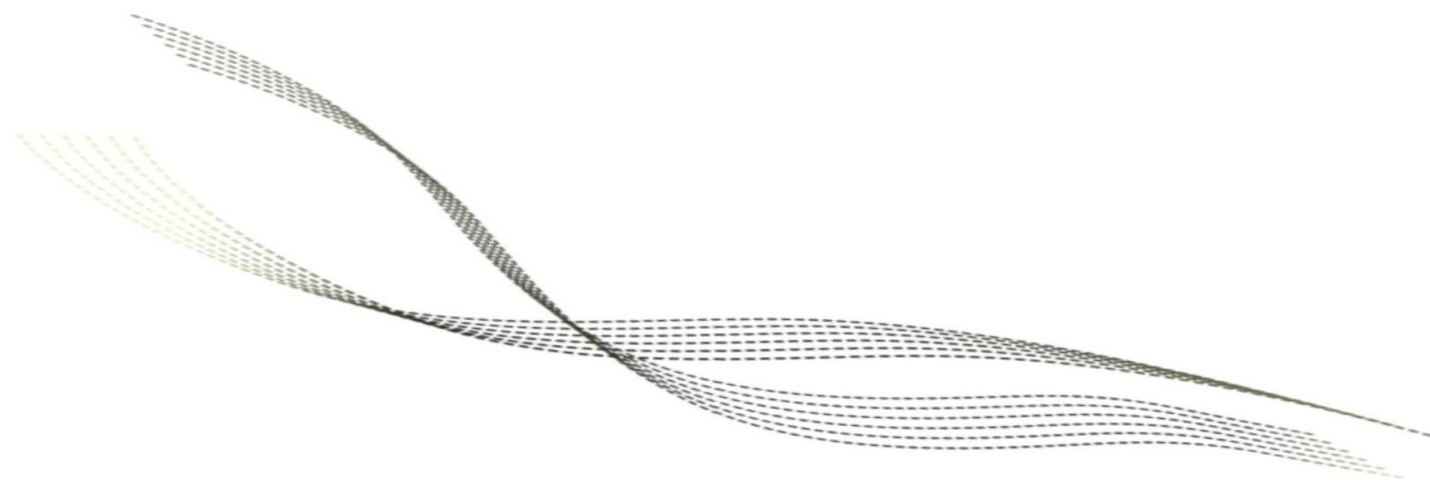


AI 芯片基础知识



并购家 免费行业报告网

IPOIPO.CN 每日更新 不用注册会员 无广告 永久免费

目录

CONTENTS

01

芯片的分类

02

CPU和GPU

03

ASIC和FPGA

04

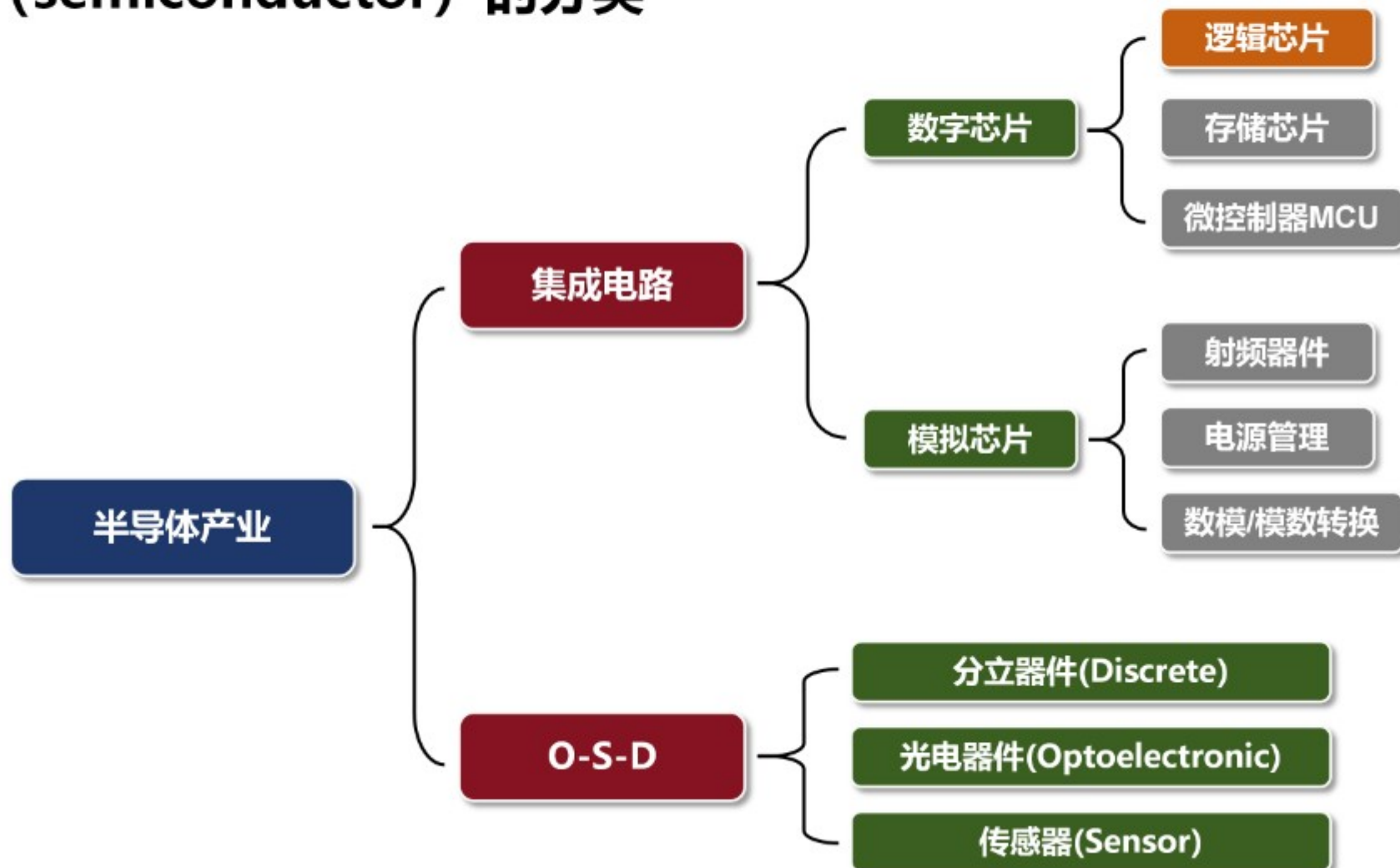
总结对比

PART 01

芯片的分类

■ 芯片的分类

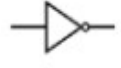
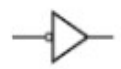
□ 半导体 (semiconductor) 的分类



■ 芯片的分类

□ 逻辑芯片 (Logic Chip)

- 逻辑芯片是一类用于执行特定逻辑运算的集成电路，是现代电子系统的核心组件之一。
- 逻辑芯片利用晶体管来构建各种逻辑门电路（例如：与门 AND、或门 OR、非门 NOT、异或门 XOR 等），进而组成更为复杂的电路，实现不同类别的逻辑运算。
- 逻辑芯片能够实现数据处理、控制和其他各种功能，广泛应用于消费电子、工业制造、教育医疗、国防军事等各个领域。

名称	图形符号
与门	
或门	
非门	 
与非门	
或非门	

■ 芯片的分类

□ 逻辑芯片的分类

- 根据功能和用途的不同，逻辑芯片可以分为以下几大类别：



PART 02

CPU和GPU

■ CPU 和 GPU

□ CPU 的定义

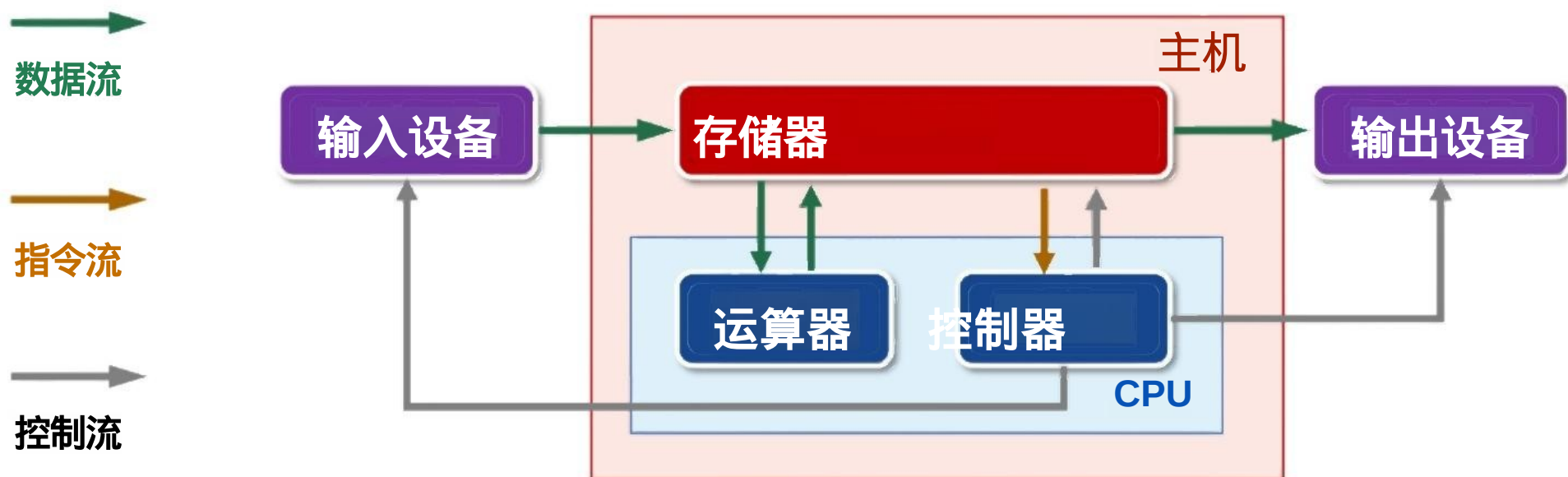
- CPU, 就是 Central Processing Unit, 中央处理器。
- CPU 是计算机系统的运算和控制核心, 是信息处理、程序运行的最终执行单元。
- CPU 的主要作用是解释计算机指令以及处理计算机软件中的数据。
- CPU 是计算机中负责读取指令, 对指令译码并执行指令的核心部件。



