

An 88.36 TOPS/W

每瓦特执行 88.36 兆次 整数/浮点 计算

Bit level weight compressed

weight 使用 bit 级压缩 — Active-Bit-Allocate-Format

Large Language Model Accelerator

LLM 加速器

独立处理自回归模型 GPT, LLaMA 的高内存计算需求

with cluster Aligned INT-FP-GEMM

整数、浮点混合计算 加速办法 — CAMP-PE, exponent 使用 activation 分解 转为整数计算

and Bi-Dimensional Workflow Reformulation

双维度重构 线性/非线性混合 — Bi-dimensional MAC-SFU Reformulation

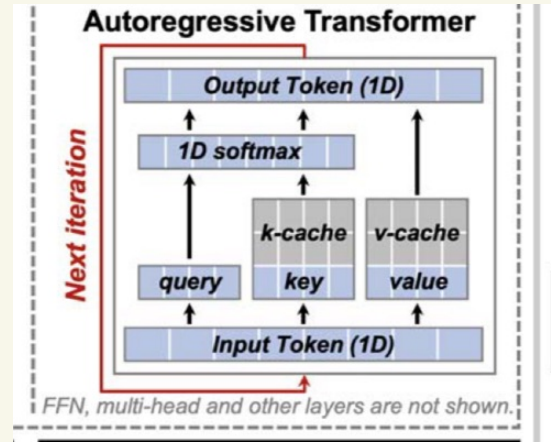
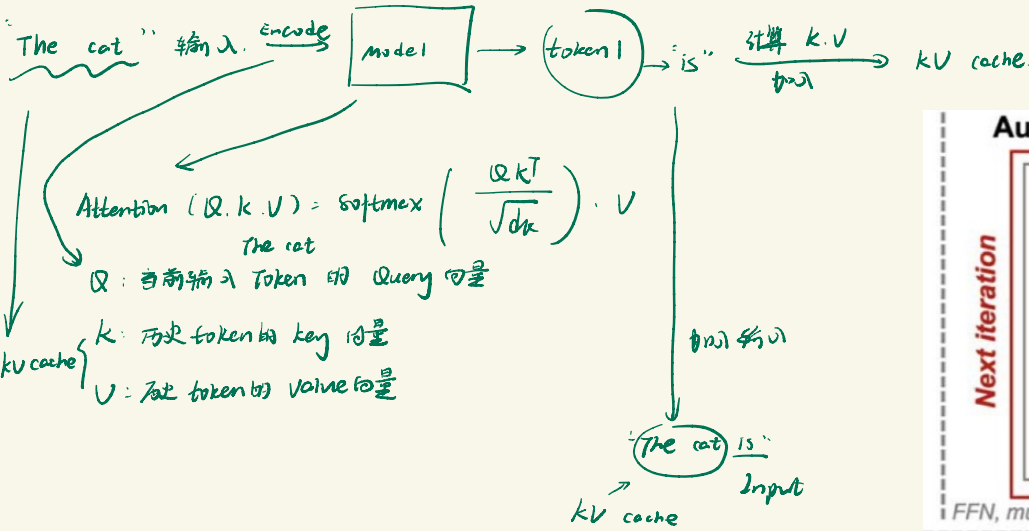
非线性计算 $\xrightarrow{\text{Taylor 展开}}$ MAC Op.

auto regressive: 生成与之前的生成有关

transformer $\rightarrow$ self Attention. 理解上 72

inference: auto regressive 方式 g. token. 基于历史生成记录 生成新 token

KV cache 的作用:



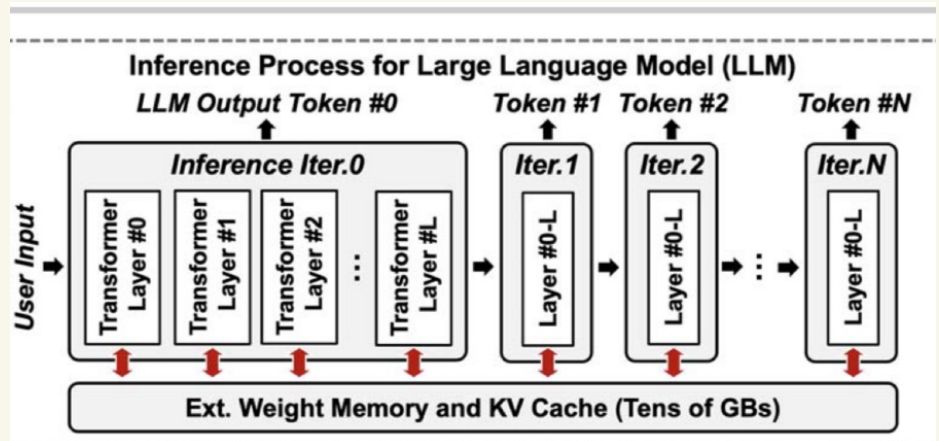
对于 Transformer, KV cache 会存储每层的 KV

inference: current token $\rightarrow$ Query

all layers' K, V $\rightarrow$ K, V

attention compute.

new token store into KV cache.



