

AI Computing Design Trends for LLMs in the Generative AI Era

Bor-Sung Liang, Ph.D

Dec 2024

MEDIATEK

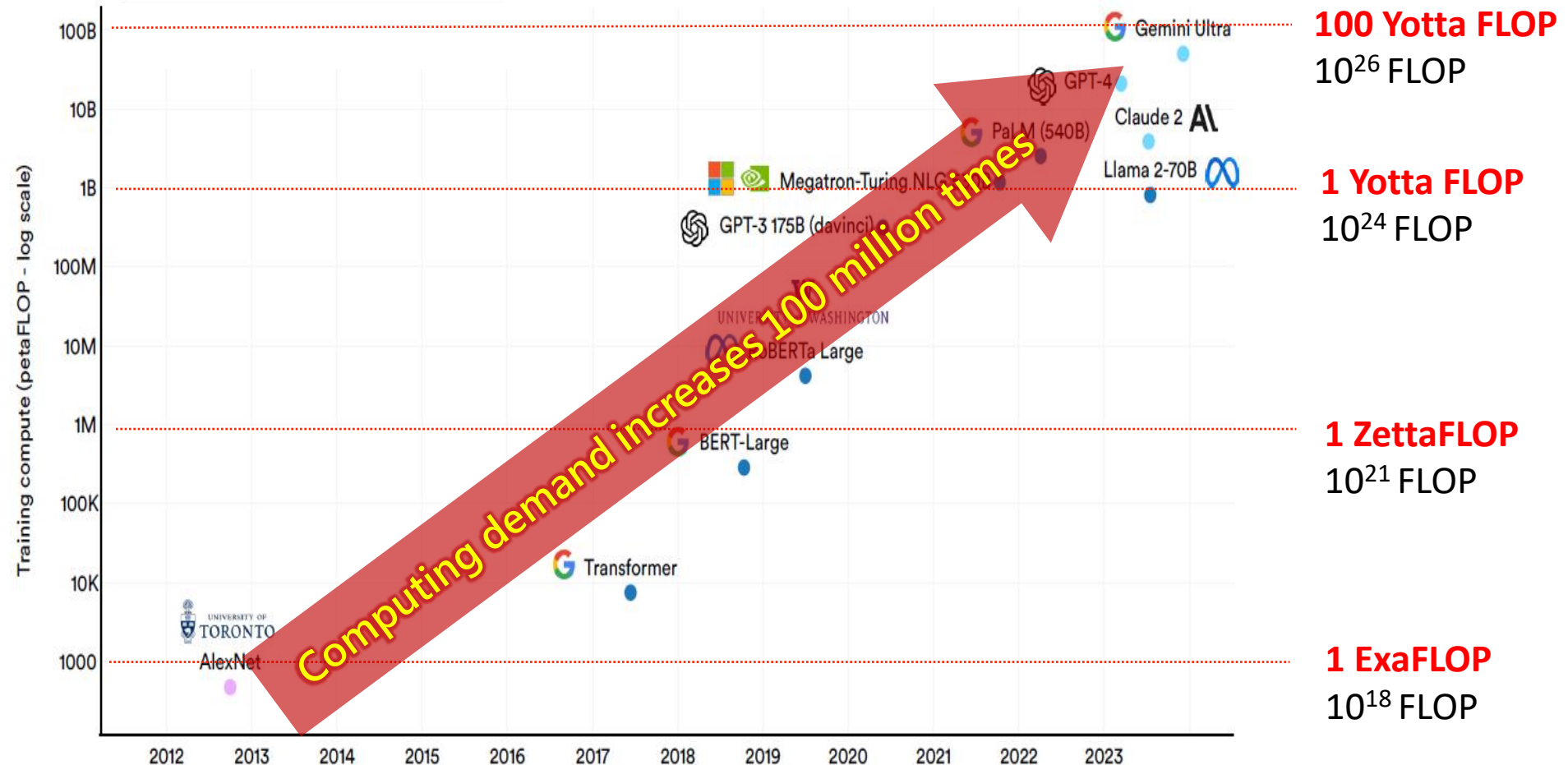
bs.liang@mediatek.com

Bor-Sung Liang
bsliang@gmail.com

The computing power required to train AI models has increased by 100 million times in ten years

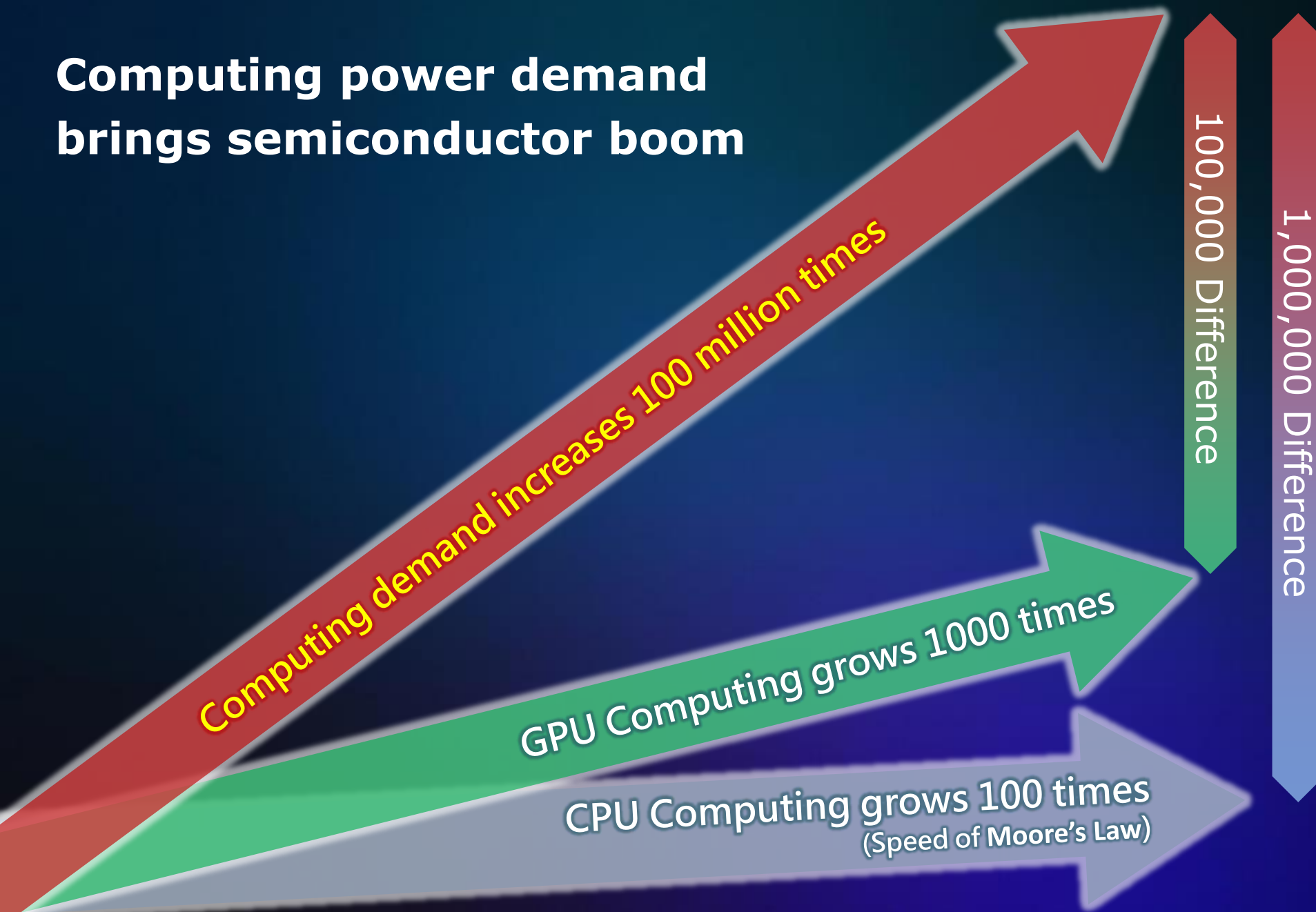
Training compute of notable machine learning models by domain, 2012–23

Source: Epoch, 2023 | Chart: 2024 AI Index report



Source: Stanford HAI, "AI Index 2024", <https://aiindex.stanford.edu/report/>

Computing power demand brings semiconductor boom



Large number of computing chips to provide computing power

→ semiconductor boom

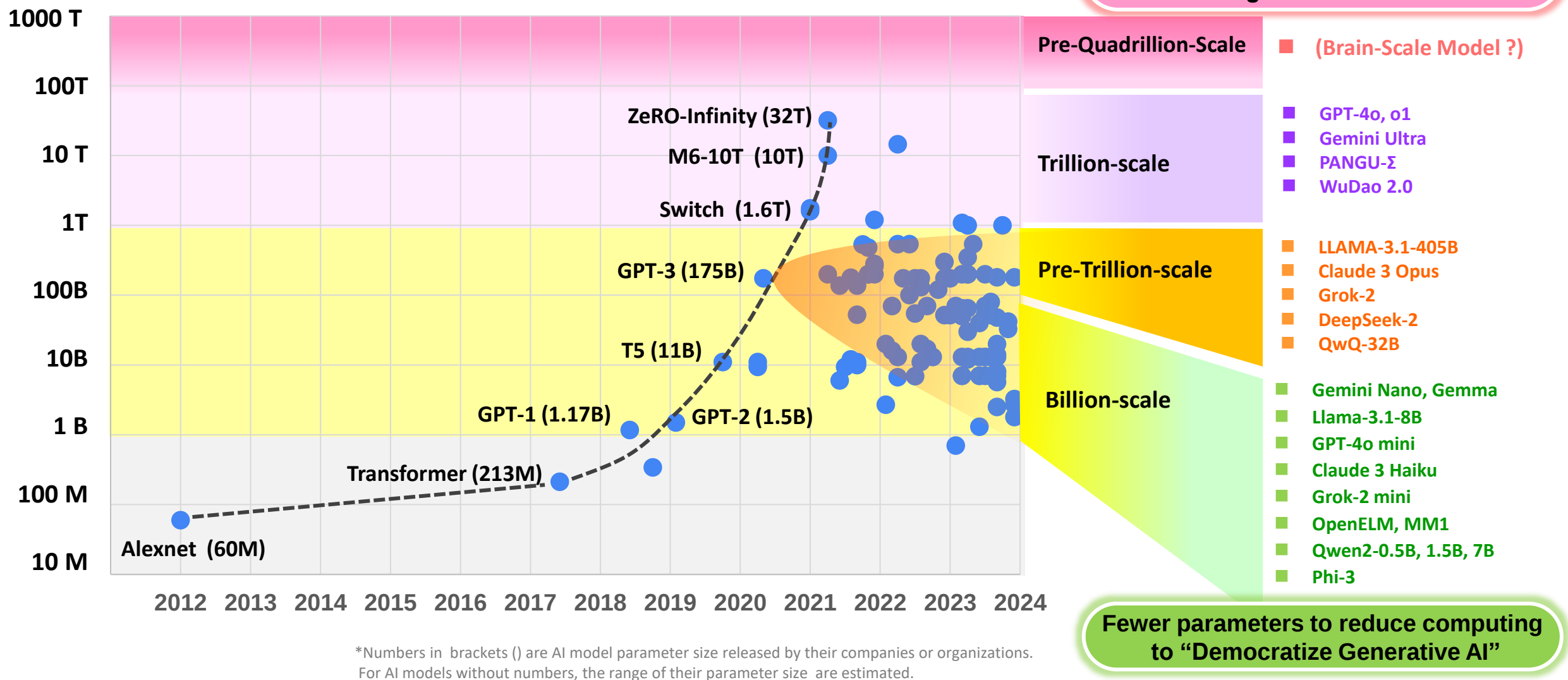
New Golden Age For AI & IC

- Huge demand on IC/Semiconductors
- Computing Efficiency
- Energy Efficiency



Trends of LLM Models

Trends of AI Models in a Decade



Trends of LLM Models for Inference

Full-Scale LLM

100s Bn ~ Tn
Parameters

Professor / Ph.D



Graduate Student



High School Student



Middle-size LLM

10s~100s Billion
Parameters

Small-size LLM

Billion Parameters

- Knowledge Distillation
- Model Reduction
 - Number Representation
 - Sparsity / Pruning
- Mixture of Experts (MoE)