■ CPU 和 GPU

冯 · 诺依曼架构

- 现代计算机，都是基于 1940 年代诞生的冯 · 诺依曼架构。
- 在这个架构中，包括了运算器（也叫逻辑运算单元，ALU）、控制器（CU）、存储器、输入设备、输出设备等组成部分。运算器和控制器这两个核心功能，都由 CPU 负责承担。



■ CPU 和 GPU

冯 · 诺依曼架构

- 处理流程：数据先存在存储器。然后，控制器会从存储器拿到相应数据，再交给运算器进行运算。运算完成后，再把结果返回到存储器。



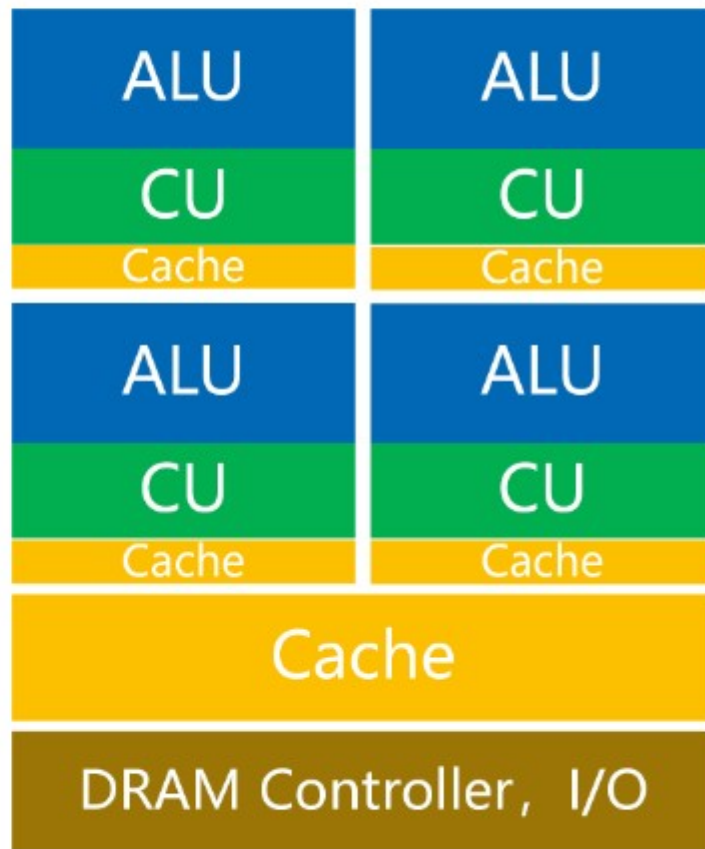
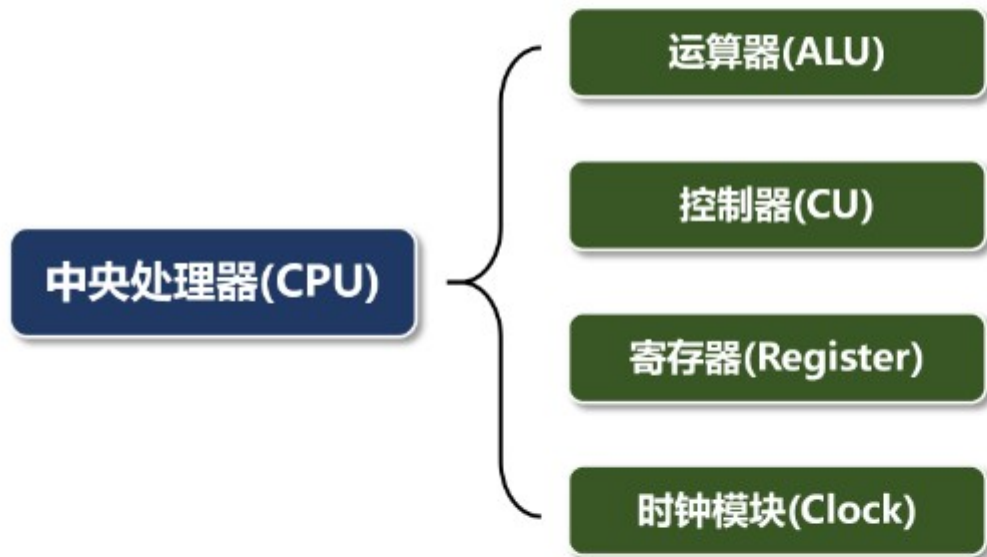
■ CPU 和 GPU

□ CPU 的主要组成

- **运算器 (也叫算术逻辑单元 Arithmetic Logic Unit,ALU):** 包括加法器、减法器、乘法器、除法等 , 负责执行所有数学计算和逻辑判断。
- **控制器 (控制单元 Control Unit):** 负责协调整个 CPU 的操作 , 包括取指、解码、执行和写回四个阶段。 负责从内存中读取指令、解码指令、执行指令。还负责生成各种控制信号来指导其他硬件组件的工作。
- **寄存器 (高速缓存):** 是 CPU 中的高速存储器 , 存储最近使用过的数据或即将使用的指令。通常分为 L1、 L2 和 L3 三级缓存。它的 CPU 与内存 (RAM) 之间的“缓冲”, 速度比一般的内存更快 , 避免内存“拖 累” CPU 的工作。
- **时钟模块 :** 负责管理 CPU 的时间 , 为 CPU 提供稳定的时基。它通过周期性地发出信号 , 驱动 CPU 中的所有 操作 , 调度各个模块的工作。