Mixed precision - General Matrix Multiplication

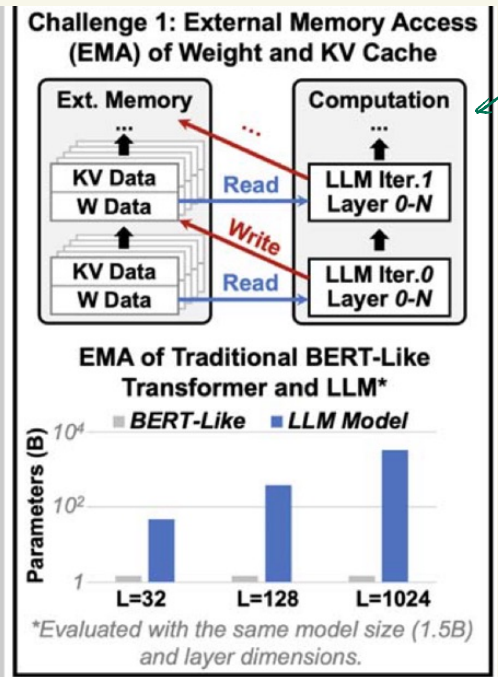
accuracy → weight → INT
activation → FP ?

weight → FP , → Energy Consumption 4.6x.

Special Function Unit.

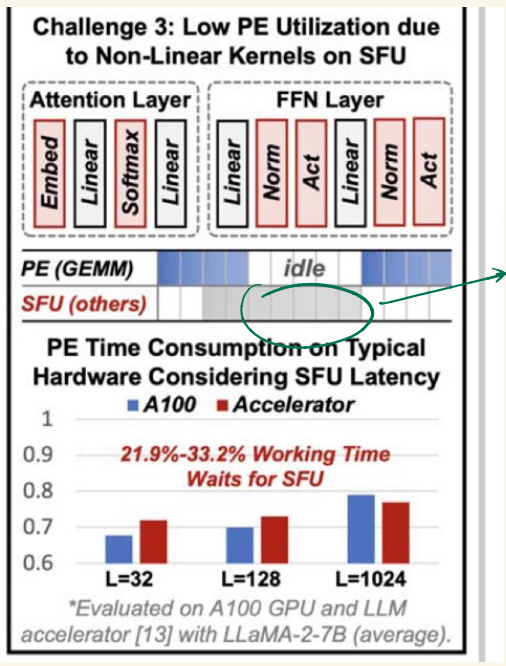
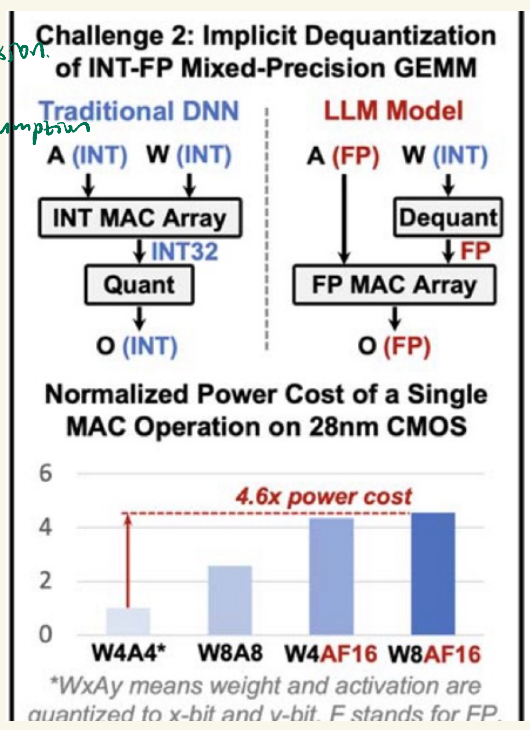
非线性函数计算慢, GEMM要等待 SFU 计算. 这时间浪费

