و ارتباط کلی با Accelerator را بر عهده دارد. Accelerator بصورت یک واحد محاسبات تخصصی عمل کرده و از طریق یک رابط استاندارد با میزبان ارتباط برقرار خواهد کرد. مساحت مورد نیاز برای طراحی ترجیحاً باید کمتر از ۴ میلی‌متر مربع و الزاماً کمتر از ۹ میلی‌متر مربع باشد و طرح‌های پیشنهادی باید موارد زیر را دربر گیرند:

طراحی RTL کامل

توسعه نرم‌افزار پایه

مشخصات کلی تراشه

- Technology Node: ۴۰nm
- Frequency Target: ۲۰۰-۴۰۰ MHz
- Compute Precision: INT^۸ (primary)
- Accumulator: INT^{۳۲}
- Peak Throughput: ~۰.۵-۱ TOPS
- On-chip SRAM: ۱۲۸-۲۵۶ KB
- Power Budget: <۱W
- Die Size Target: ۲-۵ mm²

معماری داده (Dataflow)

- Row-Stationary (ترجیحی)؛
- Weight-Stationary (در صورت توجیه PPA).

ارائه تحلیل مقایسه‌ای (energy vs reuse)

مشخصات بلوک‌های سخت‌افزاری تراشه

- MAC Array (Core Block)
 - Type: Systolic Array
 - Size: 32×32
 - Precision: $INT^8 \times INT^8$ to INT^{32}
 - Features:
 - pipeline
 - clock gating
 - configurable tiling
 - local reuse support
 - Additional Requirements:
 - systolic dataflow must be explicitly defined (input/weight/output flow direction)
 - pipeline depth per MAC stage must be determined
 - support for partial sum accumulation across tiles
 - possibility of stall-free streaming if data is provided
 - inter-PE communication latency ≤ 1 cycle
 - The memory subsystem must support sufficient bandwidth to avoid MAC starvation.
 - Designers must provide bandwidth analysis demonstrating no underutilization due to memory bottlenecks.
- Buffer Subsystem
 - Buffers:
 - Input Buffer
 - Weight Buffer
 - Output Buffer
 - Features:
 - double buffering (ping-pong)
 - banked SRAM
 - parallel access
- SRAM
 - SRAM should be implemented as multi-bank single-port or dual-port memory.
 - logical multi-port access must be achieved via banking and arbitration.
 - bank conflict handling must be determined.

Total: ۱۲۸-۲۵۶ KB

Type: multi-bank SRAM

SRAM latency target: ≤ 4 cycles, design must specify read/write latency and banking strategy explicitly.

Support:

- simultaneous read/write
- low latency

- DMA Engine

Channels: ≥ 2

Features:

- burst transfer (AXI4)
- address generation
- configurable stride
- interrupt (optional)
- support for memory-mapped and streaming modes
- descriptor-based operation (optional but preferred)
- latency hiding via prefetch
- back-pressure handling from compute unit

- AXI Interface

AXI4 (data path)

AXI-Lite (control)

Features:

- burst support
- high throughput
- AXI Interface Requirements:
- AXI4 data width: حداقل ۶۴ یا ۱۲۸ بیت
- burst length: حداقل ۱۶
- target bandwidth: ۲-۴ GB/s حداقل
- QoS considerations to prevent starvation

- Control Unit

Type: FSM + Loop Engine

A clear programming model must be defined:

- Instruction-based
- Register-based
- Microcode-based

The choice must be justified in terms of flexibility vs hardware overhead.

Capabilities:

- nested loops (ND):
 - ✓ channel in/out
 - ✓ spatial
 - ✓ kernel
- programmable via registers
- Control Unit must support:

- ✓ programmable instruction sequence یا microcode
- ✓ loop nesting تا حداقل ۶ سطح
- ✓ runtime reconfiguration without reset
- Activation Unit
 - ReLU (mandatory)
 - optional: clamp / leaky ReLU
- Quantization Unit
 - $INT^{32} \rightarrow INT^8$
 - scaling + shift
 - support for per-channel quantization (ترجیحی)
 - rounding mode must be determined (nearest / stochastic)
 - overflow handling: saturation (mandatory)
- PMU (Performance Monitoring Unit)
 - cycle counter
 - MAC utilization
 - stall tracking
 - bandwidth usage
- DFT
 - scan chain
 - MBIST for SRAM
 - test mode control

هر یک از بخش‌های فوق باید با ارائه deliverable‌های مشخص (مطابق بخش ۷) و گزارش‌های قابل ارزیابی تکمیل شوند. صرف انجام فعالیت بدون ارائه مستندات و نتایج قابل اندازه‌گیری قابل قبول نیست.