■ CPU和GPU

□ CPU的主要组成

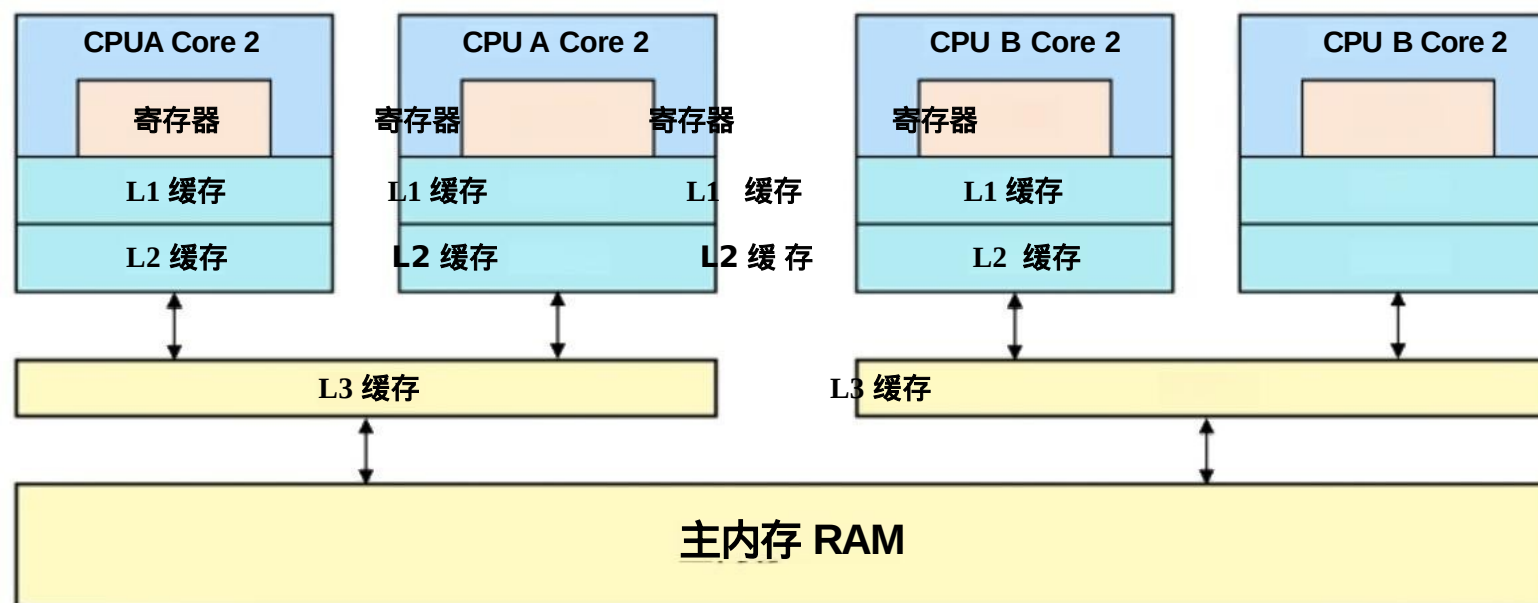


■ CPU 和

GPU

□多核

CPU



**CPU 多核硬件架构示
例**

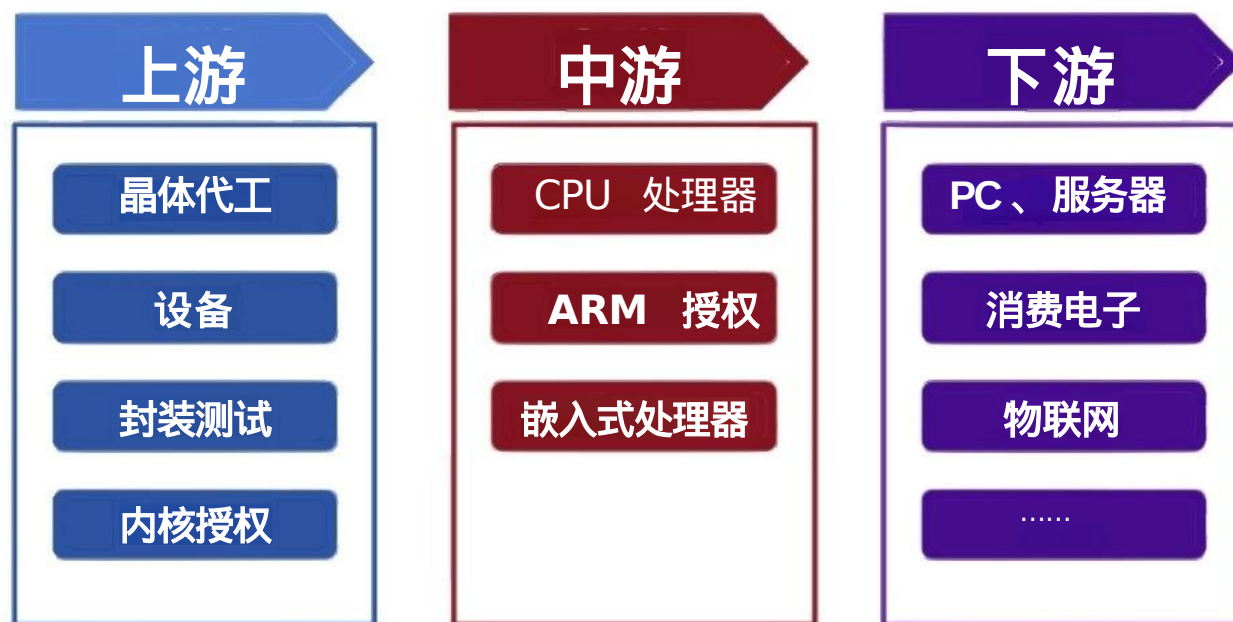
CPU 和 GPU

□ CPU 的类别（按指令集）

指令集名称	类型	推出时间	推出公司 / 机构	采用此指令集主要授权商
x86	CISC 复杂指令集	1978 年	美国 Intel、美国 AMD	兆芯、众志、海光等
ARM	RISC 精简指令集	1985 年	英国 Arm (被日本软银公司收购)	苹果、三星、AMD、TI、东芝、微芯、高通、联发科、展讯、飞腾、海思、瑞芯微、晶晨、全志等
MIPS (终止更新)		1980 年代	美国 MIPS	瑞昱、炬力等
SPARC (终止更新)		1985 年	美国 SUN (被甲骨文收购)	德州仪器、Cypress(被 Infineon 收购)、富士通等
PowerPC (被迫开源)		1991 年	美国 IBM	曾用：苹果、任天堂、微软、索尼、中晟
Alpha (终止更新)		1992 年	美国 DEC (被惠普并购)	申威
RISC-V (开源)		2010 年	美国加州大学伯克利分校 (RISC-V 基金会运营)	GreenWaves、Imagination、平头哥、晶心科技、芯源股份、中天微、睿思芯科、香山处理器

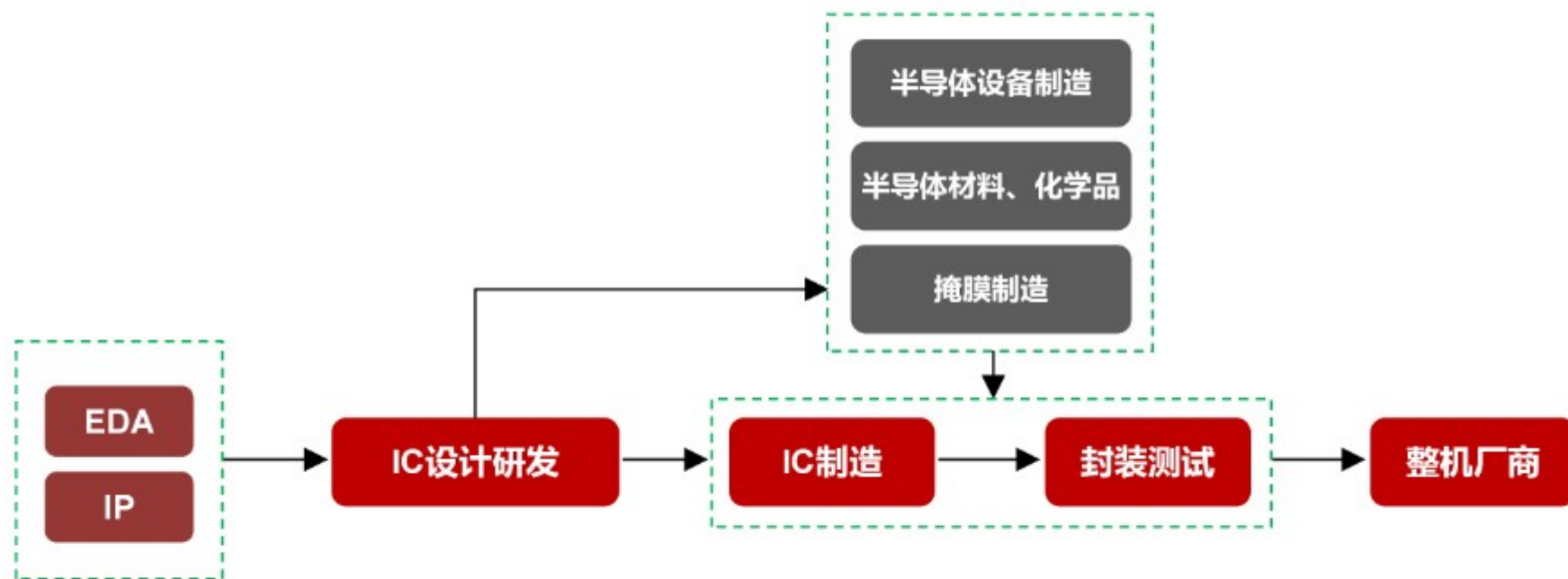
CPU 和 GPU

半导体产业链



■ CPU和GPU

□ 半导体产业链

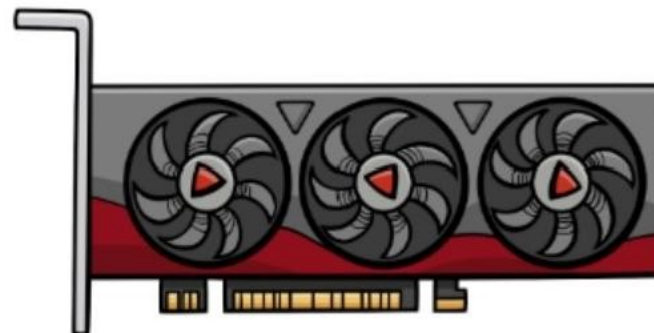


■ CPU 和

GPU

□ GPU 的定义

- GPU, Graphics Processing Unit, 图形处理单元。
- GPU 最初是为了加速计算机图形渲染而设计的专用处理器。它能够高效地执行大量并行计算任务，在 3D 图形渲染、视频编码解码、科学计算等领域表现出色。
- 随着技术的发展，GPU 的应用范围已经远远超出了传统的图形处理，成为通用计算的重要组成部分。