Trends of LLM Models for Inference

Training

Full-Scale LLM

100s Bn ~ Tn
Parameters



- Knowledge Distillation
- Model Reduction
 - Number Representation
 - Sparsity / Pruning
- Mixture of Experts (MoE)

Middle-size LLM

10s~100s Billion
Parameters

Small-size LLM

Billion Parameters

Inference
on
Cloud
Servers

Inference
on
Edge
Devices

Inference

Trends of LLM Models for Inference

Full-Scale LLM

100s Bn ~ Tn
Parameters



- Knowledge Distillation
- Model Reduction
 - Number Representation
 - Sparsity / Pruning
- Mixture of Experts (MoE)

How to improve
performance?



Prompt

Guidance from teachers,
to know the key points & role
(Tips, Hints)

In-context
Learning

Provide background information
to extent capability
(Reference manual, Cheat Sheets)

RAG

Prepare FAQ and standard procedure
to answer various problems
(FAQ · SOP)

RLHF
RLAIF

Hire a coach (Human or AI)
to practice and to improve perform.
(Mock exams, practice games)

AI Agents

Memory, Planning and Tool Skills
Multi-Agents Collaboration
(Domain expert, Task Force)

Chain of
Thought

Self-reflection, Awareness,
Reasoning to Improving Perform.
(Self-review, Reasoning)

RLHF= Reinforcement Learning from Human Feedback
RLAIF = Reinforcement Learning by AI Feedback
RAG = Retrieval-Augmented Generation

Training

Inference

Trends of LLM Models for Inference

Training

Inference

Full-Scale LLM

100s Bn ~ Tn
Parameters



- Knowledge Distillation
- Model Reduction
 - Number Representation
 - Sparsity / Pruning
- Mixture of Experts (MoE)

Middle-size LLM

10s~100s Billion
Parameters

Small-size LLM

Billion Parameters

Prompt

RAG

AI Agent

In-context
Learning

RLHF
RLAIF

Chain of
Thought

Prompt

RAG

AI Agent

In-context
Learning

RLHF
RLAIF

Chain of
Thought

Prompt

RAG

AI Agent

In-context
Learning

RLHF
RLAIF

Chain of
Thought

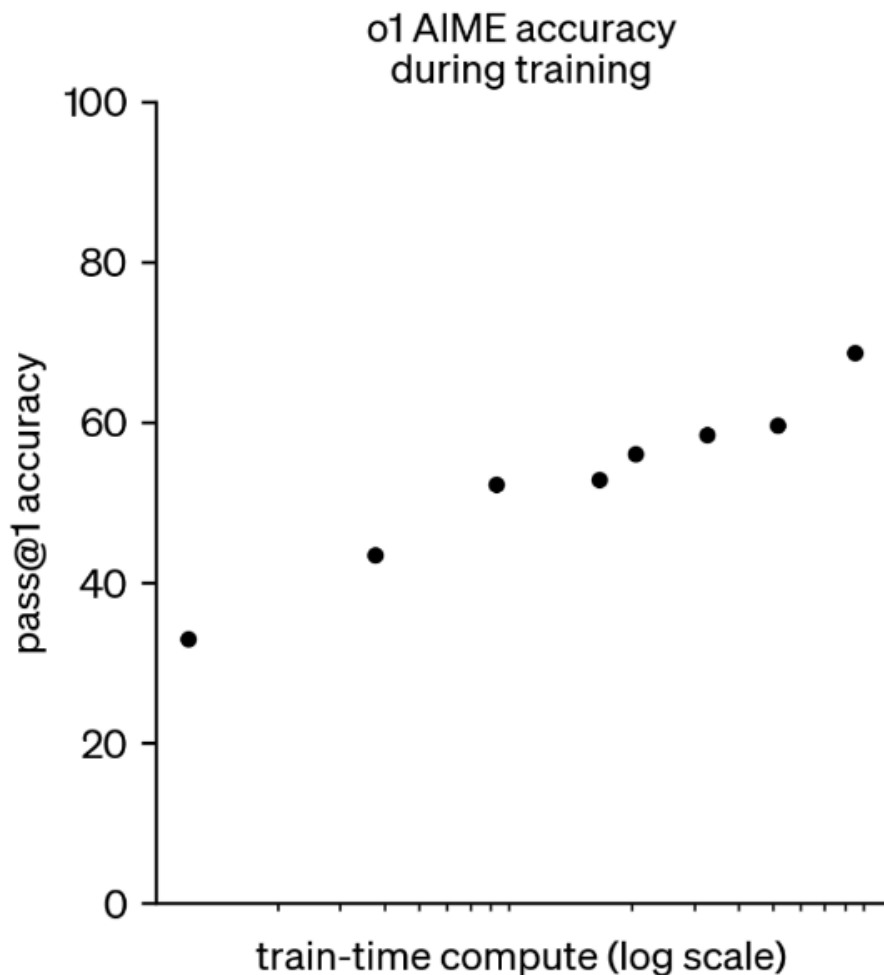
Inference
on
Cloud
Servers

Inference
on
Edge
Devices

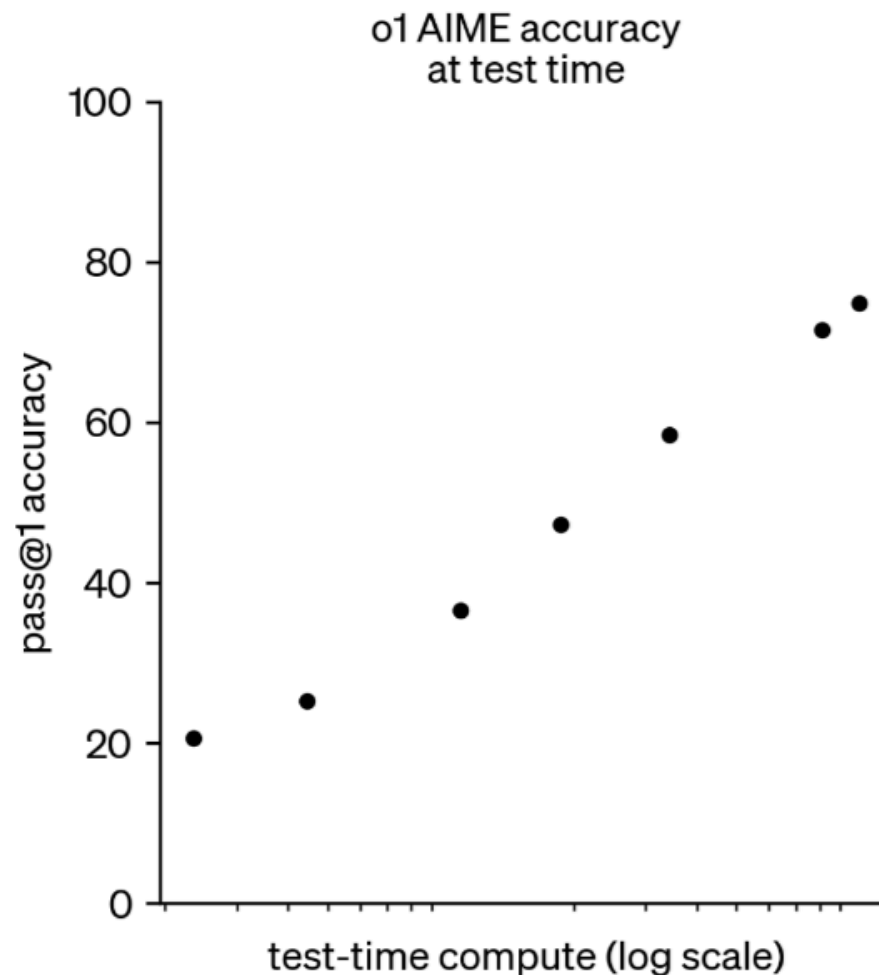
RLHF= Reinforcement Learning from Human Feedback
RLAIF = Reinforcement Learning by AI Feedback
RAG = Retrieval-Augmented Generation

Two Paradigms of Scaling to Increase LLM Performance

Training Time Compute



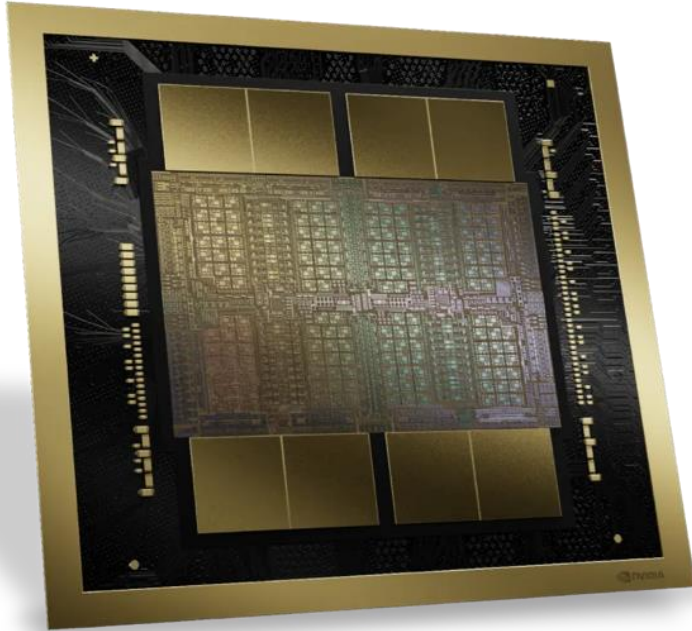
Test-time Compute (Inference)



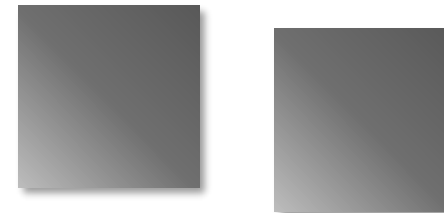
Source: OpenAI, 2024/9/12, "Learning to Reason with LLMs", <https://openai.com/index/learning-to-reason-with-llms/>