۵- قابلیت چیپلتی داشتن (Chiplet-Ready Requirements)

به منظور قابل تقسیم بودن تراشه طراحی شده به چند چیپلت، در طرح‌های پیشنهادی رعایت موارد زیر ضروری است:

شرایط طراحی

- partition-based design
- clean interface boundary
- decoupled subsystems

راهبرد ارتباطی

Internal:

- AXI^۴
- AXI-Stream (optional)

Future-ready:

- Transport Layer (AXI to packet abstraction)
- Adapter placeholder for UCI PHY
- latency budget for internal interface should be determined
- packetization boundary definition (transaction layer abstraction)
- clock domain crossing strategy must be defined
- die-to-die can be written to UCI protocol stack (PHY-independent design)
- interface should be designed in a manner that without changing the compute logic in the next versions we can add:
 - CHI adapter;
 - UCI PHY.

تقسیم‌بندی پیشنهادی

Chiplet A: Compute (MAC + SRAM);
Chiplet B: IO + DMA + Host Interface.

۶- نرم‌افزارهای نیازمند توسعه

راه‌انداز (driver)

Language: C

Features:

- register configuration;
- DMA control;
- execution control.

کتابخانه زمان اجرا (runtime library)

APIs:

- conv^۲d();
- load_weights();
- run_layer().

مبدل مدل (model convertor)

Language: Python

Quantization accuracy loss must be reported: Top-۱ accuracy degradation $\leq 3\%$ نسبت به مدل

FP۳۲

Functions:

- PyTorch $\rightarrow$ INT^۸
- weight extraction
- config generation

Model Converter must support:

- ONNX input (mandatory)
- PyTorch export
- graph optimization (layer fusion)
- memory scheduling (tiling generation)

فرم افزار آزمون و ارزیابی (test and evaluation software)

- golden model (NumPy);
- bit-accurate comparison;
- validation scripts.

۷- شرایط و الزامات تست، ارزیابی و تأیید

تأیید RTL (RTL Verification)

- unit-level test bench;
- system-level simulation;
- random + directed tests.

تأیید عملکردی (Functional Validation)

- CNN inference correctness;
- comparison with PyTorch.

تأیید کارایی (performance validation)

- throughput measurement;
- utilization analysis.

تأیید بعد از ساخت (post silicon validation)

- bring-up:
 - JTAG debug interface (mandatory);
 - Possibility of complete register read back;
 - internal signal observability (via scan یا debug mux).
- hardware testing;
- real model execution.

معیارهای موفقیت این طرح عبارت خواهند بود از:

- اجرای صحیح CNN؛
- دستیابی به throughput هدف؛
- tape-out موفق؛

- silicon functional:
 - Functional Coverage $\geq 80\%$
 - Code Coverage $\geq 80\%$
 - assertion-based verification
 - formal verification for vital blocks (optional but preferred)

۸- خروجی‌ها و موارد تحویلی مورد انتظار

خروجی‌های شبیه‌سازی

الف) RTL کامل و تست شده:

- کد RTL به صورت Synthesizable Verilog/System Verilog
- رعایت Coding Style استاندارد؛
- پوشش Functional Coverage حداقل 80% .

ب) محیط کامل شبیه‌سازی:

- تست‌بنچ‌ها (UVM یا معادل آن)؛
- تست‌های Latency، Throughput، Corner Case، Functional؛
- مدل‌های رفتار (Behavioral Models)؛
- اسکریپت‌های شبیه‌سازی (Questa/VCS/Xcelium).

ج) گزارش PPA در فناوری ۴۰ نانومتر:

- گزارش فرکانس، توان، سطح اشغال؛
- نتیجه سنتز از Synopsys/Cadence؛
- تحلیل Timing (STA) کامل.

فرم‌افزارها

- Driver;
- Runtime library;
- Model converter.

ملزومات نمونه‌سازی

الف) GDSII نهایی:

- فایل GDSII به همراه DRC clean و LVS clean؛

- DRC/LVS Report کامل.

ب) فایل های ماسک و Sign-off :

- LEF/DEF;
- Layer map;
- Extraction (PEX);

- Timing sign-off پس از PEX.

ج) مستندات لازم برای نمونه سازی (MPW) :

- فرم های PDK-specific ؛
- پکیج تست (Test Chip Package Plan) ؛
- لیست ریسک ها و نکات ساخت ؛

- Liberty files (timing models);
- SDF back-annotated;
- power intent (UPF/CPF);
- IR drop analysis report;
- EM analysis report.