→ EMA 过高, 混合精度难实现. 非线性运算限制



auto regressive
KV cache.
↓
EMA

mix precision.
energy consumption
complex

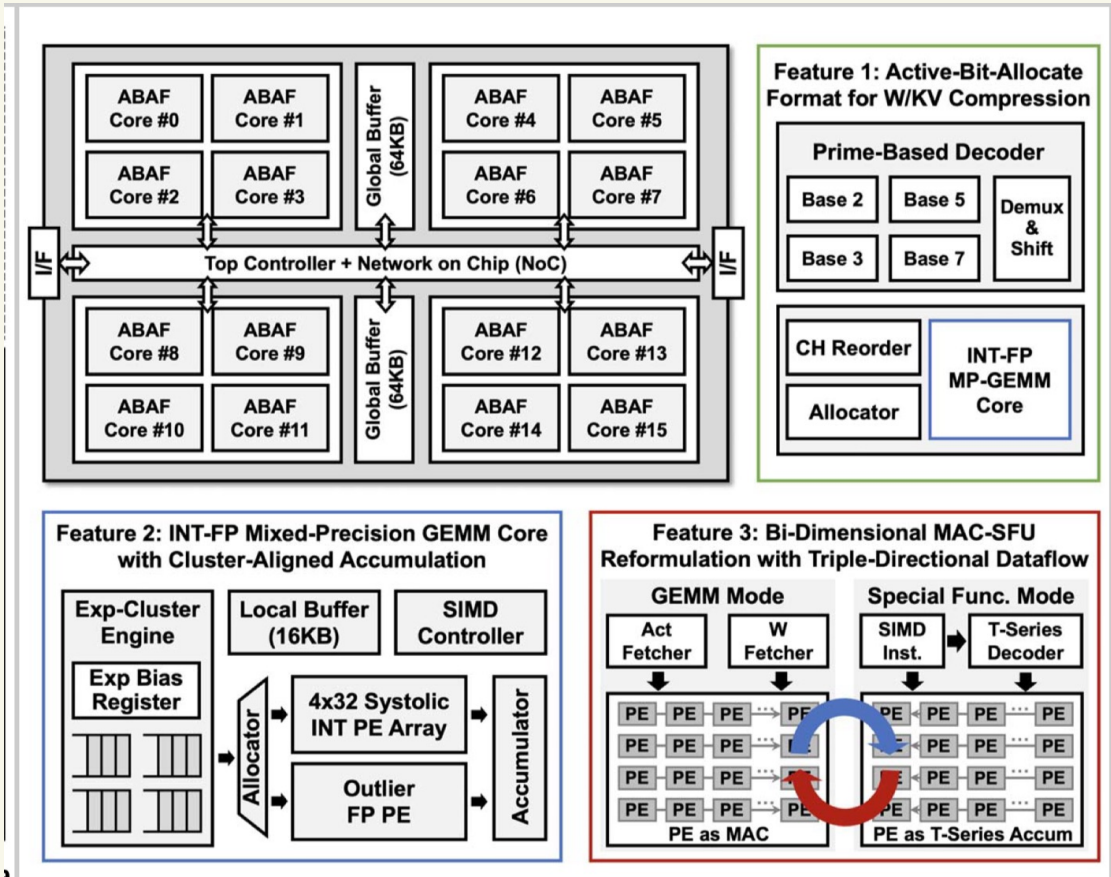


high latency

→ Accelerator: ① Active-Bit Allocate Format. ABAF
 reduce EMA for unimportant kv-cache / weight Fetch. / 95% weight is unimp.
 approach: restrict: 11-bits less than 3, + use low-width combination number.
 codec-free.
 reduce parameter size.
 带宽压力↓.

② Cluster-Aligned Mixed-precision - Processing Element
 CAMP-PE 混合精度聚类处理
 split operands into clusters by exponent.
 turn into INT GEMM. low energy consumption

③ Bi-dimensional MAC-SFU Reformulation.
 horner + taylor → nonlinear → MAC Op.
 PE wast ↓
 utilization ↑
 data transportation between linear and unlinear Unit.



network on chip.
 用于芯片内部通信的结构化网络架构
 取代 bus 或点对点连接
 广泛用于 微处理器 阵列器
 SoC, FPGA
 NoC:
 Node: 每个拉/存储/处理单元接入 NoC 为
 Route/Switch: 每个 node 有个, 决定 packet 路径
 Link: Node 之间的链路
 packet: 数据打包为小块传输
 拓扑结构: Mesh, Torus, Ring, Tree, Crossbar

architecture of all chip:
 ABAF x16.