AI Computing : Data Center vs Edge Devices

NVIDIA® B200



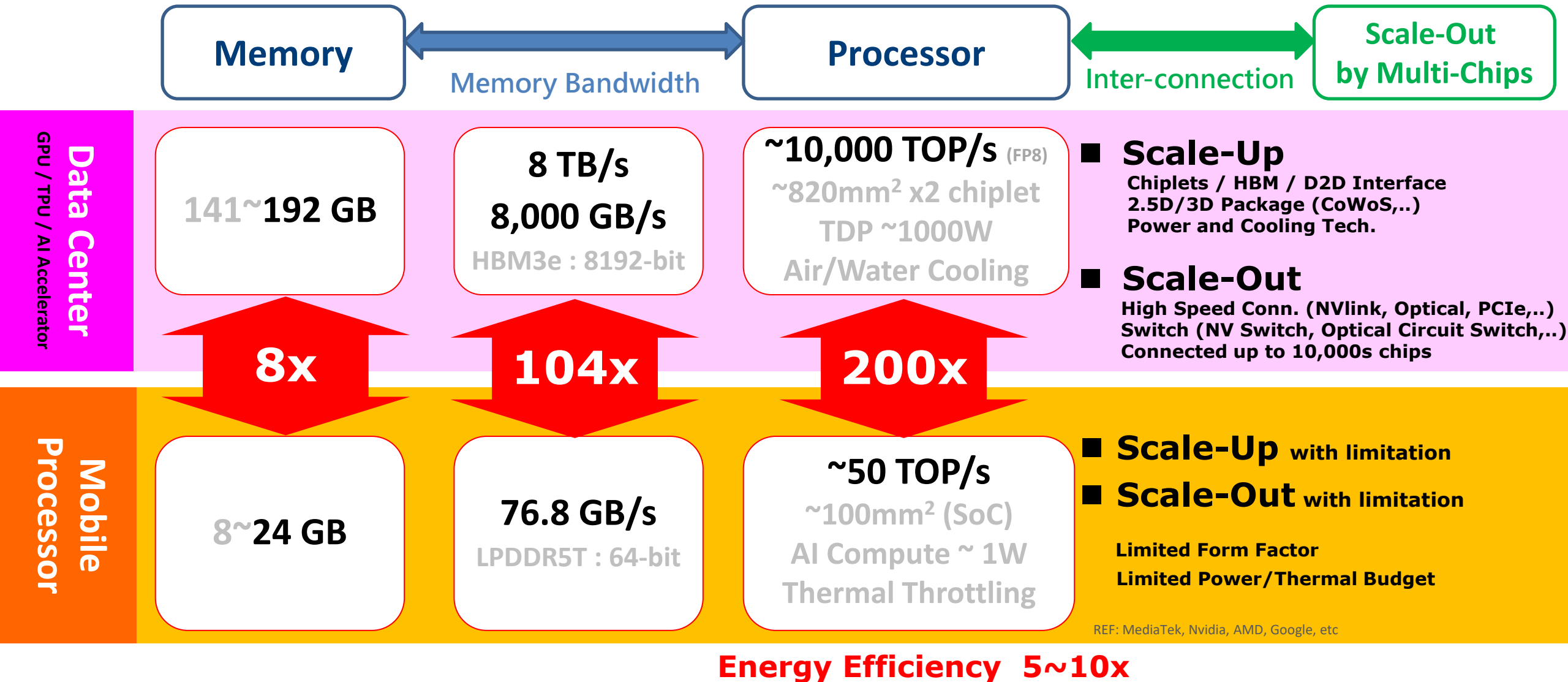
Die Size > 1600mm² (2 chiplet)
TDP ~ 1000W
10,000 TOP/s (FP8 , with Sparsity)
20,000 TOP/s (FP4, with Sparsity)



Die Size ~100mm²
TDP ~10W

~50 TOP/s

AI Computing : Data Center vs Edge Devices



AI Computing in Edge Devices

AI Inference on Mobile Processors for Large Language Models (LLMs)

On-Device

AI Model Parameter Numbers
Enlarged from 10s Million



10s M*

Face Unlock
Voice Assistant



100s M*

AI Camera
Photo Beautify
Image Recognition



1~10B*

10s B*

100s B*

1s T*

Natural Language Communication
Image/Video/Voice Understanding
Text/Image/Content Generation
AI Agent, AI Robot

Large Language Model

AI Model Parameter Numbers
Up to 100s Billion ~ 1s Trillion

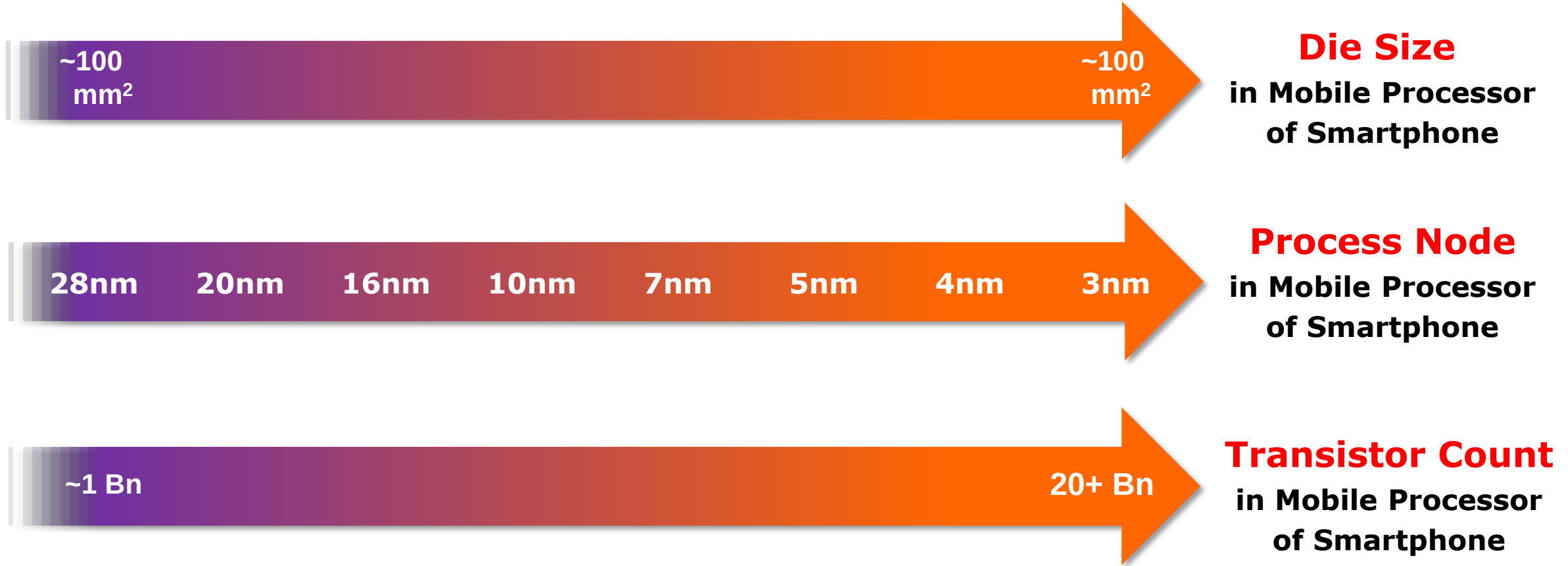
ChatGPT, GPT-4o, Gemini
Llama, Claude, Grok, PaLM
Mistral, Ernie, PANGU

* parameter numbers in AI models

Trends of Mobile Processor in a Decade

2013

2023



Trends of Computing Architecture in Mobile Processor

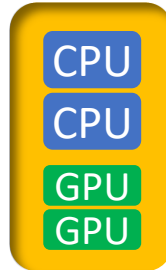
2007

2023

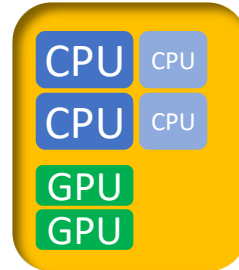
Single
Core



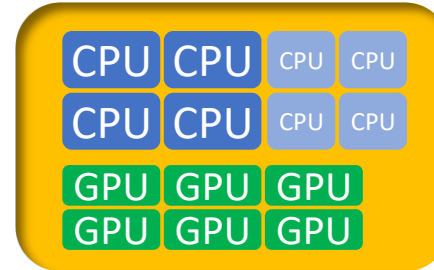
Multi
Core



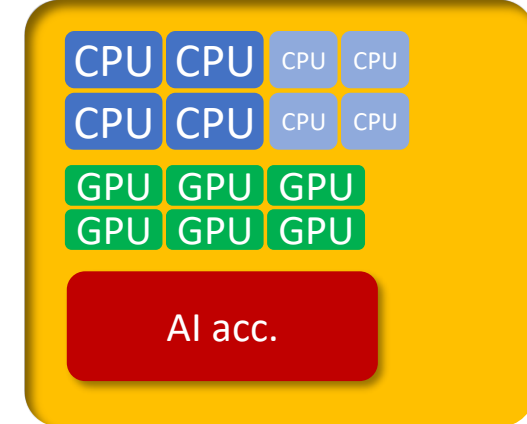
Big-Little
Multi Core



...More Cores...



AI Accelerator



- CPU Center Processor Unit
- GPU Graphics Processor Unit
- AI acc. AI Accelerator

*The number of CPU, GPU, and AI accelerator cores is for reference. Each application processor will have different architecture and number depending on the actual situation.

Model Size for different Parameter and Number Representation

AI Training

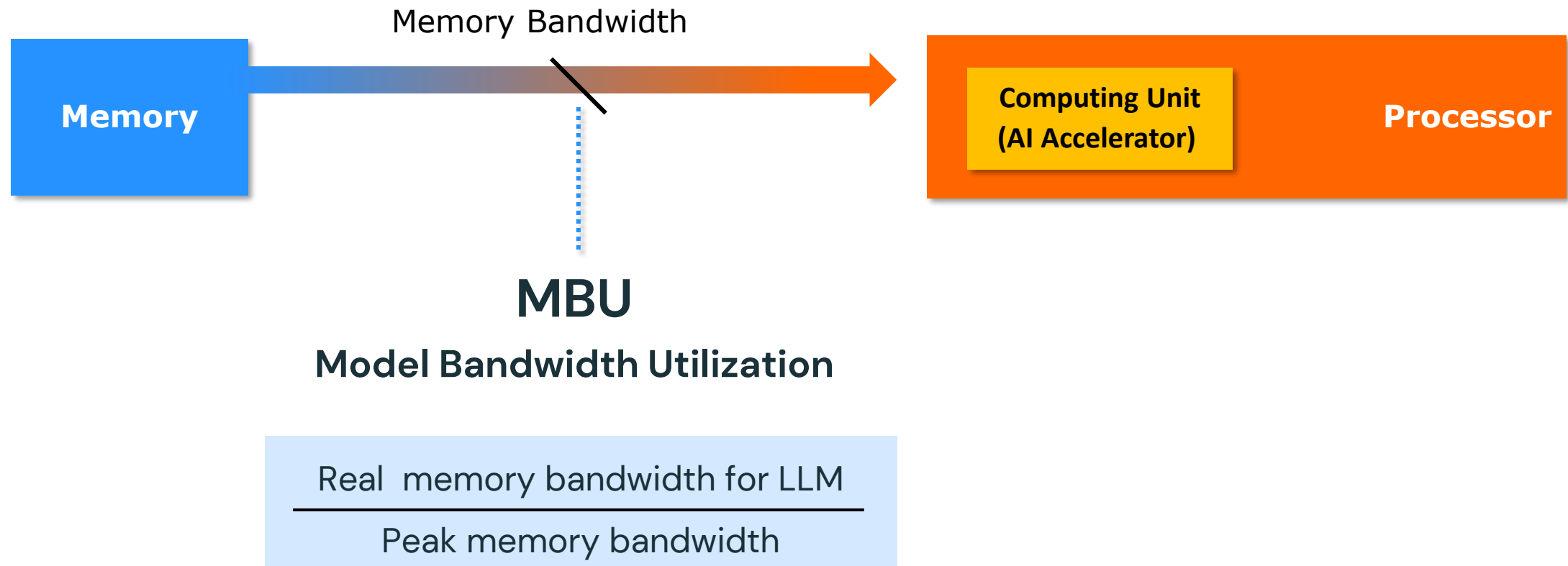
AI Inference

Model Parameters		FP32	FP16, BF16	INT8, FP8	INT4, FP4
		32bit (4Byte)	16bit (2Byte)	8bit (1Byte)	4bit (0.5Byte)
Single digital Billion-Scale Parameter LLM	70 Bn	280 GB	140 GB	70 GB	35 GB
	13 Bn	52 GB	26 GB	13 GB	7.5 GB
	7 Bn	28 GB	14 GB	7 GB	3.5 GB
	3.25 Bn	13 GB	6.5 GB	3.25 GB	1.625 GB
	1.8 Bn	7.2 GB	3.6 GB	1.8 GB	0.9 GB