■ CPU 和 GPU

□ GPU 的特点

- 高度并行架构

GPU 拥有成百上千个简单的处理核心，可以同时处理多个数据流，具有强大的并行计算能力。

- 专为浮点运算优化

由于图形渲染涉及大量的矩阵乘法和向量运算，因此 GPU 在浮点运算方面进行了特别优化，提供了比 CPU 更高的吞吐量。

- 内存带宽高

GPU 通常配备有高速的专用显存 (VRAM)，其带宽远高于普通系统内存，这有助于快速读取和写入大量图像数据。

- 低延迟响应

在图形渲染过程中，GPU 需要即时生成每一帧画面，因此它被设计成能够在极短的时间内完成复杂的计算任务。

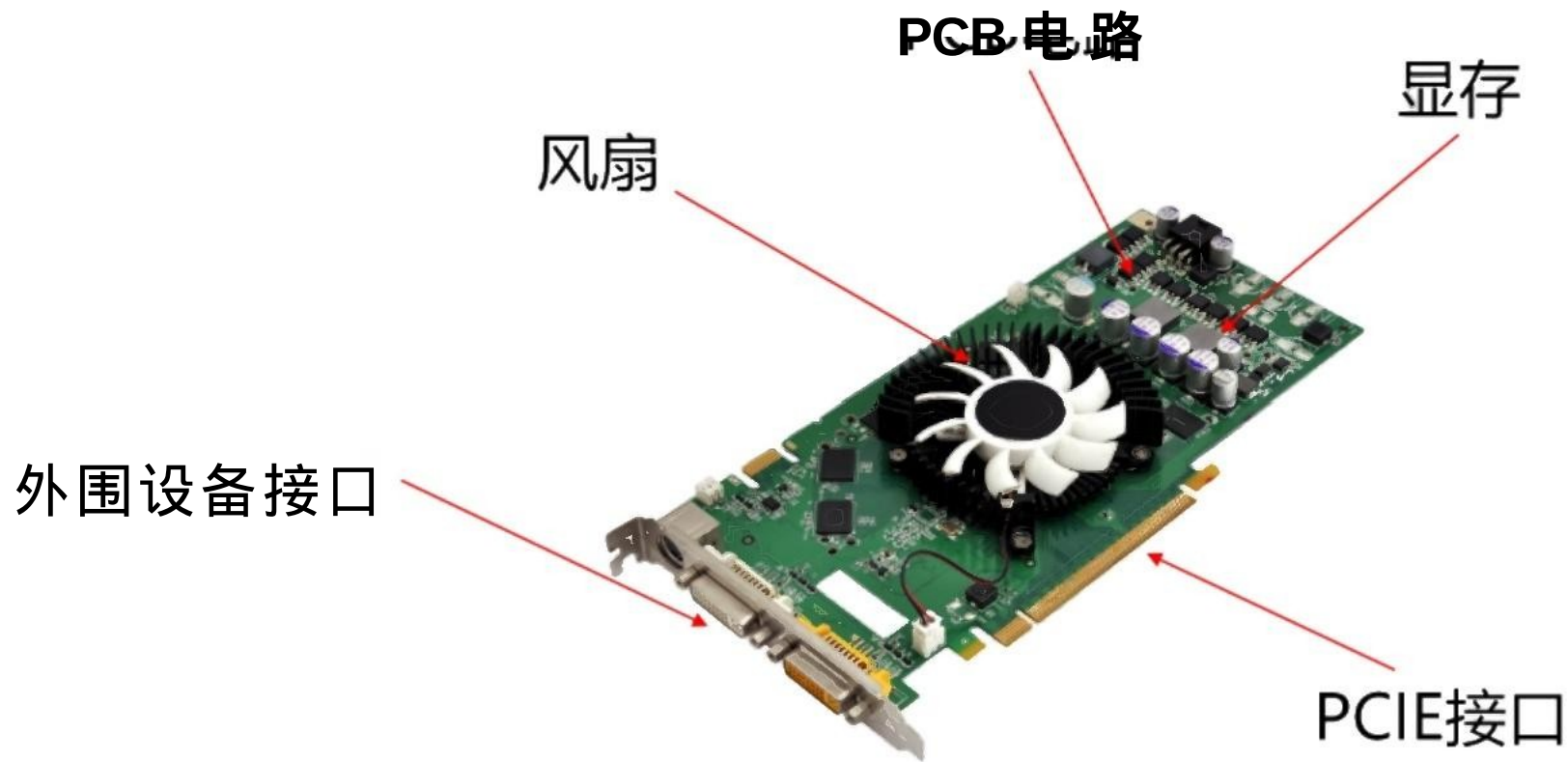
- 支持多种编程接口

现代 GPU 不仅限于使用 OpenGL、DirectX 等图形 API，还广泛支持 CUDA、OpenCL、Vulkan 等通用计算 API。

■ CPU 和 GPU

显卡的组成

- 显卡除了 GPU 之外，还包括显存、VRM 稳压模块、MRAM 芯片、总线、风扇、外围设备接口等。



■ CPU和GPU

□ 图形渲染的流程





CPU 和 GPU

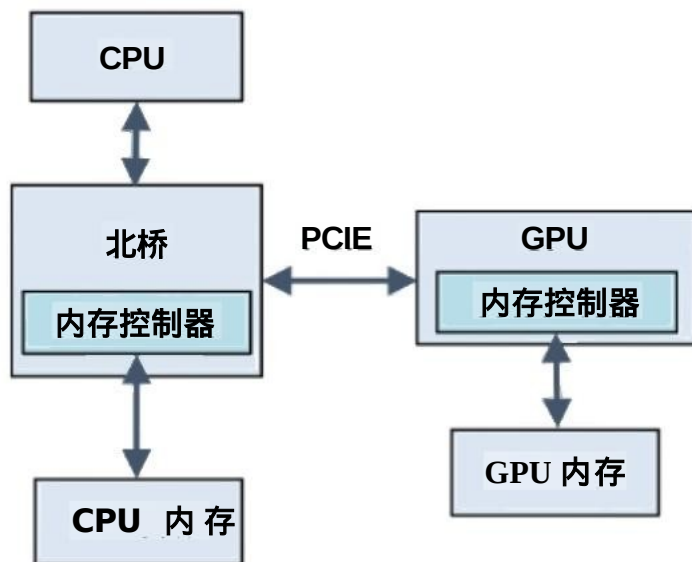
显卡产业的发展历程

- 1962 年，麻省理工学院博士伊凡·苏泽兰 (Ivan Sutherland) 奠定了计算机图形学基础。
- 1984 年，SGI 公司推出了面向专业领域的高端图形工作站，俗称图形加速器，是首个专门的图形处理硬件。
- 1994 年，3DLabs 发布 GLINT300SX，是 PC 最早的 3D 硬件加速图形芯片，从此开启 3D 显卡时代。
- 1995 年，3Dfx 发布 Voodoo 图形加速卡，是真正意义第一款消费级 3D 显卡。
- 1999 年 8 月，NVIDIA（英伟达）公司发布图形芯片 GeForce256，首次提出 GPU 的概念。
- 1999 年，NVIDIA 崛起，击败并收购 3Dfx。
- 2006 年，AMD 以 54 亿美元收购 ATI。

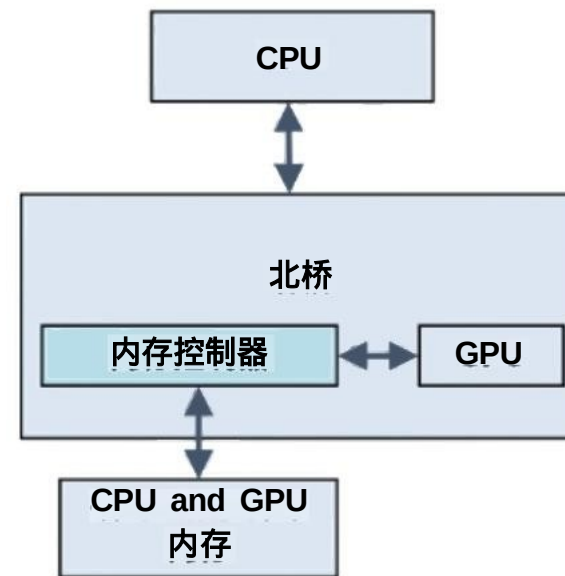
■ CPU 和 GPU

□ GPU （显卡）的分类

- 独立 GPU(dGPU,discrete/dedicated GPU), 常说的独立显卡（独显）。
- 集成 GPU(iGPU,integrated GPU), 常说的集成显卡（集显）。



独立 GPU



集成 GPU

■ CPU 和 GPU

□ GPU （显卡）的主要参数

- 制程：GPU 的制造工艺和设计规则，代表不同电路特性，通常以生产精度 nm 表示
- 图形处理器单元数量：包含了光栅单元 ROP，纹理单元 TMU 的数量，数量越多可执行指令越多
- CUDA 核数：CUDA 是执行函数的重要部件，CUDA 核数越多，性能运行越好
- Tensor 核数：指张量处理单元的数量，Tensor Core 核数越多，性能越好
- 核心频率：指显示核心的工作频率，能反映显示核心的性能优良
- 显存容量：显存容量越大，GPU 能够处理的数据量越大
- 显存位宽：指显存在单位时钟周期内所传送数据的位数，位数越大瞬间传送数据量越大
- 显存带宽：等于显存频率 X 显存位宽 /8, 与显存频率、位宽成正比
- 显存频率：反映显存速度，以 MHz 为衡量单位，越高端的显存，频率越高

■ CPU 和 GPU

■ 英伟达消费级显卡经典型号



时间	发布型号	制程
1995	STG-2000X	500nm
1998	RIVA 128	350nm
1999	Riva TNT2	250nm
1999	GeForce 256	220nm
2001	GeForce 3	180nm
2002	GeForce 4 Ti 4200	150nm
2004	GeForce 6800	130nm
2006	GeForce 8800 GTX	90nm
2010	GeForce GTX 480	40nm
2013	GeForce GTX Titan	28nm
2014	GeForce GTX 970	28nm
2016	GeForce GTX 1080	16nm
2018	GeForce RTX 2080	12nm
2020	GeForce RTX 3090	三星 8nm
2022	GeForce RTX40 系列	台积电 5nm
2025	GeForce RTX 50 系列	台积电 3nm

■ CPU和GPU

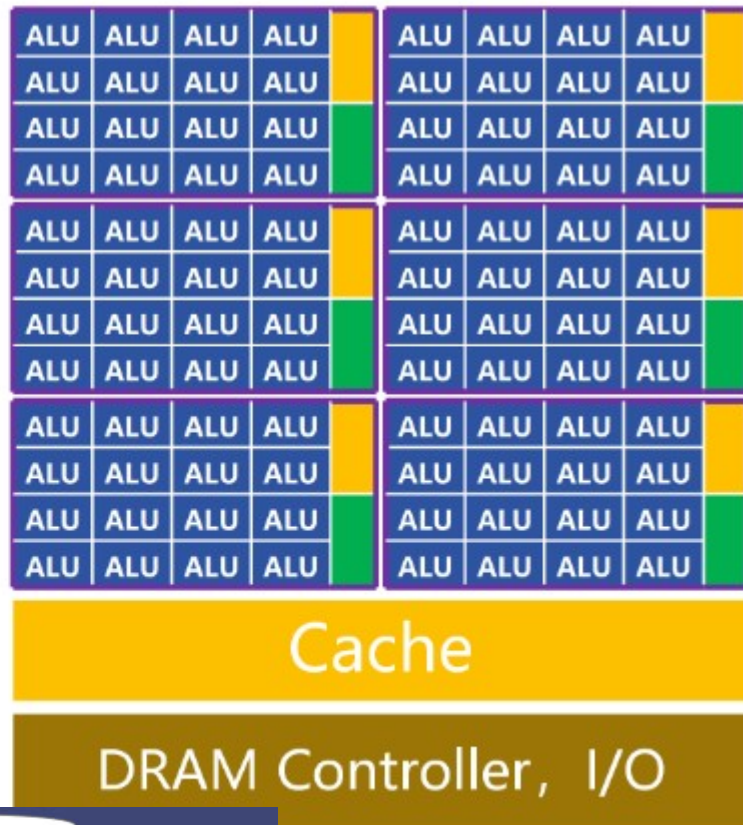
□ GPU（显卡）的组成

- GPU的核，称为流式多处理器（Stream Multi-processor, SM），是一个独立的任务处理单元。
- 在整个GPU中，会划分为多个流式处理区。每个处理区，包含数百个内核。