Global Buffer 128KB.

Network on chip + top Controller.

ABAF 包含 16 个 tile.

① 使用 ABAF 压缩权重与 KV. 通过基于质数的电路实现无需解解码的解压缩.

prime based decoder

CH recoder + allocator $\xrightarrow{\text{output}}$ MP-GEMM

② features: exponent cluster 指数对齐引擎 将有相似指数的 activations 合并

1 个 CAMP-PE local Buffer

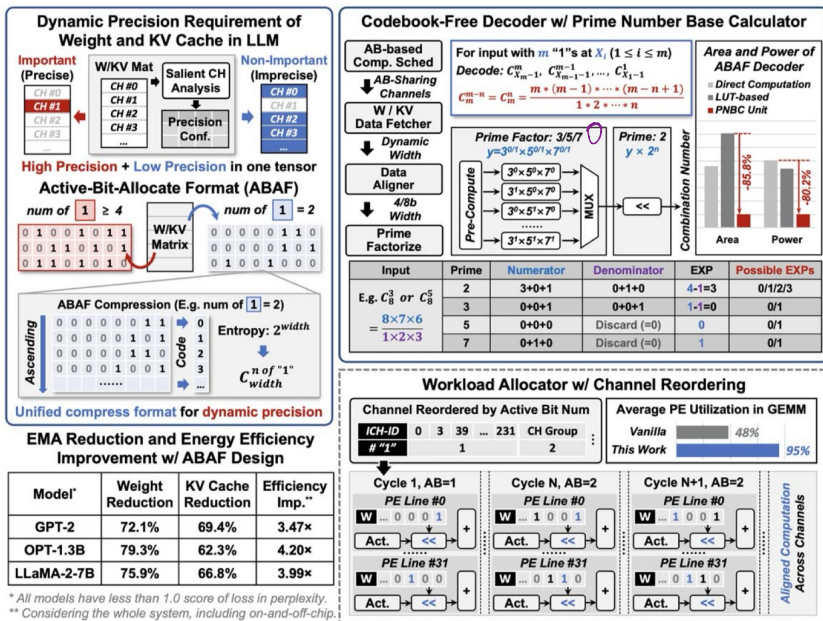
SIMD Controller

4x32 INT PE.

Outlier accumulator 异常值累加器

每个 bit serial MAC 逐位串行相加.

③ BMSR 覆盖整个核心. 将非线性计算 $\rightarrow$ GEMM.



ABAF 具体设计: 为 1 的称为 active bits.

在不重要的通道中, 只能有 1-3 个 ABAFs.

数字由索引以及更多的 ABAFs.

降低信息熵

$$C_m^n = C_{m-1}^{n-1} + C_{m-1}^n$$

① 使用 Prime Number Based Calculator 解码.

② 使用 Codebook 用传统方法

对于 8bit, 只有包含质数 2, 3, 5, 7

2 乘以 3 次, 其余最高 1 次

用 multiplier 迭出结果 + shifter 移位最终结果

用于快速还原原始组合

Work flow:

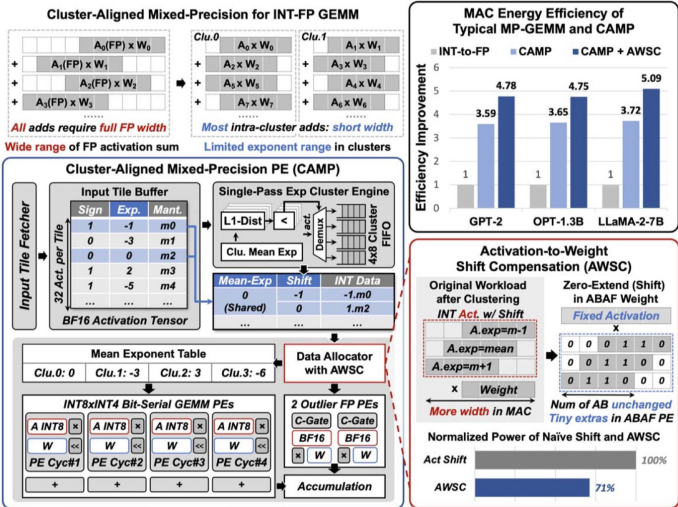


Figure 23.8.4: Cluster-aligned mixed-precision PE (CAMP) design for efficient integer-float MP- GEMM implementation.

work flow

CAMP-PE 设计:

将 activation 和权重 ABAF 构建 K/V 值 根据 L1 距离分为 4 组
分别存储在 4 个深度为 8 的 FIFO 中。 每组的指数差异小于 3
转为整数操作。

指数对齐? → MAC 宽度增加。

activation-to-weight shift compensation.

activation → zero-extension.