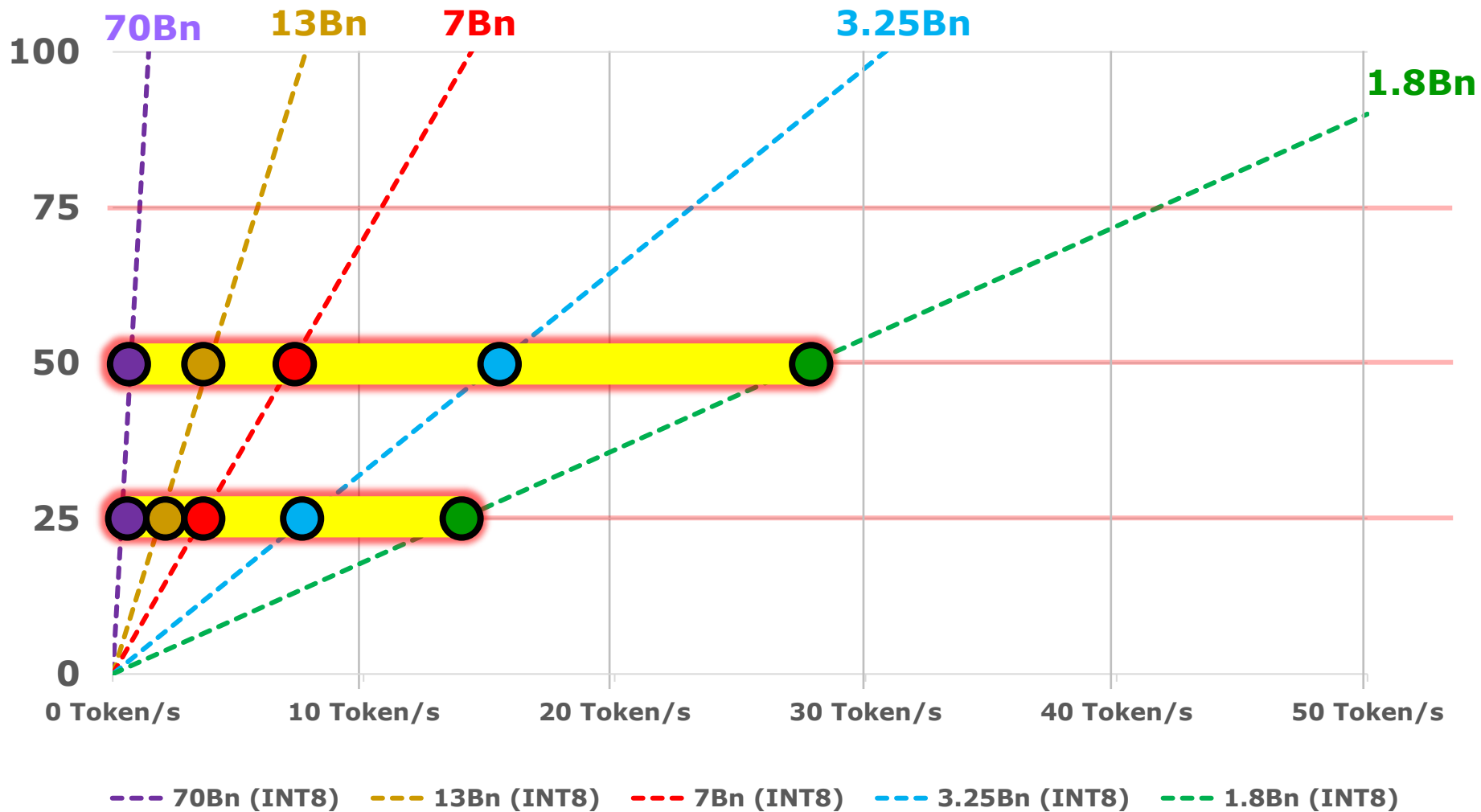
DRAM Size : >128GB 128~24GB 32~8GB 8~4GB 4~2GB <2GB

Bandwidth Utilization Ratio



Token Speed under Memory Bandwidth Limit (LLM @ INT8/FP8)

Memory Bandwidth (GB/s)



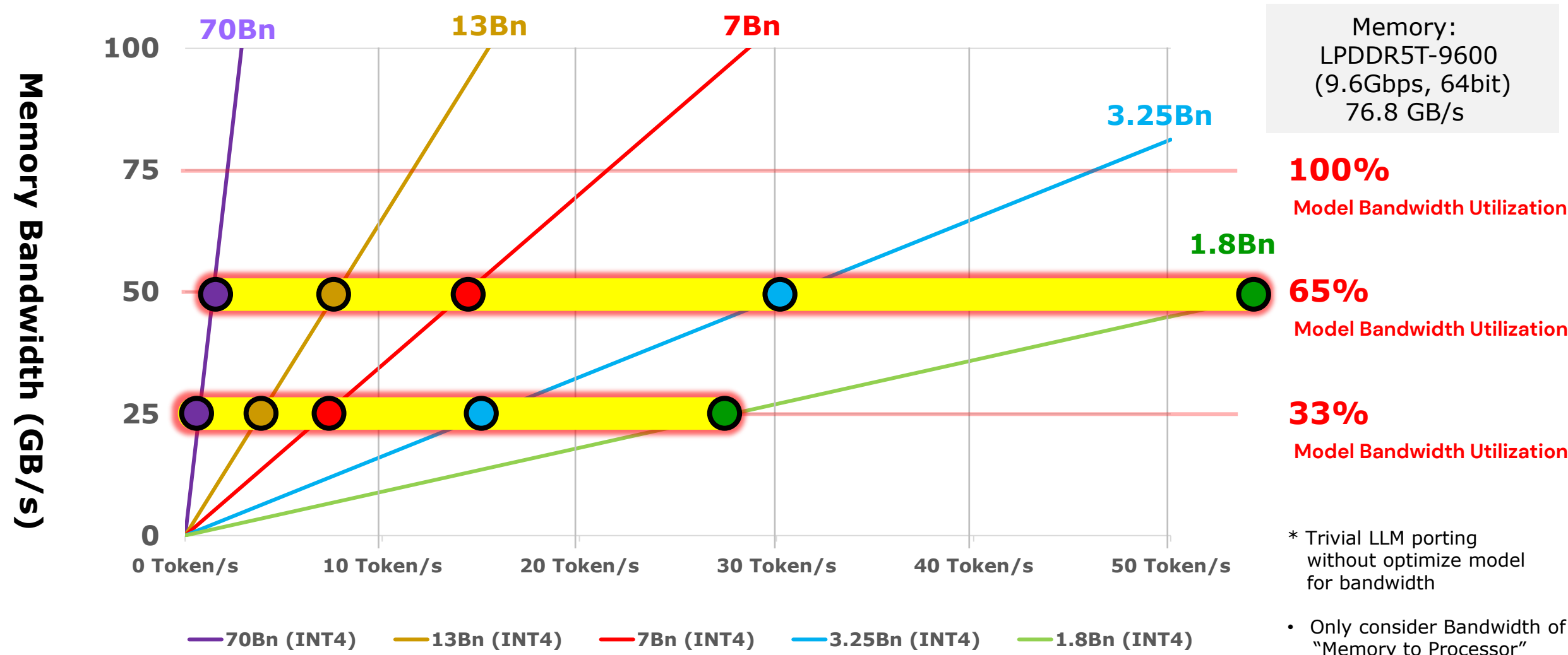
Memory:
LPDDR5T-9600
(9.6Gbps, 64bit)
76.8 GB/s

- 100%**
Model Bandwidth Utilization
- 65%**
Model Bandwidth Utilization
- 33%**
Model Bandwidth Utilization

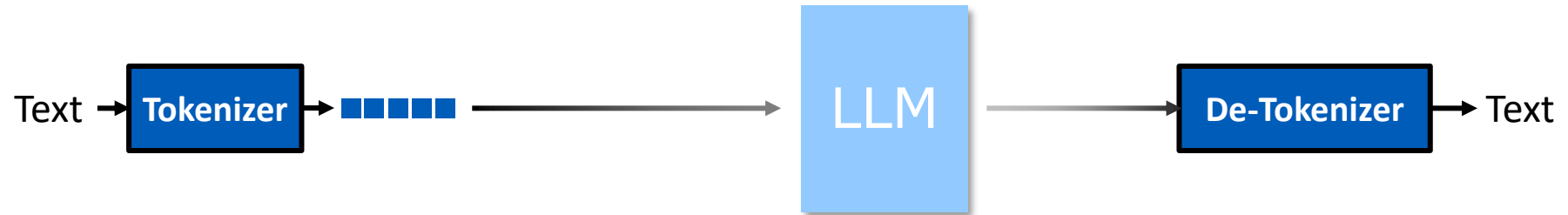
* Trivial LLM porting without optimize model for bandwidth

- Only consider Bandwidth of "Memory to Processor"

Token Speed under Memory Bandwidth Limit (LLM @ INT4/FP4)