流式处理区

流式处理区

流式处理区



■ CPU 和 GPU

□ GPGPU

- GPGPU,General Purpose computing on GPU, 基于 GPU 的通用计算。
- GPGPU 利用 GPU 的计算能力，在非图形处理领域进行更通用、更广泛的科学计算。
- GPGPU 在传统 GPU 的基础上，进行了进一步的优化设计，使之更适合高性能并行计算。

	主要执行任务	功能	国内主要公司
GPU	图形渲染	图形渲染、图形计算	景嘉微、摩尔线程、象帝先、芯动科技、 格兰菲、励算、深流微、芯瞳、绘智微
GPGPU	并行计算	AI 相关计算，科学计算和通用计算	壁仞、沐曦、登临、天数智芯、红山微电子、瀚博

CPU 和

GPU

英伟达 GPU 架构 演进

架构代号	中文代号	年代	工艺制程	晶体管数量	代表型号
Tesla	特斯拉	2008	90nm	约 6.84 亿	G80
Fermi	费米	2010	40/28nm	30 亿	Quadro 7000
Kepler	开普勒	2012	28nm	71 亿	K80、K40M
Maxwell	麦克斯韦	2014	28nm	80 亿	M5000、M4000
Pascal	帕斯卡	2016	16nm	153 亿	P100、GTX1080、P6000
Volta	伏特	2017	12nm	211 亿	V100、TiTan V
Turing	图灵	2018	12nm	186 亿	T4、2080TI、RTX 5000
Ampere	安培	2020	7nm	283 亿	A100、A30、3090
Hopper	赫柏	2022	5nm	800 亿	H100
Blackwell	布莱克威尔	2024	5nm	2080 亿	B200、B100

■ CPU 和 GPU

□英伟达 GPU 架构演进

- 2007 年，Tesla 架构：是第一代真正用于并行运算的 GPU 架构，标志用于计算的 GPU 产品线正式独立。
- 2010 年，Fermi 架构：首个完整 GPU 架构，是第一个可支持与共享存储结合纯 cache 层次的 GPU 架构。
- 2012 年，Kepler 架构：首次在 GPU 中引入了动态并行技术。
- 2014 年，Maxwell 架构：可解决视觉计算领域中最复杂的光照和图形难题，优化功耗。
- 2016 年，Pascal 架构：采用了 HBM2 的 CoWoS 技术。首次引入了 3D 内存及 NVLink 高速互联总线。
- 2017 年，Volta 架构：首次引入 Tensor（张量）运算单元。
- 2018 年，Turing 架构：架构最大的变革，引入了 RTX 追光技术总线。
- 2020 年，Ampere 架构：包含 540 亿个晶体管，大幅提升了人工智能和高效能运算。
- 2022 年，Hopper 架构：第一个真正的异构加速平台，适用于高性能计算 (HPC) 和 AI 工作负载。
- 2024 年，Blackwell 架构：包含 100 亿个晶体管，性能提升 25 倍。

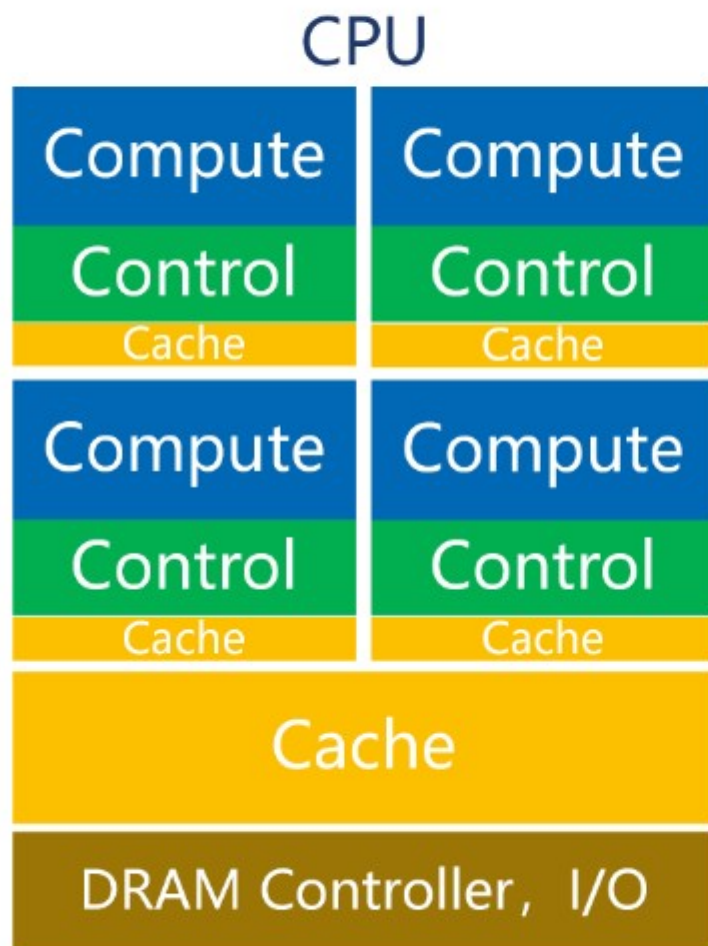
■ CPU 和 GPU

口部分国产 GPU 设计厂商及产品

厂商名称	代表型号
景嘉微电子	JM5 系列、JM7 系列和 JM9 系列
壁仞科技	BR100
摩尔线程	MTT S60、MTT S80、MTT S3000
燧原科技	邃思 2.0
寒武纪 -U	思元 220、思元 290、思元 370 等
沐曦集成电路	MXN 系列 GPU(曦思)、MXC 系列 GPU(曦云)、MXG 系列 GPU(曦彩)
昆仑芯	昆仑芯 2 代芯片
芯动科技	风华 1 号、风华 2 号

■ CPU和GPU

□ CPU和GPU的区别



■ CPU 和 GPU

□ CPU 和 GPU 的区别

- CPU 的内核 (包括了 ALU) 数量比较少，最多只有几十个。但是，CPU 有大量的缓存 (Cache) 和复杂的 控制器 (CU)。
- CPU 是一个通用处理器。作为计算机的主核心，它的任务非常复杂，既要应对不同类型的数据计算，还要 响应人机交互。复杂的条件和分支，还有任务之间的同步协调，会带来大量的分支跳转和中断处理工作。
- CPU 需要更大的缓存，保存各种任务状态，以降低任务切换时的时延。它也需要更复杂的控制器，进行逻辑控制和调度。
- CPU 的强项是管理和调度。真正干活的功能，反而不强 (ALU 占比大约 5%~20%)。

■ CPU 和 GPU

□ CPU 和 GPU 的区别

- GPU 的内核数，远远超过 CPU, 可以达到几千个甚至上万个（也因此被称为“众核”）。
- GPU 为图形处理而生，任务非常明确且单一。它要做的，就是图形渲染。图形是由海量像素点组成的，属于类型高度统一、相互无依赖的大规模数据。
- GPU 的任务，是在最短的时间里，完成大量同质化数据的并行运算。所谓调度和协调的“杂活”，反而很少。
- GPU 的控制器功能简单，缓存也比较少。它的 ALU 占比，可以达到 80% 以上。
- 虽然 GPU 单核的处理能力弱于 CPU，但是数量庞大，非常适合高强度并行计算。同等晶体管规模条件下，它的算力，反而比 CPU 更强。