没有增加计算负担。因为 ABAF 不变。

L1 distance: $\equiv |x_i - x_j|$

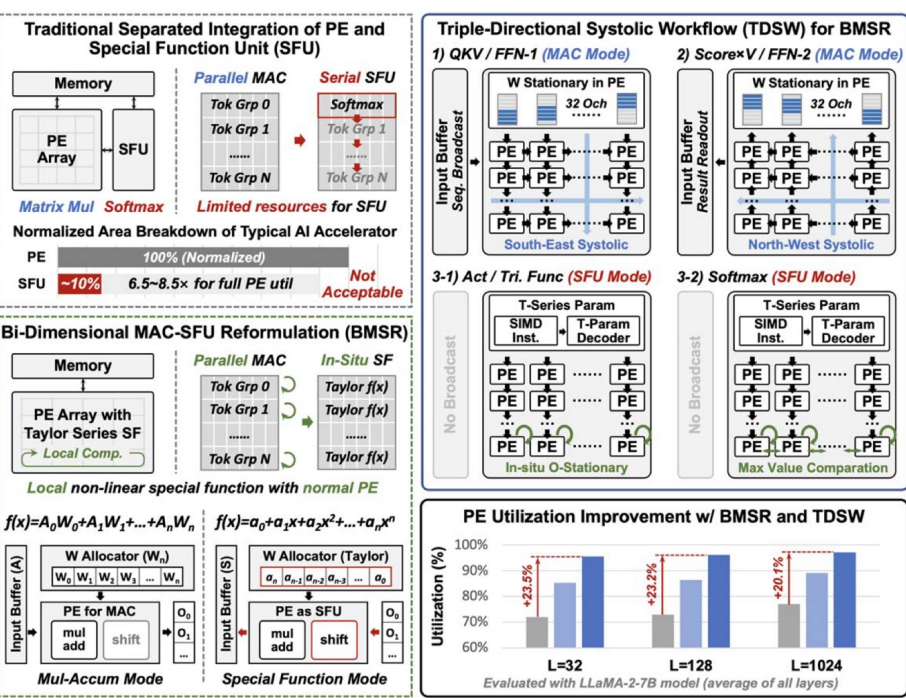


Figure 23.8.5: Bi-dimensional MAC-SFU reformulation design and triple-directional systolic workflow for high-utilization nonlinear function.

work flow

BMSR 设计:

taylor 展开的计算形式与 GEMM 基本一致。
只需要个 2 选 1

可以非线性计算由 PE 执行不用 SFU

LLM 有多种计算核 → 数据移动多

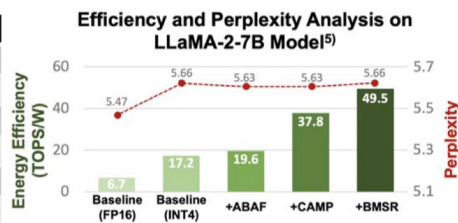
Triple-Directional-Systolic-workflow

数据传输方向 } in-situ accumulation.
south east
north west

PE 间插入比较器

TDSW: 资料搬移开销

Model	GPT-2-Small	OPT-1.3B	LLaMA-2-7B
Dataset	Wikitext-2 Language Modeling		
Accuracy and Loss ¹⁾	30.85 (+0.90)	15.50 (+0.88)	5.66 (+0.19)
Precision ²⁾	A: BF16, W:INT8 with ABAF		
Latency (ms)	40.3 (6.72) ³⁾	419.4 (68.9)	2139 (356.5)
Efficiency (uJ/Token) ⁴⁾	3.11	36.2	194.1
Energy Saving	4.00x	4.62x	4.32x



1) Perplexity on Wikitext-2 dataset, lower is better. Accuracy baseline is based on FP16 weight precision. 2) ABAF is used for weight dynamic compression. The best compression results of weight within 1.0 score of perplexity degradation limitation are used. 3) Evaluated with 6 chips working in parallel for actual working scenario, which has an total area of 21.12 mm², similar to SOTA LLM processor [23,24]. 4) System level evaluation with EMA included. 5) Evaluated on 1.0V, 460MHz, averaged on the generation process of 1024 tokens. Considering only on-chip performance and energy cost.