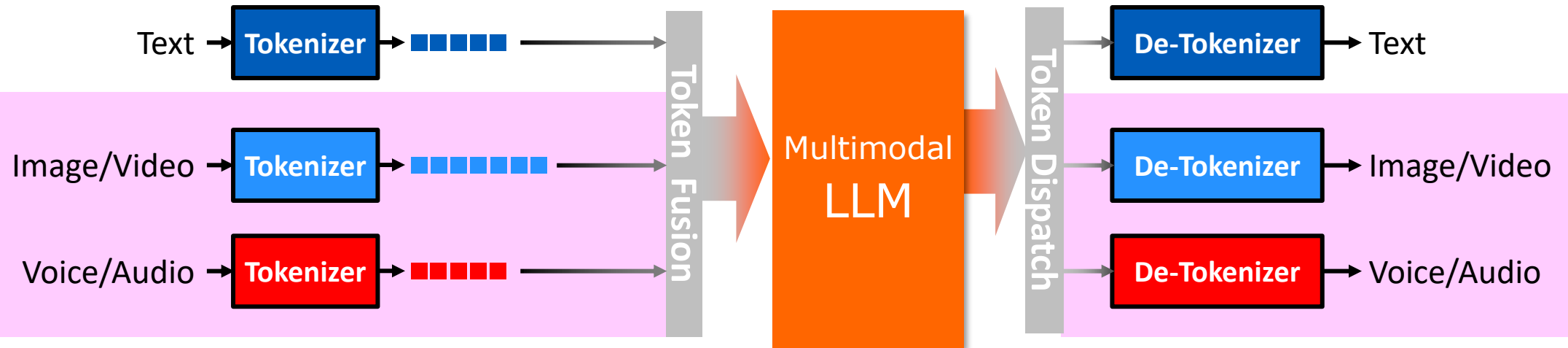


Token Process of LLM



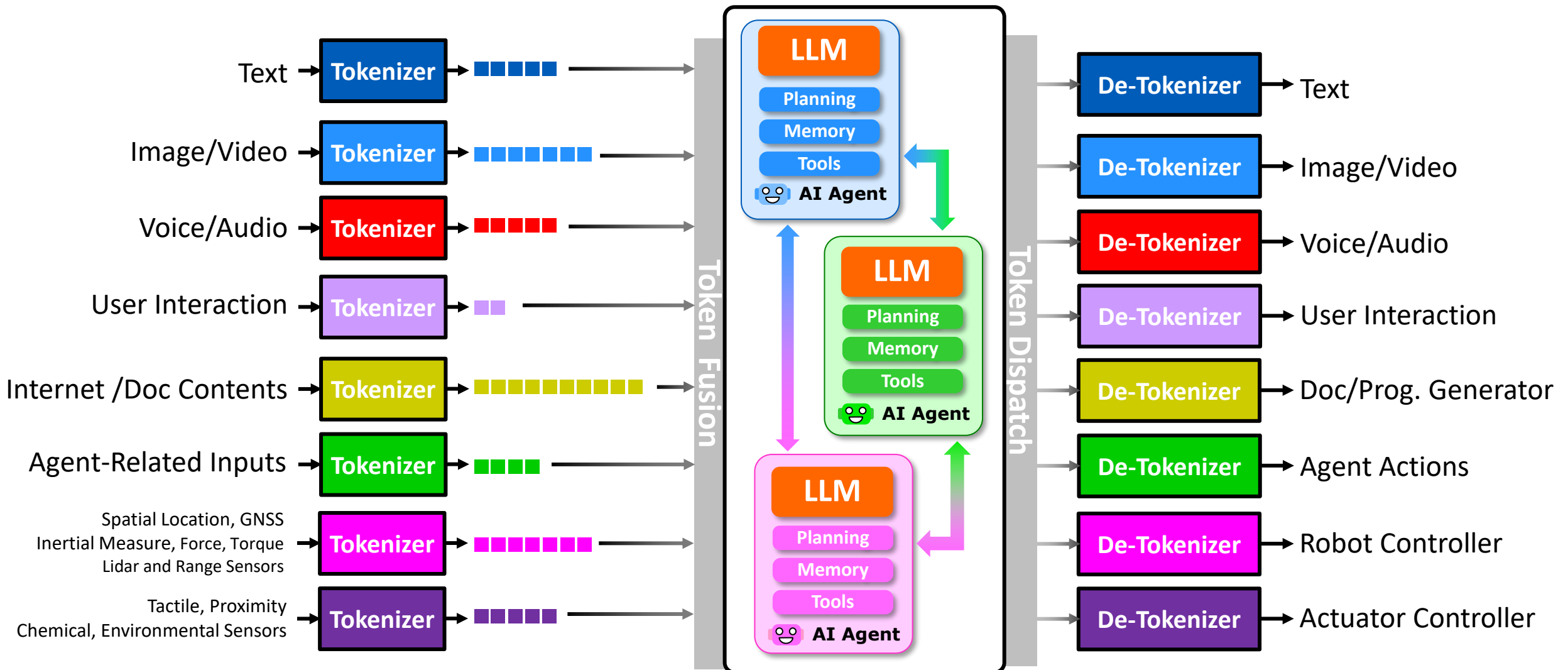
Mobile Devices
w/ text-based LLM

Token Process of Multimodal LLM



Mobile Devices
w/ Multimodal Gen AI

Token Process for AI Robotics with Multi AI Agents



Token Speed

How Fast can Humans Read? → related to how fast should LLM generate text

Read word-by-word

Read Normally 250 words/min

(Adults Silent-reading Speed 238 words/min)

The IEEE International 3D Systems Integration Conference (3DIC) will be held at the Hotel Metropolitan Sendai and Sendai Kokusai Hotel in Sendai, Japan in September 25-27, 2024. This conference combines the previous ASET and IEEE EDS Japan sponsored International 3D System Integration Conference, held in Tokyo in 2007 & 2008 and the Fraunhofer and IEEE CPMT sponsored International 3D System Integration Workshop held in 2003 & 2007 in Munich.

So much has changed since the first combined conference in San Francisco in 2009, 3D integrated circuits have moved from academic curiosity to solid commercial reality. The underlying technologies continually improve and evolve, incorporating new methods and resources. Today's hot topics include chiplets, photonic interconnect, micro-die handling, and exotic substrates. This 2024 conference will focus on the research and science of 3DICs. It will cover relevant 3D/Chiplet topics from manufacturing, processes, materials, and equipment to circuit designs, design methodology, and applications.



Skim Reading looking *only* for the general or main ideas
Scan Reading looking *only* for a specific information

Read Proficiently 1000 words/min

(Skimming speeds of 1000+ words/min)

The IEEE International 3D Systems Integration Conference (3DIC) will be held at the Hotel Metropolitan Sendai and Sendai Kokusai Hotel in Sendai, Japan in September 25-27, 2024. This conference combines the previous ASET and IEEE EDS Japan sponsored International 3D System Integration Conference, held in Tokyo in 2007 & 2008 and the Fraunhofer and IEEE CPMT sponsored International 3D System Integration Workshop held in 2003 & 2007 in Munich.

So much has changed since the first combined conference in San Francisco in 2009, 3D integrated circuits have moved from academic curiosity to solid commercial reality. The underlying technologies continually improve and evolve, incorporating new methods and resources. Today's hot topics include chiplets, photonic interconnect, micro-die handling, and exotic substrates. This 2024 conference will focus on the research and science of 3DICs. It will cover relevant 3D/Chiplet topics from manufacturing, processes, materials, and equipment to circuit designs, design methodology, and applications.

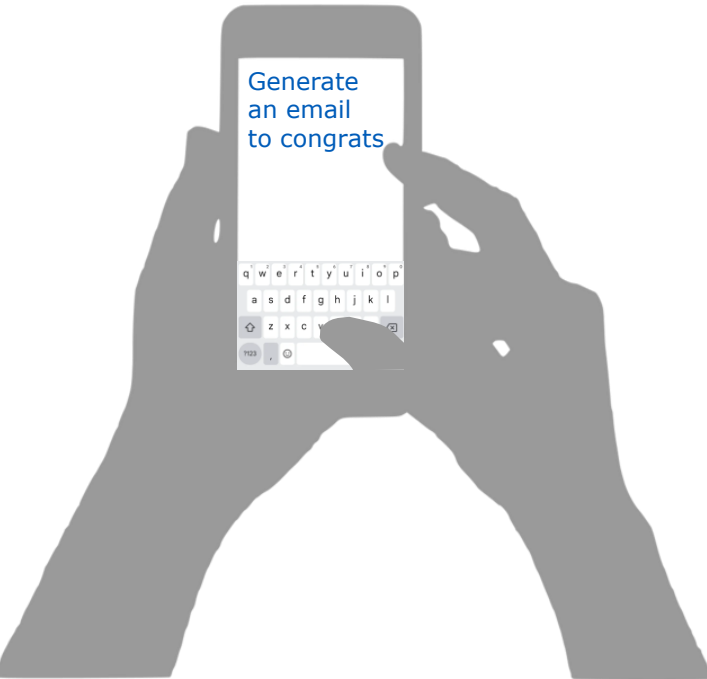


LLM in Mobile Devices

LLM (Large Language Model)

Multimodal LLM

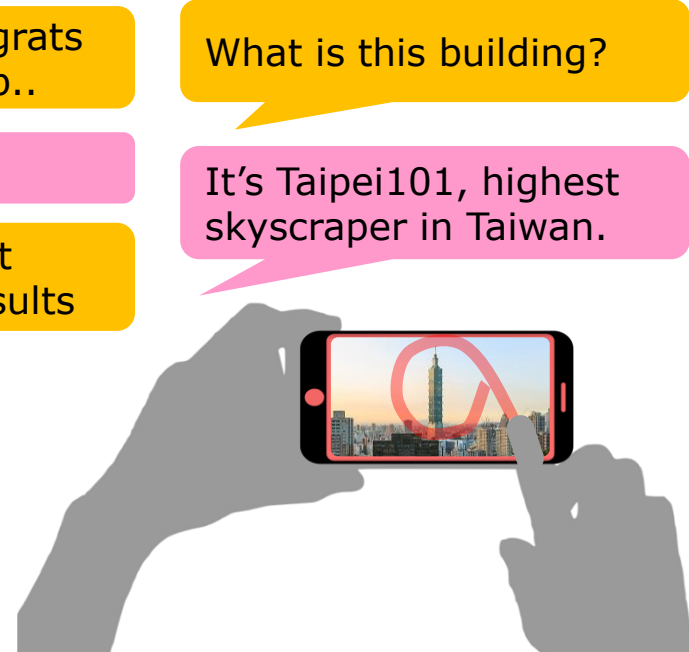
□ Typing Text



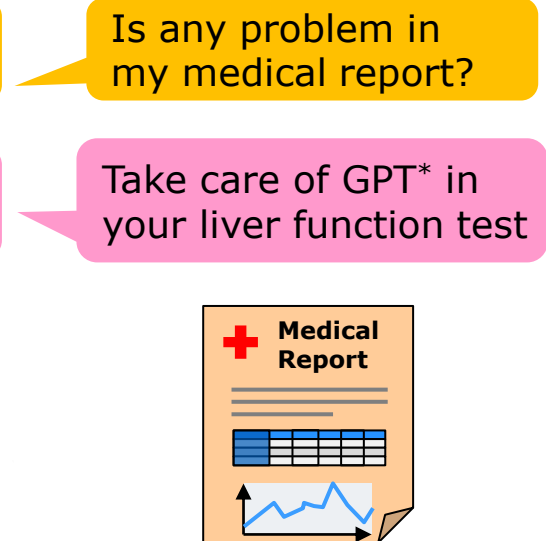
□ Voice
□ Nature Language Communication



□ Camera
□ Visual indicator



□ Image
□ Doc, Table



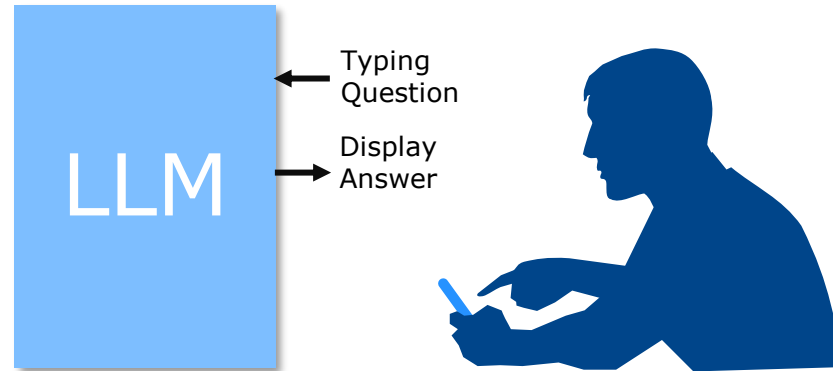
* GPT (glutamyl pyruvic transaminase) in liver function test