■ CPU 和

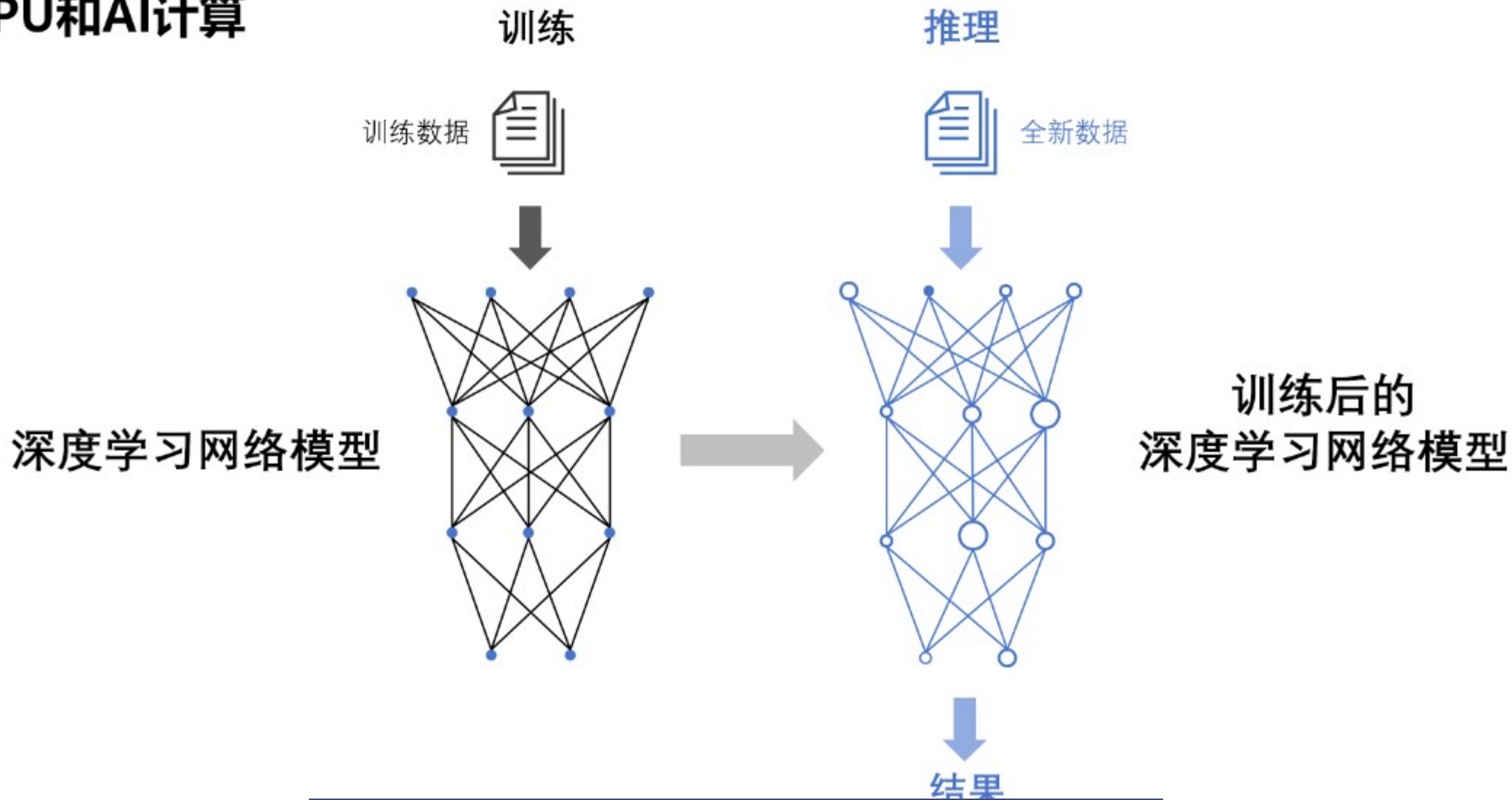
GPU

□ GPU 和 AI 计算

- AI 计算和图形计算一样，也包含了大量的高强度并行计算任务。
- 深度学习是目前最主流的人工智能算法，包括训练 (training) 和推理 (inference) 两个环节。
- 在训练环节，通过投喂大量的数据，训练出一个复杂的神经网络模型。在推理环节，利用训练好的模型，使用大量数据推理出各种结论。
- 它们所采用的具体算法，包括矩阵相乘、卷积、循环层、梯度运算等，分解为大量并行任务，可以有效缩短任务完成的时间。
- GPU 凭借自身强悍的并行计算能力以及内存带宽，可以很好地应对训练和推理任务，已经成为业界在深度学习领域的首选解决方案。
- 目前，大部分企业的 AI 训练，采用的是英伟达的 GPU 集群。如果进行合理优化，一块 GPU 卡，可以提供相当于数十甚至上百台 CPU 服务器的算力。

■ CPU和GPU

□ GPU和AI计算



CPU 和

GPU

英伟达主流 AI 算卡参数

项目	A100	H100	L40S	H200
架构	Ampere	Hopper	AdaLovelace	Hopper
发布时间	2020	2022	2023	2024
FP16 TensorCore	312 TFLOP	756.5 TFLOPS	366.5 TFLOPS	1979 TFLOPS
INT8 TensorCore	624 TOPS	1513 TOPS	733 TOPS	3958 TOPS
FP64	9.7 TFLOPS	34 TFLOPS	25.7 TFLOPS	34 TFLOPS
FP32	19.5 TFLOPS	67 TFLOPS	91.6TFLOPS	67 TFLOPS
GPU 内存	80 GB HBM2e	80 GB	48 GB GDDR6, 带有 ECC	141 GB HBM3e
GPU 内存带宽	2.039 Tbps	3.35 Tbps	0.864 Tbps	4.8 Tbps
最高 TDP	400 W	700 W	350W	700 W
互联技术	NVLink:600GB/s PCIeGen4:64GB/s	NVLink:900GB/s PCIeGen5:128GB/s	PCIeGen4x16: 64GB/s bidirectional	NVLink:900GB/s PCIe Gen5:128GB/s

■ CPU 和 GPU

□ CUDA

- CUDA(Compute Unified Device Architecture) 是英伟达推出的一种并行计算平台和编程模型。
- CUDA 通过提供一系列的工具、库和 API, 使开发人员可以编写能够在 NVIDIA GPU 上高效运行的代码。
- CUDA 广泛应用于科学计算、机器学习、深度学习、图像处理、视频编码等多个领域。
- CUDA 和 cuDNN(CUDA 深度神经网络库) 已经成为训练复杂神经网络不可或缺的一部分。



NVIDIA.
CUDA

PART 03

ASIC和FPGA

■ ASIC 和

FPGA

□ ASIC 的定义

- ASIC(Application Specific Integrated Circuit, 专用集成电路), 就是一种专用于特定任务的芯片。
- ASIC 的官方定义：应特定用户的要求，或特定电子系统的需要，专门设计、制造的集成电路。
- ASIC 芯片面向专项任务，计算能力和计算效率都严格匹配于任务算法。
- ASIC 芯片的核心数量、逻辑计算单元和控制单元比例，以及缓存等，都是精确定制的。



■ ASIC 和

FPGA

□ ASIC 的优点

- 高性能：由于 ASIC 是为特定应用定制的，它们可以在这些应用中提供比通用芯片更高的性能。
- 低功耗：ASIC 可以通过优化电路设计来降低功耗，这在移动设备和嵌入式系统中尤为重要。
- 高成本效益：尽管 ASIC 的设计和制造初期成本较高，但单位成本会随着产量增加而显著下降。
- 小尺寸：ASIC 可以集成更多功能于更小的芯片面积内，有助于减小产品体积。
- 高安全性：由于 ASIC 是专为特定用途设计的，它能够提供更强的安全特性，防止逆向工程和攻击。

■ ASIC 和

FPGA

□ ASIC 的缺点

- 成本高昂，技术难度大。对芯片进行定制设计，对一家企业的研发技术水平要求极高，且耗资极为巨大。
- 研发一款 ASIC 芯片，首先要经过代码设计、综合、后端等复杂的设计流程，再经过几个月的生产加工以及封装测试，才能拿到芯片来搭建系统。
- 研发 ASIC 需要“流片 (Tape-out)”。像流水线一样，通过一系列工艺步骤制造芯片，就是流片。简单来说，就是试生产。14nm 工艺，流片一次需要 300 万美元左右。5nm 工艺，更是高达 4725 万美元。流片一旦失败，将损耗大量的经费，耽误大量的时间和精力。

■ ASIC 和

FPGA

□ ASIC 的应用领域

- 消费电子产品：例如智能手机中的基带处理器、图像信号处理器等。
- 网络通信：包括路由器、交换机中的交换芯片等。
- 加密货币挖矿：比特币等加密货币的挖矿机常采用 ASIC 来提高哈希率并减少电力消耗。
- 汽车电子：用于高级驾驶辅助系统 (ADAS)、 动力总成控制单元等。
- 医疗设备：如超声波成像仪、心电图仪等需要高效处理大量传感器数据的设备。
- 人工智能 / 机器学习加速器：一些公司开发了专门用于 AI 推理和训练的 ASIC，如 Google 的 TPU（张量处理单元）。

■ ASIC 和

FPGA

□ TPU

- TPU, 全称 Tensor Processing Unit, 张量处理单元。
- TPU 专为加速 TensorFlow 框架中的张量运算而设计，显著提高了深度学习模型训练和推理的速度与效率。
- 所谓“张量 (tensor)”，是一个包含多个数字（多维数组）的数学实体。
- 目前，几乎所有的机器学习系统，都使用张量作为基本数据结构。所以，张量处理单元，我们可以简单理解为“AI 处理单元”。



■ ASIC 和

FPGA

□ TPU

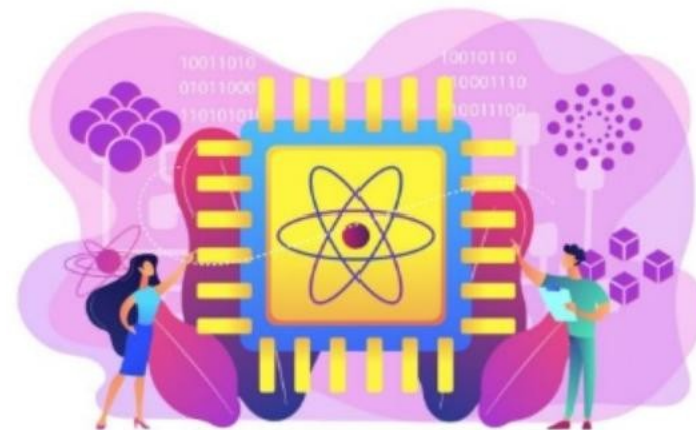
- 2015 年，为了更好地完成自己的深度学习任务，提升 AI 算力，Google 推出了一款专门用于神经网络训练的芯片，也就是 TPU v1。
- 相比传统的 CPU 和 GPU，在神经网络计算方面，TPU v1 可以获得 15~30 倍的性能提升，能效提升更是达到 30~80 倍，给行业带来了很大震动。
- 2017 年和 2018 年，Google 又推出了能力更强的 TPU v2 和 TPU v3，用于 AI 训练和推理。
- 2021 年，Google 推出了 TPU v4，采用 7nm 工艺，晶体管数达到 220 亿，性能相较上代提升了 10 倍，比英伟达的 A100 还强 1.7 倍。

ASIC 和

FPGA

及其它 ASIC

- 除了 Google 之外，还有很多头部企业这几年也在研发 ASIC。
- 英特尔公司在 2019 年底收购了以色列 AI 芯片公司 Habana Labs, 2022 年，发布了 Gaudi 2 ASIC 芯片。
- 2022 年底，IBM 研究院发布了 AI ASIC 芯片 AIU。
- 三星早几年也推出过 ASIC，当时做的是矿机专用芯片。



■ ASIC 和

FPGA

□ NPU

- NPU,Neural Processing Unit, 神经网络处理单元。
- 在电路层模拟人类神经元和突触，并用深度学习指令集处理数据。
- NPU 专门用于神经网络推理，能够实现高效的卷积、池化等操作。
- 经常集成在手机 SoC 芯片中，提供端侧的 AI 计算能力。



