



企业级 AI 应用开发：从技术选型到生产落地

黛忻

阿里云 Serverless AI 团队



Contents

目录

- 01** 企业级 AI 应用开发运行时选型
- 02** Serverless AI 运行时关键技术
- 03** 客户案例 – Serverless + AI 让应用开发更简单

企业级 AI 应用开发运行时选型

01

AI 原生范式对基础设施提出全新的要求

构建支持 AI Agent 的高效基础设施



Agent-Centric

基础设施的核心服务对象从“人类用户”转变为“自主 Agent”，以 Agent 而非服务或 API 为中心

以 Agent 为中心



State-First

状态是 Agent 的“记忆”与“人格”载体，基础设施必须原生支持状态的持久化、低延迟访问与跨环境迁移

状态优先



Task-Driven Orchestration

基础设施主动协调 Agent 完成目标，而非被动响应请求，Agent 和 Agent 或者 Agent 和工具之间的协作依靠事件驱动和动态弹性

任务驱动协作



Embrace Uncertainty

承认 LLM 输出的非确定性，通过基础设施能力降低风险，而非追求绝对可控，从“防御性编程”转向“容错自愈”

接受不确定性

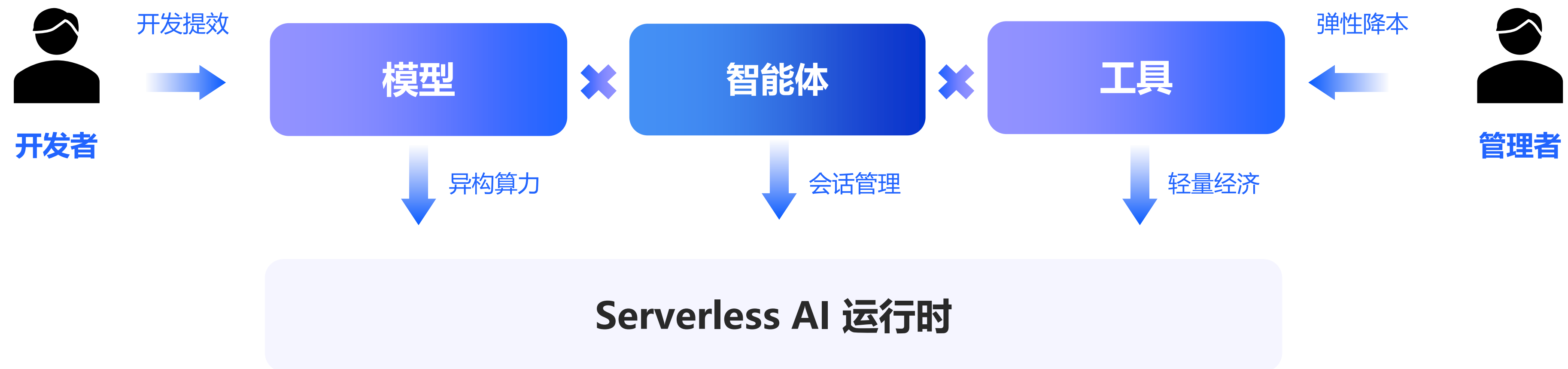
AI 时代开发者关注业务创新而非基础设施

Serverless 是 AI 原生架构的最短实现路径



从 Serverless 到 Serverless AI

Serverless AI 运行时是 AI 原生应用的最佳选择



Serverless AI 运行时关键技术

02

函数计算 FC: Serverless AI 运行时

0 运维、轻量、经济、弹性



100倍



启动效率

冷启动速度：FC 毫秒~秒级，
虚机数分钟，容器 30+秒~数分钟

5倍



规格粒度

最小规格：FC 0.05C128MB，
虚机 1C512MB，容器 0.25C512MB

0

不使用不计费

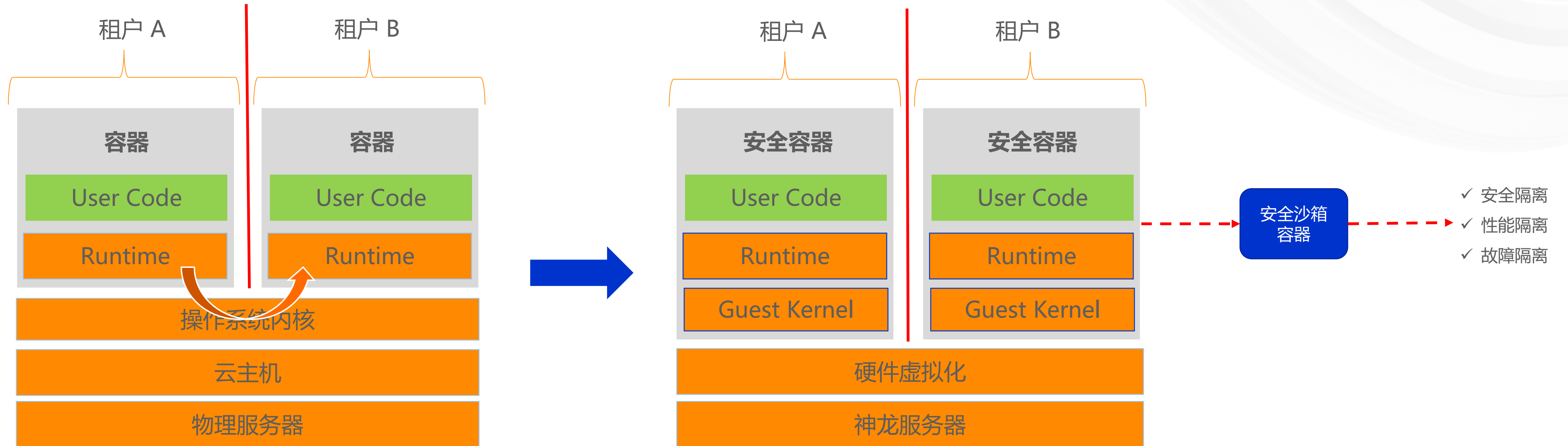
按请求调度，毫/秒计费，低峰自动缩 0
虚机包月浪费多，容器为集群持续付费
FC 不为 3AZ 容灾额外付费，虚机/容器则需额外付费！