Types of Tokens

Simple LLM

Single Person

One Simple LLM Task



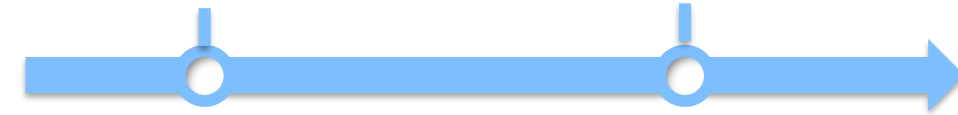
Only Tokens for Text

Read Normally
250 words/min

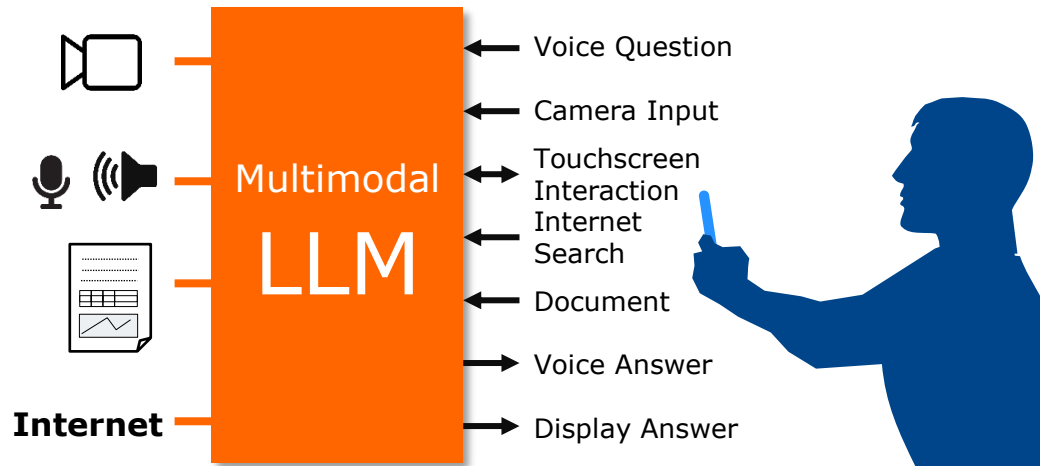
~8.33 Token/sec

Read Proficiently
1000 words/min

~33.3 Token/sec



Multimodal Interaction



More Various Tokens Concurrently

Tokens for Text

Tokens for Image/Video

Tokens for Voice /Audio

Tokens for User Interactive

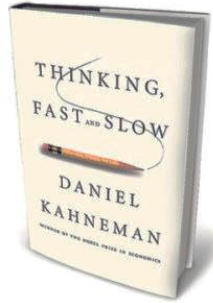
Tokens for Internet Contents

Tokens for Document Contents

LLM

Collaboration with Mobile OS and Cloud

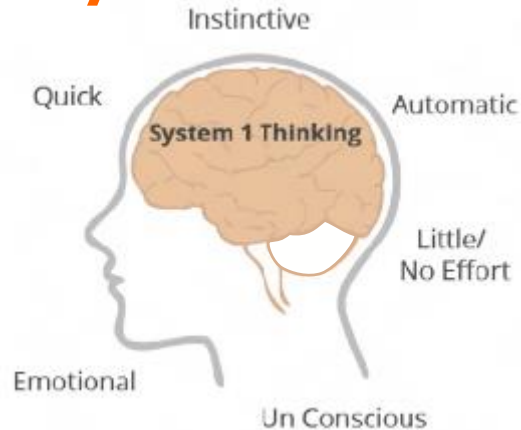
Current LLM: from "System 1" to "System 2"



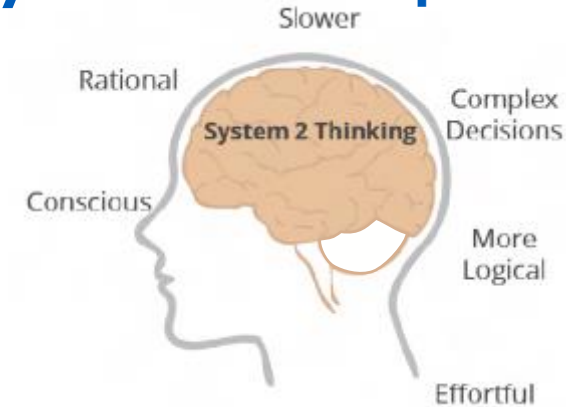
Cerebrum

Current LLM Status

System 1 : Instinctive



System 2 : Deep Thinking



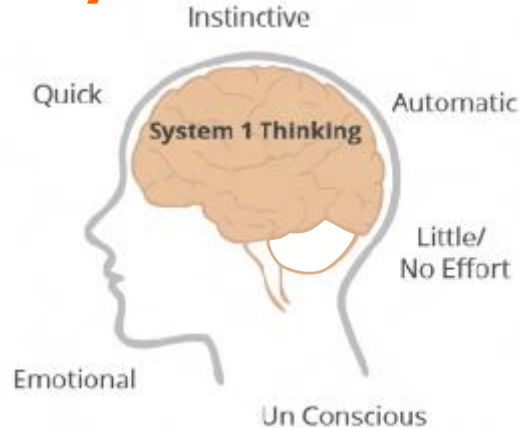
Inspired from System 1 and 2 :

Can we collaborate both Device-side and Cloud-side for System 1 and 2 ?

LLM – Thinking Fast, Deep Thinking and Swift Response

LLM on Device

System 1 : Instinctive



Cerebrum

Thinking Fast

Fast

Quick Response
Multimodal Signals
Multi-Tasks Concurrently

Cerebellum

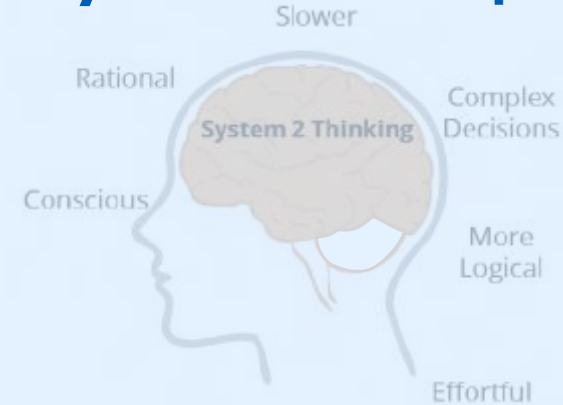
Brainstem



Swift
Response

LLM in Cloud

System 2 : Deep Thinking



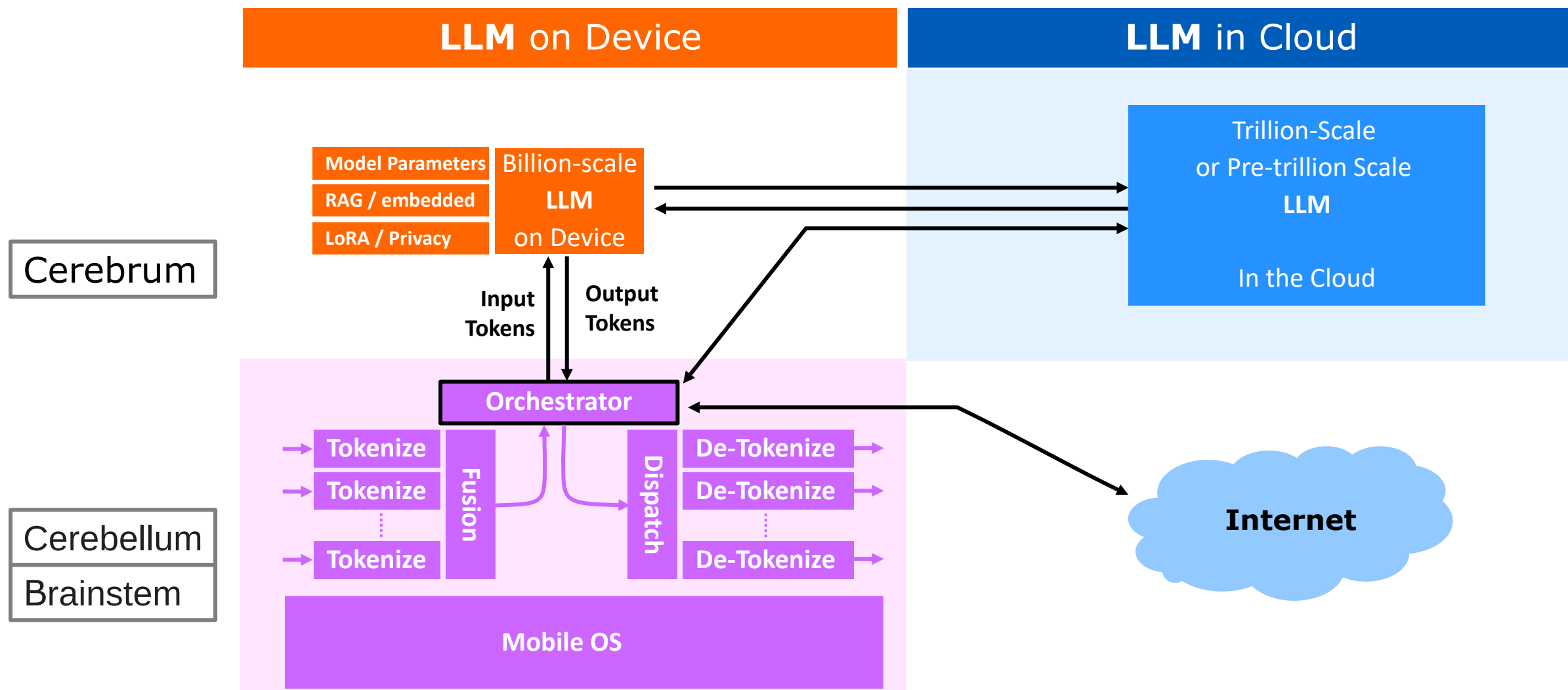
Deep
Thinking

~~Slower~~
Fast

provide sufficient
computing resources
to speed-up

→ Cloud Computing

LLM on Device - Collaborate with Cloud and Mobile OS



LLM as Mobile Software on Mobile Processor

LLM on Device

LLM in Cloud

Mobile SW

Model Parameters
RAG / embedded
LoRA / Privacy

Billion-scale
LLM
on Device

Orchestrator

Tokenize
Tokenize
...