ملزومات باندینگ

اندازه، موقعیت، گام و نقشه ی چیدمان تمام باندپدها، نقشه اتصالات (Bonding Diagram) و ... جز مواردی است که طراح باید پس از اتمام طراحی در اختیار مرکز ملی پبشران طراحی تراشه قرار دهد.

ملزومات تست و ارزیابی

پس از تحویل دای، تیم باید موارد زیر را ارائه کند:

الف) برد تست (Test Board) :

- طراحی شماتیک؛
- فایل PCB؛
- نرم افزارهای کنترل تراشه و آزمون آن؛
- Interface مناسب.

ب) برنامه تست سیلیکون:

- پارامترهای Power، Latency، Throughput؛
- Functional test با برد تست؛

- تست Corner (ولتاژ و دما) در حد امکانات آزمایشگاهی و بدون الزام به qualification صنعتی؛
- ثبت تمام باگ‌ها و نتایج تست.
- (ج) گزارش کامل Silicon Validation :
- مقایسه عملکرد واقعی با پیش‌بینی PPA ؛
- تحلیل اختلافات؛
- پیشنهاد برای Tape-out بعدی.

فایل‌های مستندسازی و تدوین دانش فنی

الف) مستند معماری:

- توصیف کامل AI accelerator ؛
- بلوک دیاگرام‌ها؛
- جریان داده و محاسبات؛
- Interface diagrams

ب) دستورالعمل استفاده و کاربری:

- نحوه راه‌اندازی؛
- ریجستر مپ؛
- ساختار دستورات؛
- محدودیت‌ها.

پ) کد متن نرم‌افزار (Source Code).

۹- زمان‌بندی اجرای طرح

برای اجرای طرح ۱۸ ماه زمان پیش‌بینی شده است:

- مدت زمان لازم برای طراحی حداکثر ۱۰ ماه:
- ✓ معماری (۱ ماه): شامل spec definition و architecture design؛
- ✓ طراحی RTL و verification (test bench and integration) (۷ ماه): شامل MAC, DMA, AXI, control؛
- ✓ طراحی فیزیکی (۲ ماه): شامل synthesis و P&R و timing closure.
- نمونه‌سازی (Tape out): ۶ ماه؛
- باندینگ، تست و ارزیابی و گزارش نویسی: ۲ ماه.

۱۰- حمایت‌ها

کممک هزینه طراحی

کممک هزینه نیروی انسانی جهت پژوهش و طراحی تا سقف ۵۰ میلیارد ریال و طبق دستورالعمل بنیاد علم ایران در سال ۱۴۰۵ به شرح زیر خواهد بود:

- هیئت علمی: ۵۰ میلیون تومان در ماه (مجموع بودجه اعضای هیئت علمی در طرح نباید بیش از ۲۰ درصد بودجه طرح باشد)؛
- دانشجوی پسا دکتري: ۳۰ میلیون تومان در ماه؛
- دانشجوی دکتري: ۱۵ میلیون تومان در ماه؛
- دانشجوی کارشناسی ارشد: ۸ میلیون تومان در ماه.

خدمات نمونه‌سازی (Tape out)

ارائه خدمات دسترسی به کارخانه ساخت تراشه (Foundry) و دریافت تکنولوژی فایل (PDK) جهت ایجاد فایل GDS II و نمونه‌سازی تراشه‌ها از طریق مرکز ملی پیشران طراحی تراشه و با حمایت برنامه ملی میکروالکترونیک معاونت علمی، فناوری و اقتصاد دانش‌بنیان خواهد بود.

خدمات باندینگ و تست

این خدمات نیز با حمایت برنامه ملی میکروالکترونیک معاونت علمی، فناوری و اقتصاد دانش‌بنیان و از طریق مرکز ملی پیشران طراحی تراشه ارائه خواهد شد.

۱۱- مالکیت معنوی

مالکیت کلیه دستاوردها و خروجی‌های فنی و علمی پروژه (شامل کدهای HDL، بلوک‌های IP، طرح فیزیکی GDSII، اسرار فنی، مستندات و ...) متعلق به طراح خواهد بود، لکن مرکز ملی پیشران طراحی تراشه اجازه بهره‌برداری غیرانحصاری، دائمی، غیرقابل فسخ و رایگان از تمامی این دستاوردها و خروجی‌ها را خواهد داشت و می‌تواند از آن‌ها در پروژه‌های دیگر استفاده نماید.