■ ASIC 和

FPGA

□ DPU

- DPU,Data Processing Unit, 数据处理单元。
- 被设计用于加速数据中心内的特定任务，特别是那些与网络、存储和安全相关的任务。
- 用于卸载、加速和隔离关键基础设施服务，从而提高效率、性能和安全性。
- 部分 DPU 制造商和技术：
 - 英伟达 BlueField DPU
 - Fungible DPU
 - 英特尔 Infrastructure Processing Units(IPU)

ASIC 和

FPGA

华为昇腾

- 华为昇腾 (Ascend) 系列处理器是华为自主研发的 AI 芯片，也属于 ASIC 芯片。
- 采用了先进的架构和制程技术，提供卓越的计算性能。
- 支持多种主流的 AI 框架，如 TensorFlow、PyTorch、Caffe 等，并且与华为自己的 MindSpore 框架紧密集成。

芯片	昇腾 910 Ascend910	昇腾 310 Ascend310
功能	训练	推理
架构	达芬奇	达芬奇
工艺	7nm	12nm
算力	INT8640TOPS FP16320TFLOPS	INT822TOPS FP1611TFLOPS
功耗	310W	8W
内存	HBM2E	2*LPDDR4x

■ ASIC 和 FPGA

□ FPGA 的定义

- FPGA, 英文全称 Field Programmable Gate Array, 现场可编程门阵列。
- FPGA 是在 PAL(可编程阵列逻辑)、 GAL (通用阵列逻辑) 等可编程器件的基础上发展起来的产物，属于一种半定制电路。
- 简单来说，FPGA 就是可以重构的芯片。它可以根据用户的需要，在制造后，进行无限次数的重复编程，以实现想要的数字逻辑功能。



■ ASIC 和

FPGA

□ FPGA 的组成部分

- 三种可编程电路：

- **可编程逻辑块 (Configurable Logic Blocks, CLB)**: 最重要的部分，是实现逻辑功能的基本单元，承载主要的电路

功能。它们通常规则排列成一个阵列 (逻辑单元阵列, LCA, Logic Cell Array), 散布于整个芯片中。

- **输入 / 输出模块 (I/O Blocks, IOB)**: 主要完成芯片上的逻辑与外部引脚的接口，通常排列在芯片的四周。

- **可编程互连资源 (Programmable Interconnect Resources, PIR)**: 提供了丰富的连线资源，包括纵横网状连线、

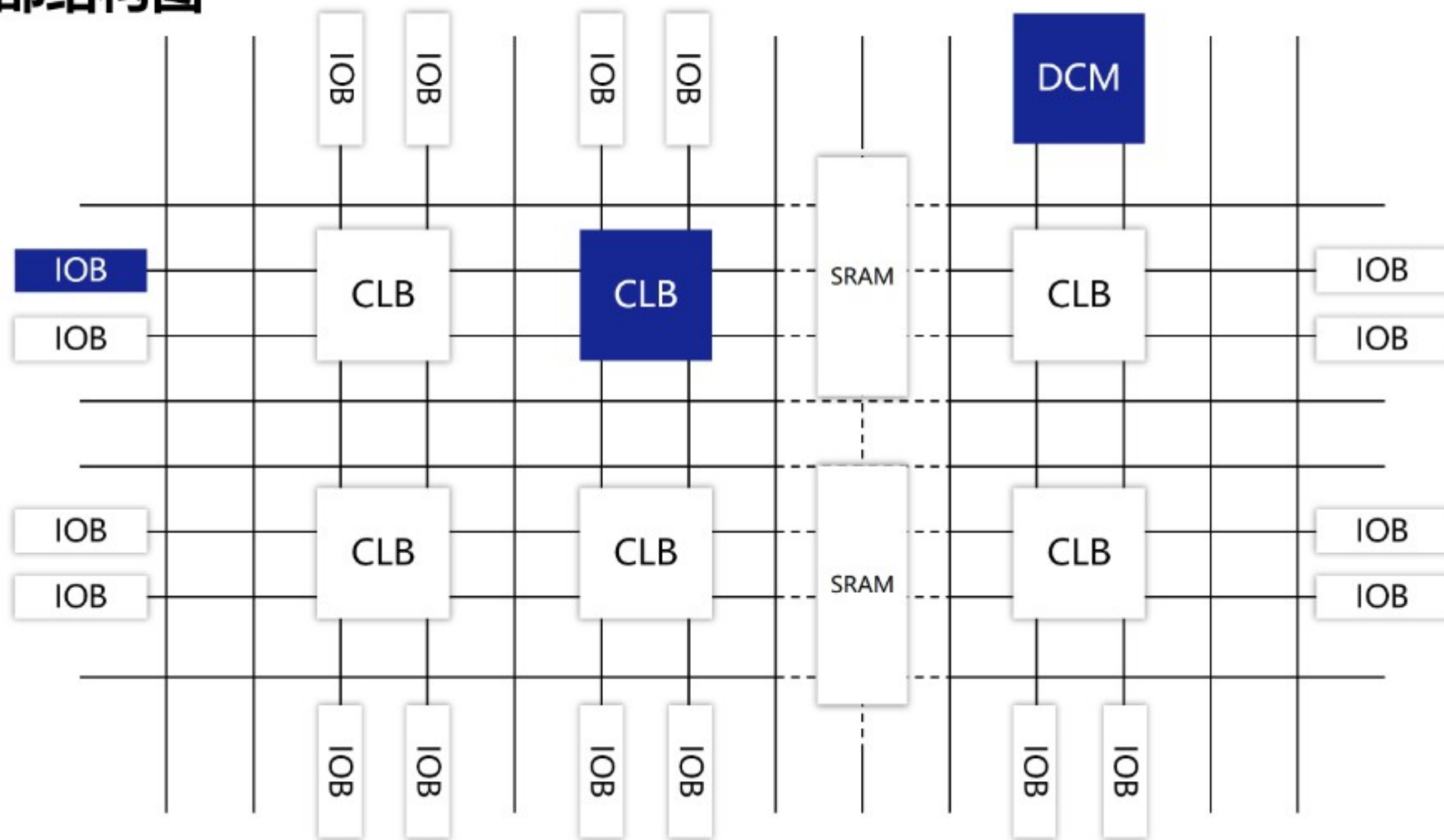
可编程开关矩阵和可编程连接点等。它们实现连接的作用，构成特定功能的电路。

- 静态存储器 SRAM:

- 用于存放内部 IOB、CLB 和 PIR 的编程数据，并形成对它们的控制，从而完成系统逻辑功能。

■ ASIC和FPGA

□ FPGA内部结构图



■ ASIC 和

FPGA

□ FPGA 的组成部分

- CLB 本身，又主要由查找表 (Look-Up Table, LUT)、多路复用器 (Multiplexer) 和触发器 (Flip-Flop) 构成。它们用于承载电路中的一个逻辑“门”，可以用来实现复杂的逻辑功能。
- 我们可以把 LUT 理解为存储了计算结果的 RAM。当用户描述了一个逻辑电路后，软件会计算所有可能的结果，并写入这个 RAM。每一个信号进行逻辑运算，就等于输入一个地址，进行查表。LUT 会找出地址对应的内容，返回结果。这种“硬件化”的运算方式，显然具有更快的运算速度。
- FPGA 的逻辑单元功能在编程时已确定，属于用硬件来实现软件算法。对于保存状态的需求，FPGA 中的寄存器和片上内存 (BRAM) 属于各自的控制逻辑，不需要仲裁和缓存。

■ ASIC 和

FPGA

□ FPGA 的使用

- 用户使用 FPGA 时，可以通过硬件描述语言 (Verilog 或 VHDL)，完成的电路设计，然后对 FPGA 进行“编程”(烧写)，将设计加载到 FPGA 上，实现对应的功能。
- 加电时，FPGA 将 EPROM(可擦编程只读存储器)中的数据读入 SRAM 中，配置完成后，FPGA 进入工作状态。掉电后，FPGA 恢复成白片，内部逻辑关系消失。如此反复，就实现了“现场”定制。
- FPGA 的功能非常强大。理论上，如果 FPGA 提供的门电路规模足够大，通过编程，就能够实现任意 ASIC 的逻辑功能。

■ ASIC 和

□ FPGA 厂商

- 海外四巨头：

- Xilinx 公司 (赛灵思):2020 年，AMD 以 350 亿美元收购了 Xilinx。
- Altera (阿尔特拉):2015 年 5 月，Intel 以 167 亿美元的天价收购了 Altera, 后来收编为 PSG (可编程解决方案事业部)
部门。2023 年 10 月，Intel 宣布计划拆分 PSG 部门，独立业务运营。
- Lattice(莱迪思)
- Microsemi (美高森美)

- 国内厂商：

- 复旦微电、紫光国微、安路科技、东土科技、高云半导体、京微齐力、京微雅格、智多晶、遨格芯等。

■ ASIC 和

FPGA

□ FPGA 产业链



■ ASIC 和

FPGA

□ ASIC 和 FPGA 的区别

-ASIC 和 FPGA, 本质上都是芯片。ASIC 是全定制芯片, 功能写死, 没办法改。而 FPGA 是半定制芯片, 功能

灵活, 可修改性强。

- 类 比 :