Fusion

Dispatch

De-Tokenize
De-Tokenize
...

HW Interface

SDK

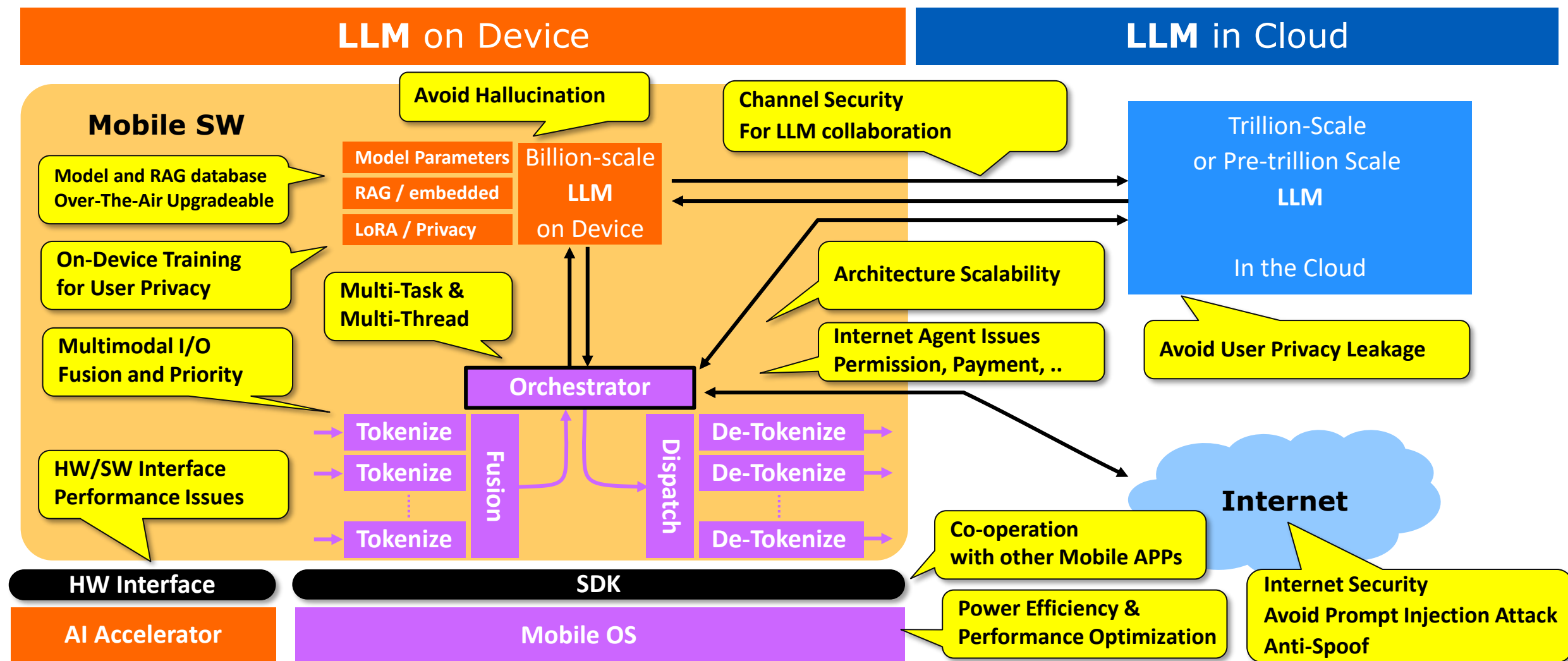
AI Accelerator

Mobile OS

Trillion-Scale
or Pre-trillion Scale
LLM
In the Cloud

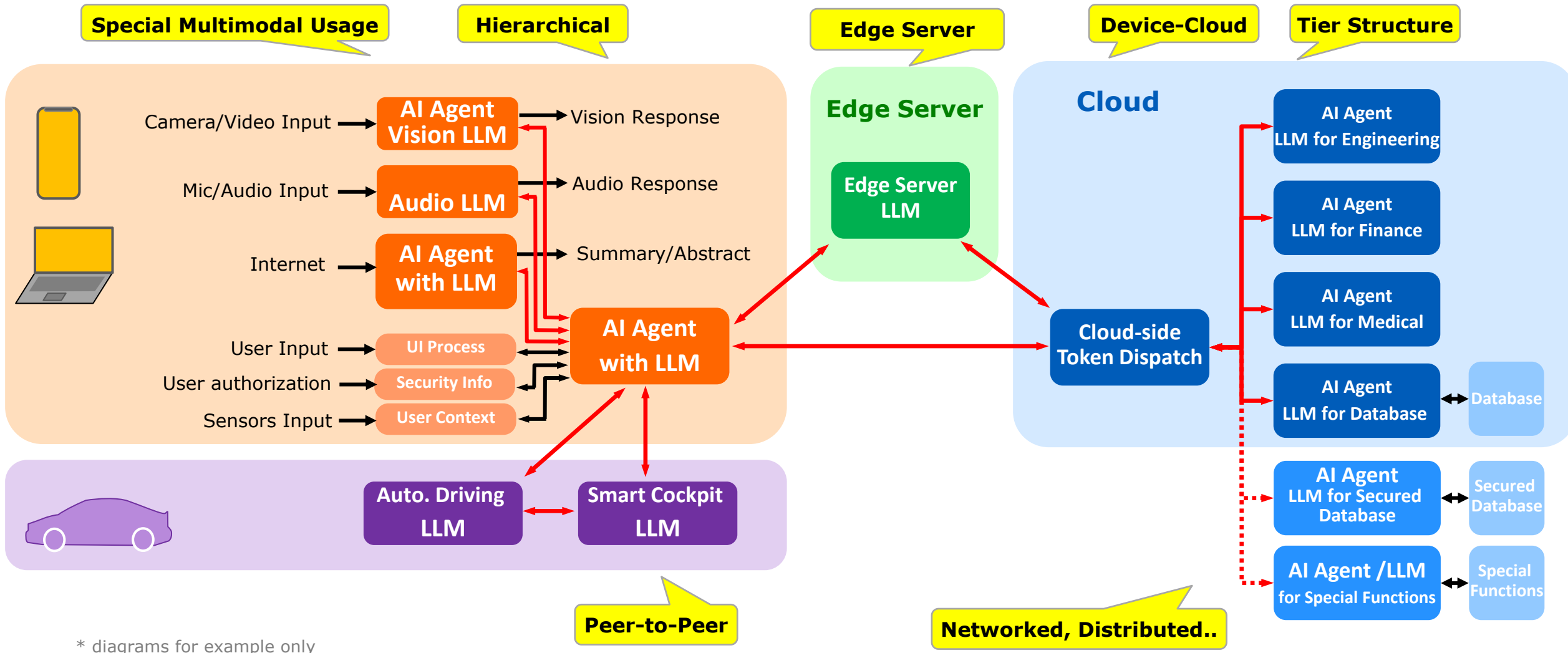
Internet

LLM on Mobile Processor : Issues



Trends of AI Computing for Data Center and Edge Devices

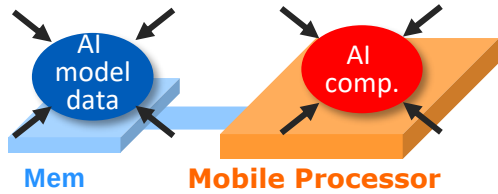
Collaboration among AI Models



Trends of AI Computing on Edge Devices

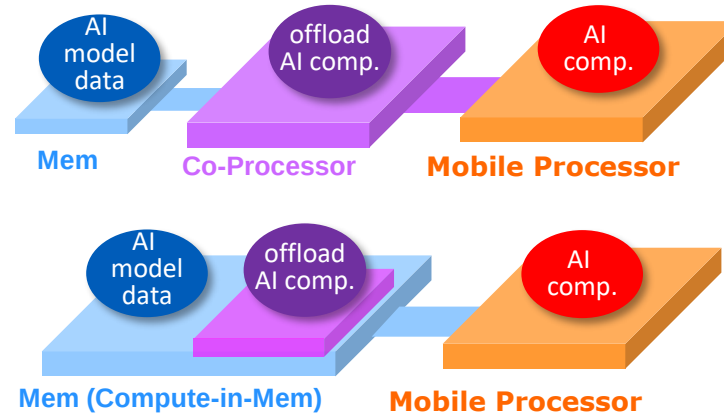
Reduce

Number Representation
Pruning/Sparsity
Smaller model by KD



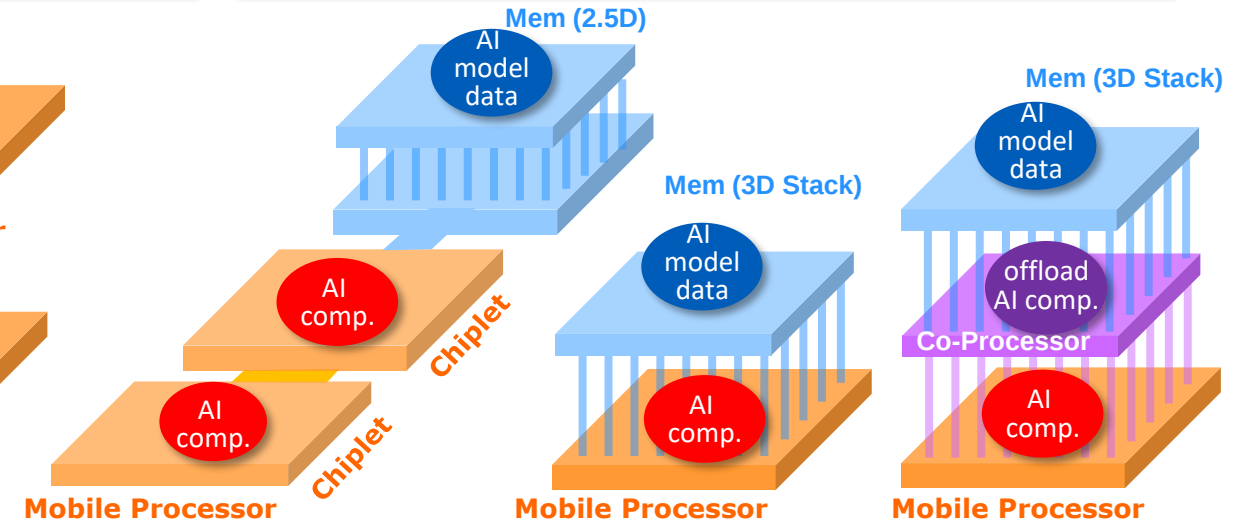
Offload

Co-processor
Compute in Memory



Scale-up, Scale-out

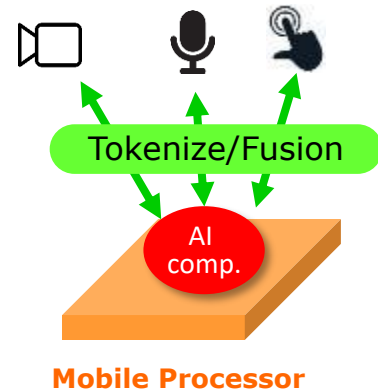
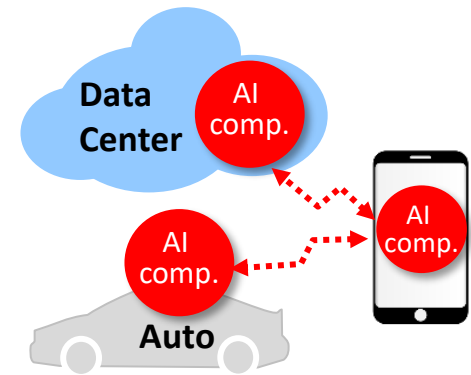
Chiplets
2.5D/3D Integration