	ISSCC'22 [13]	VLSI'22 [20]	ISSCC'23 [22]	ISSCC'23 [17]	ISSCC'24 [23]	This Work
LLM Acceleration	NO	NO	NO	NO	Specific Model ⁵⁾	General LLM Mem + Comp
Technique (nm)	28	5	12	28	28	28
Die Area (mm ²)	6.82	0.153	4.60	3.93	20.25	3.52
Supply Voltage (V)	0.56-1.1	0.46-1.05	0.62-1.0	0.64-1.03	0.7-1.1	0.63-1.0
Frequency (MHz)	50-510	152-1760	77-717	20-320	50-200	50-460
Precision	INT12	INT4/INT8	FP4/FP8	INT8	INT8/16	BF16 (A) INT8 (W)
Power (mW)	12.1-272.8	NA	10-122 (FP8)	8.27-250.65	47.5-469.2	16.95-105.0
Performance ¹⁾ (TOPS)	0.52-4.07 ²⁾	1.8 (INT8)	0.734 (FP4) 0.367 (FP8)	0.49-33.3	3.41 (INT8)	0.70 ⁷⁾ -3.58 ⁸⁾
Energy Efficiency ¹⁾ (TOPS/W)	1.91-27.56 ²⁾	39.1 (INT8) ³⁾	18.1 (FP4) 8.24 (FP8) ³⁾	1.96-53.83 ⁴⁾	22.9-47.8 (INT8) ⁶⁾	10.07 ⁷⁾ -88.36 ⁸⁾

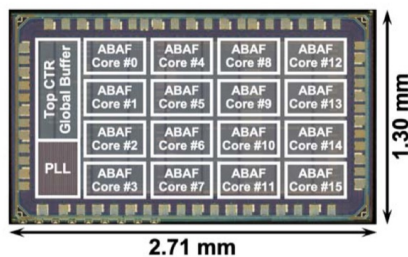
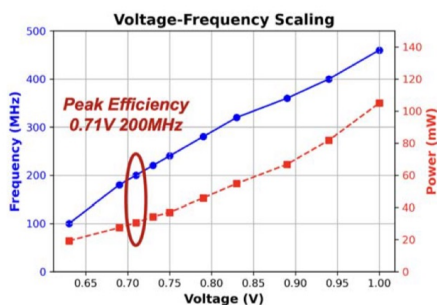
1) External memory access is excluded for fair comparison, 1 MAC = 2OPs. 2) Achieved when computing Score×V with 90% weak related tokens. 3) With 50% non-zero input density. 4) 92% sparsity in attention span. 5) This work relies on big-little network algorithm, which is not supported by the majority of LLM like LLaMA. 6) With 10% sampling ratio for zero-skipping. 7) ABAF, CAMP, and BMSR disabled, evaluated on one layer of LLaMA-2-7B. 8) The peak throughput / efficiency for a single layer (including attention and FFN) with ABAF, CAMP, and BMSR, ABAF compression ratio = 75%.

Figure 23.8.6: Measurement results and comparison with state-of-the-art processors.

测试总结

LLM generated response examples with proposed methods on OPT-1.3B

Input	Response
What is math?	Math is the study of numbers and their relationships.
What is artificial intelligence?	Artificial Intelligence, basically, it is the ability of a computer program to think.
Where is New York?	New York, NY is located in the United States.



Specification		
Technique (nm)		28
Die Area (mm ²)		3.52
Supply Voltage (V)		0.63 - 1.0
SRAM (KB)		384
Frequency (MHz)		50 - 460
Data Precision		A: BF16 W: INT8 ¹⁾
Power (mW)		16.95 - 105
Performance (TOPS) ²⁾		0.70 ³⁾ - 3.58 ⁴⁾
Area Efficiency (TOPS/mm ²) ²⁾		0.20 ³⁾ - 1.02 ⁴⁾
Energy Efficiency (TOPS/W) ⁵⁾	0.71V 200MHz	10.07 ³⁾ – 88.36 ⁴⁾
	1.0V 460MHz	6.73 ³⁾ – 59.03 ⁴⁾

1) With ABAF compression. 2) Evaluated at 1.0V, 460MHz, 1MAC = 2OPs. 3) Baseline with no optimizations. 4) The peak throughput / efficiency for a single layer (including attention and FFN) with ABAF, CAMP, and BMSR, ABAF compression ratio = 75%. 5) EMA energy cost is not included.

Figure 23.8.7: Large language model generation examples and chip summary.