·ASIC: 模具玩具。事先要进行开模, 比较费事。一旦开模之后, 就没办法修改了。如果要做新玩具, 就必须重新开模。

·FPGA: 乐高积木。上手就能搭, 花一点时间就可以搭好。如果不满意, 或者想搭新玩具, 可以拆开, 重新搭。



ASIC

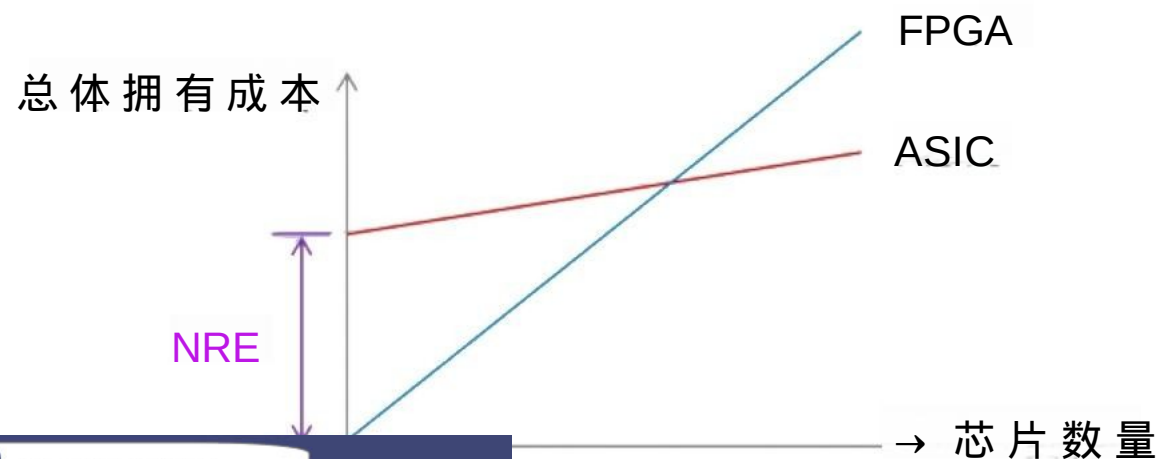


FPGA

ASIC 和 FPGA

□ ASIC 和 FPGA 的区别

- 设计：ASIC 与 FPGA 的很多设计工具是相同的。在设计流程上，FPGA 没有 ASIC 那么复杂，去掉了一些制造过程和额外的设计验证步骤，大概只有 ASIC 流程的 50%-70%。
- 流片：ASIC 需要流片。FPGA 不需要流片。
- 开发周期：开发 ASIC，可能需要几个月甚至一年以上的时间。开发 FPGA，只需要几周或几个月的时间。
- 成本：FPGA 可以在实验室或现场进行预制和编程，不需要一次性工程费用 (NRE)。但是，作为“通用玩具”，它的成本是 ASIC(压模玩具)的 10 倍。如果生产量比较低，那么，FPGA 会更便宜。如果生产量高，ASIC 的一次性工程费用被平摊，那么，ASIC 反而便宜。



■ ASIC 和

FPGA

□ ASIC 和 FPGA 的区别

- 性能和功耗：作为专用定制芯片，ASIC 比 FPGA 强。
- FPGA 是通用可编辑的芯片，冗余功能比较多。无论怎么设计，都会多出来一些部件。
- FPGA 和 ASIC，不是简单的竞争和替代关系，而是各自的定位不同。
- FPGA 现在多用于产品原型的开发、设计迭代，以及一些低产量的特定应用。它适合那些开发周期必须短的产品。FPGA 还经常用于 ASIC 的验证。
- ASIC 用于设计规模大、复杂度高的芯片，或者是成熟度高、产量比较大的产品。

■ ASIC 和

FPGA

□ FPGA 的应用场景

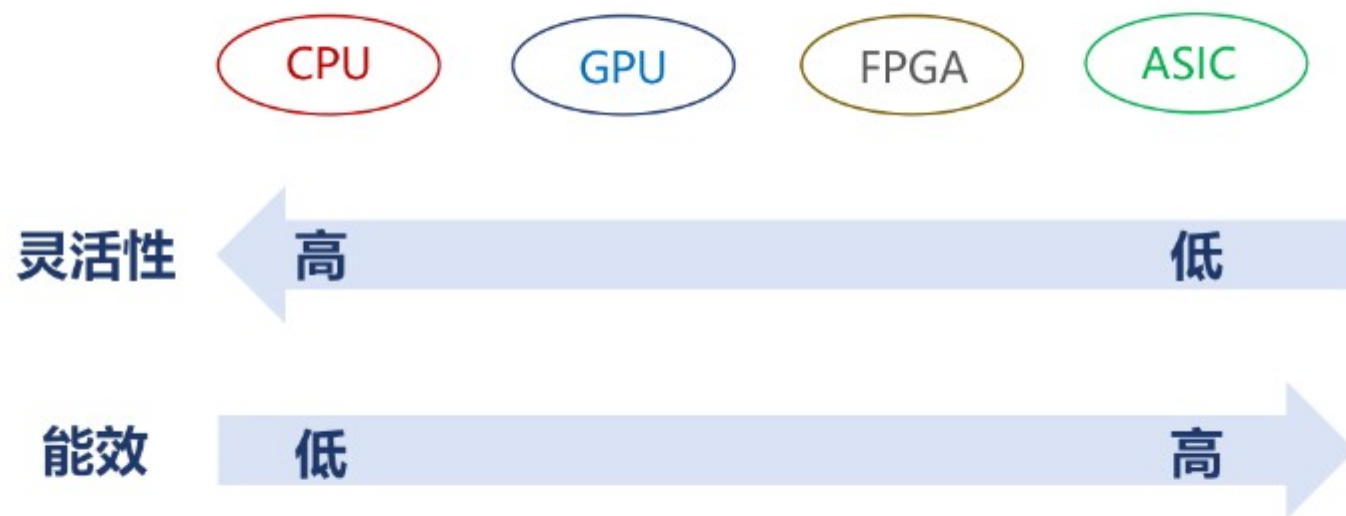
- FPGA 特别适合初学者学习和参加比赛。现在很多大学的电子类专业，都在使用 FPGA 进行教学。
- 从商业化的角度来看，FPGA 的主要应用领域是通信、国防、航空、数据中心、医疗、汽车及消费电子。
- FPGA 在通信领域用得很早。很多基站的处理芯片（基带处理、波束赋形、天线收发器等），都是用的 FPGA。核心网的编码和协议加速等，也用到它。数据中心之前在 DPU 等部件上，也用。后来，很多技术成熟了、定型了，通信设备商们就开始用 ASIC 替代，以此减少成本。

PART 04

总结对比

■ 总结对比

□ 整体对比



■ 总结对比

□ 整体对比

	CPU	GPU	FPGA	ASIC
定制化程度	通用	半通用	半定制化	全定制化
灵活性	高	高	高	低
成本	较低	高	较高	低
功耗	较高	高	较高	低
主要优点	通用性最强	计算能力强 生态成熟	灵活强较高	能效最高
主要缺点	并行算力弱	功耗较大 编程难度较大	峰值计算能力弱 编程难度较难	研发时间长 技术风险高
应用场景	较少用于 AI	云端训练和推理	云端推理 终端推理	云端训练和推理 终端推理

■ 总结对比

□ 整体对比

- 从理论和架构的角度，ASIC 和 FPGA 的性能和成本，肯定是优于 CPU 和 GPU 的。
- CPU、GPU 遵循的是冯·诺依曼体系结构，指令要经过存储、译码、执行等步骤，共享内存在使用时，要经历仲裁和缓存。
- 而 FPGA 和 ASIC 并不是冯·诺依曼架构（是哈佛架构）。以 FPGA 为例，它本质上是无指令、无需共享内存的体系结构。



冯诺依曼架构



哈佛架构

■ 总结对

比

口运算单元对比

- 从 ALU 运算单元占比来看，GPU 比 CPU 高，并行计算效率更高。
- FPGA 因为几乎没有控制模块，所有模块都是 ALU 运算单元，比 GPU 更高。

■ 总结对

比

功耗对比

- GPU 的功耗极高，单片可以达到 250W，甚至 600W(RTX5090)。而 FPGA 一般只有 30~50W。
- 这主要是因为内存读取。GPU 的内存接口 (GDDR5、HBM、HBM2) 带宽极高，大约是 FPGA 传统 DDR 接口的 4-5 倍。但就芯片本身来说，读取 DRAM 所消耗的能量，是 SRAM 的 100 倍以上。GPU 频繁读取 DRAM 的处理，产生了极高的功耗。
- 另外，FPGA 的工作主频 (500MHz 以下) 比 CPU、GPU (1~3GHz) 低，也会使得自身功耗更低。FPGA 的工作主频低，主要是受布线资源的限制。有些线要绕远，无法支持更高的时钟频率。

■ 总结对

比

口时延对比

- GPU 时延高于 FPGA。
- GPU 通常需要将不同的训练样本，划分成固定大小的 "Batch（批次）"，为了最大化达到并行性，需要将数个 Batch 都集齐，再统一进行处理。
- FPGA 的架构，是无批次 (Batch-less) 的。每处理完成一个数据包，就能马上输出，时延更有优势。

■ 总结对比

□目前 GPU 在 AI 芯片占比较大的主要原因

- 在英伟达的长期努力下，GPU 的核心数和工作频率一直在提升，芯片面积也越来越大，算力非常强劲。
- 功耗方面，GPU 依赖先进的工艺制程，以及水冷等被动散热，可以勉强支撑。
- 生态方面，英伟达推出的 CUDA 编程模型及其相关库（如 cuDNN）已经成为行业标准，极大地简化了开发者编写高效 GPU 代码的过程。几乎所有的主流深度学习框架（TensorFlow、PyTorch 等）都内置了对 GPU 的支持，这进一步促进了其普及。
- 普及性方面，由于 GPU 已经广泛存在于个人电脑、服务器甚至移动设备中，因此基于 GPU 的解决方案更容易被采纳，并且可以利用现有的硬件基础设施。

■ 总结对比

□ AI 芯片发展趋势

- ASIC 芯片加速崛起，提升市场占比。
- 随着摩尔定律逐渐接近极限，业界正在探索新的计算范式（如量子计算、类脑计算）。
- 结合 CPU、GPU、DSP、FPGA 等多种芯片的异构计算加速普及，可以根据不同任务的需求灵活调配资源，实现更高的效率和更低的功耗。
- 端侧推理 AI 芯片发展提速。
- 针对特定行业或应用领域（如自动驾驶、医疗影像、智能家居）的高度定制化 AI 芯片需求增加。