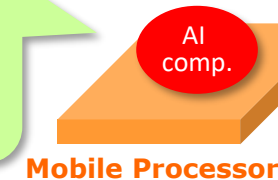
Collaboration

Multimodal Fusion/Pre-Processing

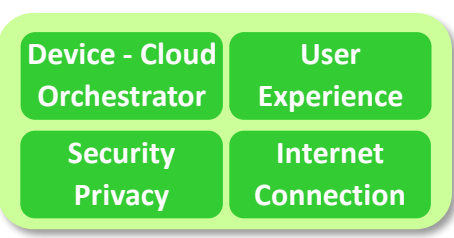
Model/System Architecture Optimization



LLM Model Related

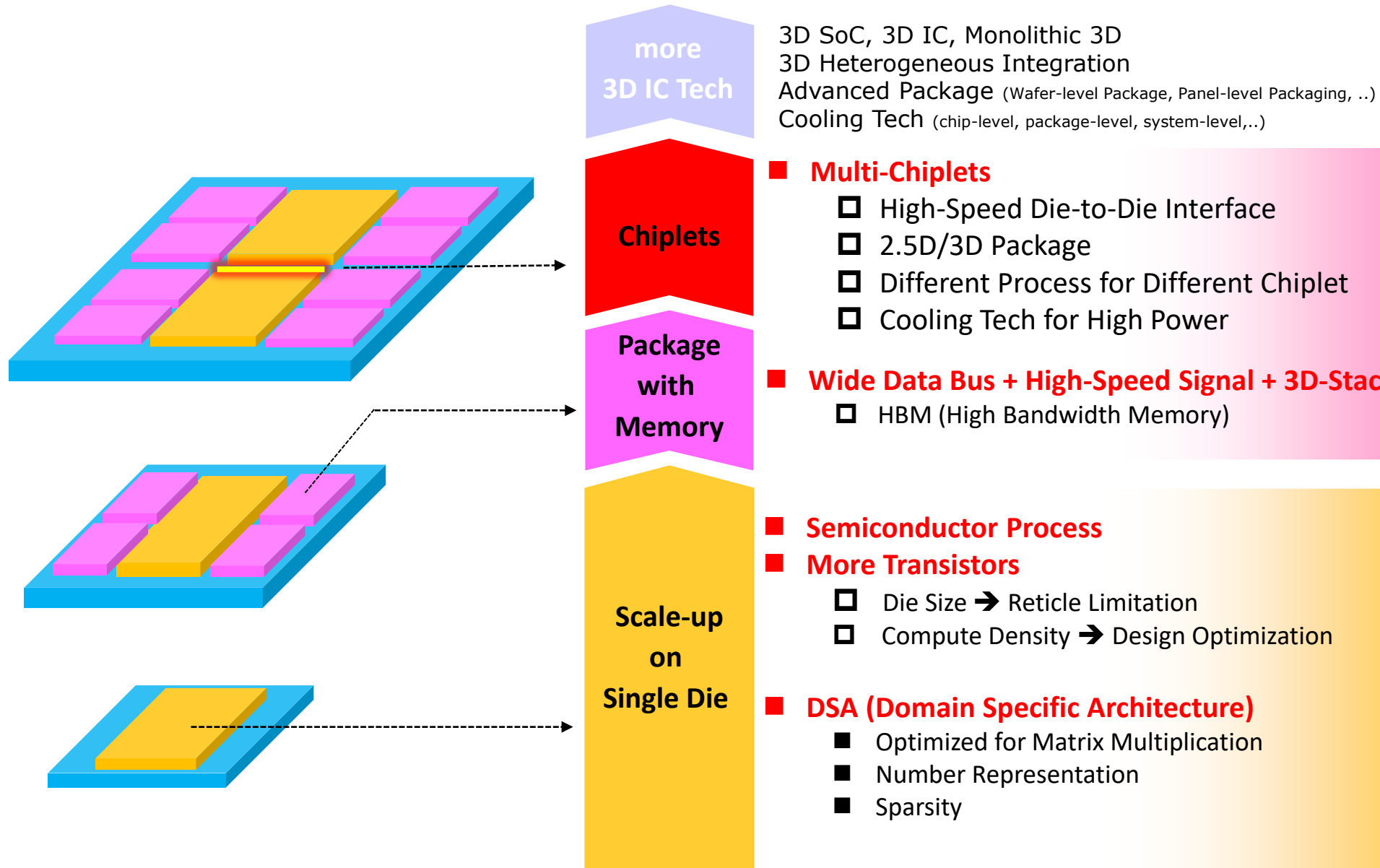


AI System Related



Trends for **Scale-up** for AI Computing in Data Center

- Data Flow Architecture
- Compute in Memory
- Compute Near Memory
- Neuromorphic Computing
- Analog Computing
- Approximate Computing
- Quantum Computing / Quantum-Inspired Computing



Increase
Compute & Bandwidth

But
Need to take care
Package, SerDes Thermal Issues

Increase
Compute

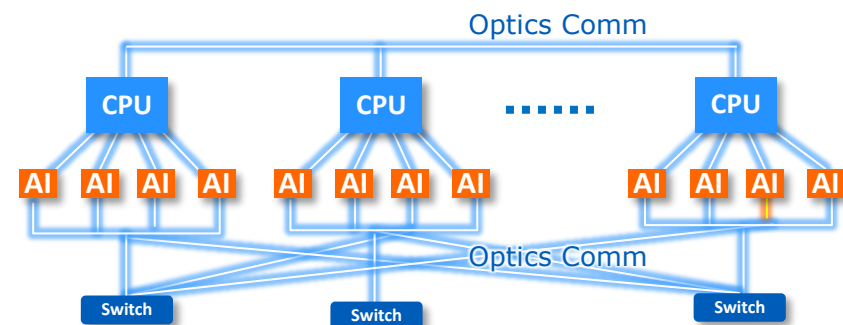
But Require
More Bandwidth and System Software Optimization

Trends for **Scale-out** for AI Computing in Data Center

Off-Board

In-Package
/On-Board

Off-Board

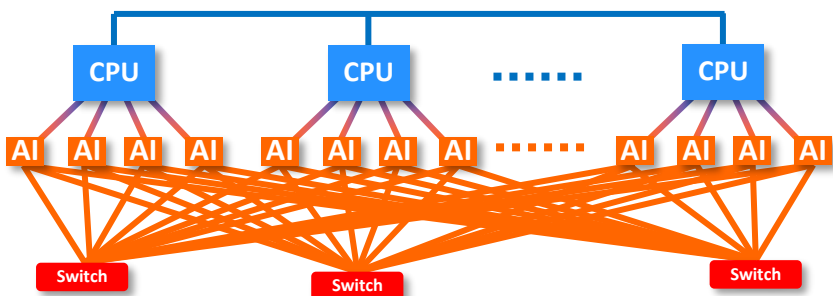


3D IC Tech

**Optical
Connection**

- ❑ High Bandwidth with Low Power
- ❑ Co-Packaged Optics (CPO)
- ❑ Packaging, Fiber Cable, Reliability
- ❑ Elec-Optical Signals, SerDes,...

**Large System
Compute
w/Low Power**
But Need to Handle
Optical System Issues

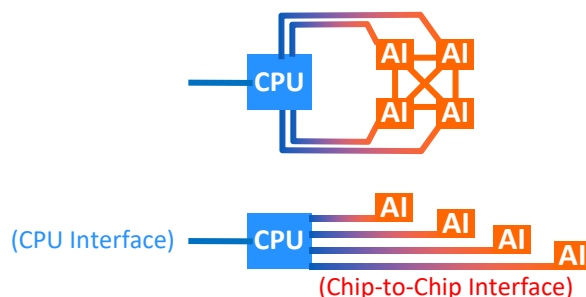


**High-Speed
Chip-to-Chip
Connection
with Switch**

- ❑ High-Speed High-Bandwidth Switch
- ❑ High-Efficiency Data Exchange in GPU (or TPU) subsystem
- ❑ Enable Larger-Scale System for LLM
- ❑ Flexibility for different Parallelism Schemes (Tensor Parallel, Pipeline Parallel, Expert Parallel, Data Parallel)

Increase
**Large System
Compute**

But Problems on
**Power,
Thermal Issues**



**High-Speed
Chip-to-Chip
Connection**

- ❑ CPU-GPU, GPU-GPU connection
- ❑ High-bandwidth, Low Latency
- ❑ Unified Memory Allocation

Increase
**Compute
Efficiency
For
Small Group
GPUs (or TPUs)**

* AI Chips can be GPU, TPU or ASIC

AI Chip Market (Forecasted in an article of *The Economist*)

- GPU or TPU with Software Stack
- Efficiency for AI Model Training
- Pre-Training, Post-Training, Evaluation
- AI Model Ability & Time-to-Market

- GPU, TPU or ASIC
- Efficiency for AI Inference
- Fine Tune, RAG, Chain-of-Thought, Agent
- Inference Capability for Million Users
& Energy Efficiency

- Mobile Processor or ASIC
- Efficiency for AI Inference on devices
- User Interaction, Privacy, Multimodal, Agent
- Inference with Smooth User Experience,
Orchestrator of Device-Cloud Collaboration
& Energy Efficiency

Training (15%)

Inference (85%)



15%
Training



45%
Data Center Inference



40%
Edge Device Inference



Data Center (60%)



Edge Device (40%)

- Scale-up and Scale-out
to **max** the scale of overall compute and bandwidth for performance
to **min** the energy of “compute per unit” and “comm per bit” for efficiency
- Heavily rely on the progress of semiconductor and package technologies
- **Affordable for data center companies. Sustainable for long-term operations**

- Hardware/Software/System co-design
to collaborate edge and cloud resources to boost AI performance
to offer smooth and swift user experience to use AI
- Heavily rely on domain specific design, and semi/package tech.
- **Affordable for end users. Penetrate to various AI applications**

• Ref: The Economist, “A Cambrian moment: AI has propelled chip architecture towards a tighter bond with software”, Sep. 16th, 2024. / Pic Source : irasutoya.com

AI Computing on Edge Devices and Data Center

Edge Devices

AI Model

Billion-level Parameter Model



Advantage of Edge Devices



Always Standby



High Energy-Efficiency



Privacy Personalization



Realtime Response

Wireless Comm



5G

6G

WiFi



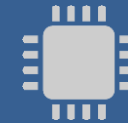
Data Center

AI Model

Pre-Trillion-level / Trillion-Level Parameter Model



Advantage of Cloud Computing



Exa-scale Computing



Scale-Out



Massive Data



Electricity