50+

内置环境

Python/Node/Java/PHP/Go/.NET 等
50+ 内置运行时环境，支持自定义运行
时和自定义镜像，方便开发者灵活定制

Serverless AI 运行时安全 —— 资源强隔离



传统容器技术

普通容器用内核提供的 namespace 和 cgroup 做资源限制和隔离（从机器上圈了一部分资源给容器用），在安全性上存在不足：

- ❑ 容器内的进程在宿主机上可以看到
- ❑ 容器和宿主机共用内核，可以对宿主机进行破坏

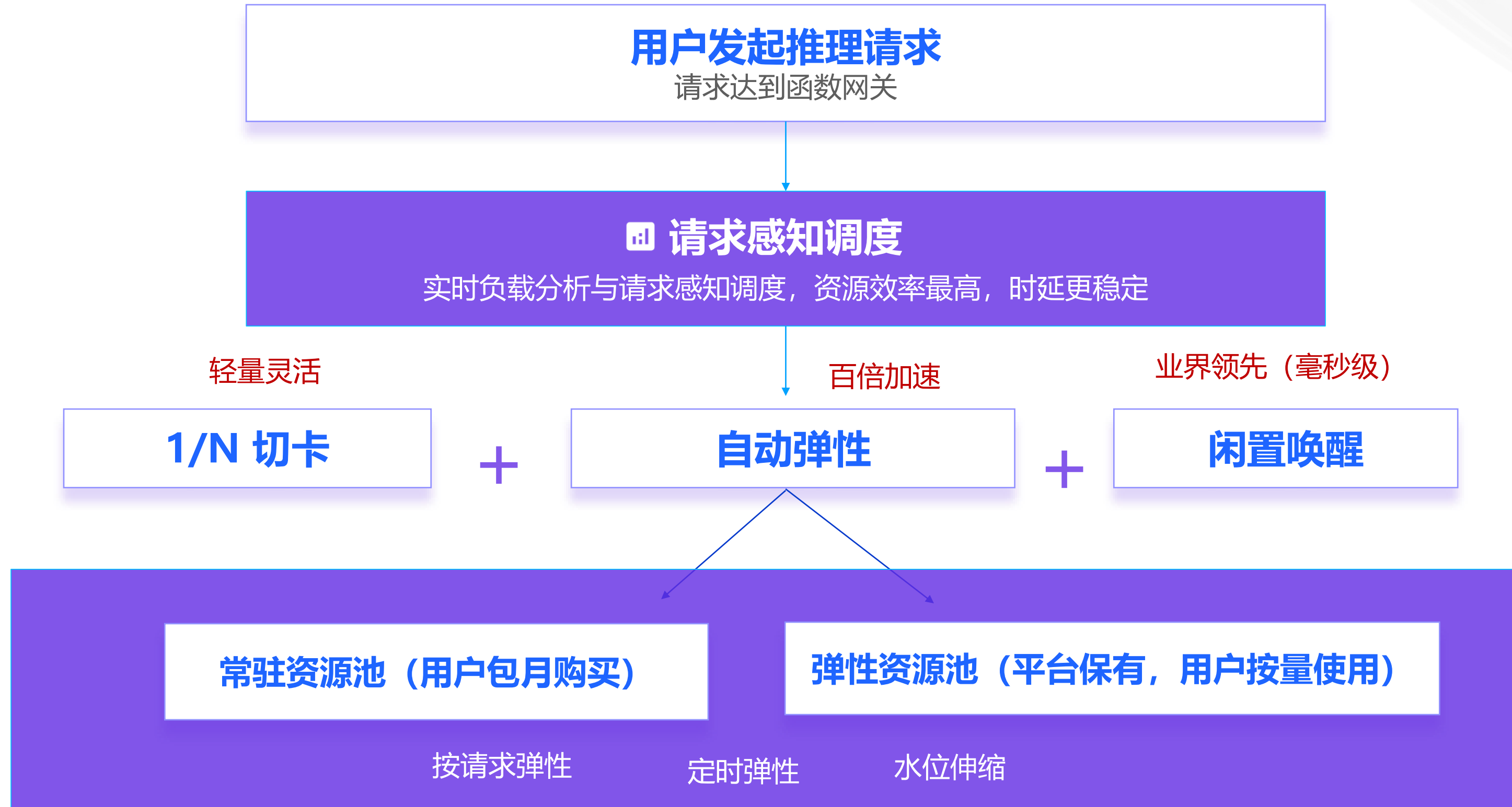
函数计算FC运行时

FC 安全容器安全加固策略（核心是**限制代码破坏范围**）：

- ✓ 安全容器提供基于**虚拟机级别**的隔离
- ✓ 函数调度**尽可能**调度到同一台神龙服务器
- ✓ 加固安全策略：端口封禁、命令行封禁等
- ✓ 组件裁减：精简不必要驱动和内核接口，启动速度更快、资源占用更少
- ✓ 实例回收：销毁重建，避免残留 /tmp目录、日志、环境变量、进程等

模型运行时关键技术

函数计算 Serverless GPU 相对虚拟/容器的核心优势：请求感知调度、毫秒级闲置唤醒、1/N卡切分使用、Serverless 混合调度



请求感知调度



RT 抖动减少 80%，可扩容到0

独家优势，相比 HPA 抖动率明显下降

用户常驻资源池 + 平台弹性资源池 混合调度



Serverless 混合调度降本50%

用户成本降低 50% 以上

毫秒级闲置唤醒



百倍启动加速

以 15% 的成本换取冷启动分钟级->毫秒级

1/N卡切分使用



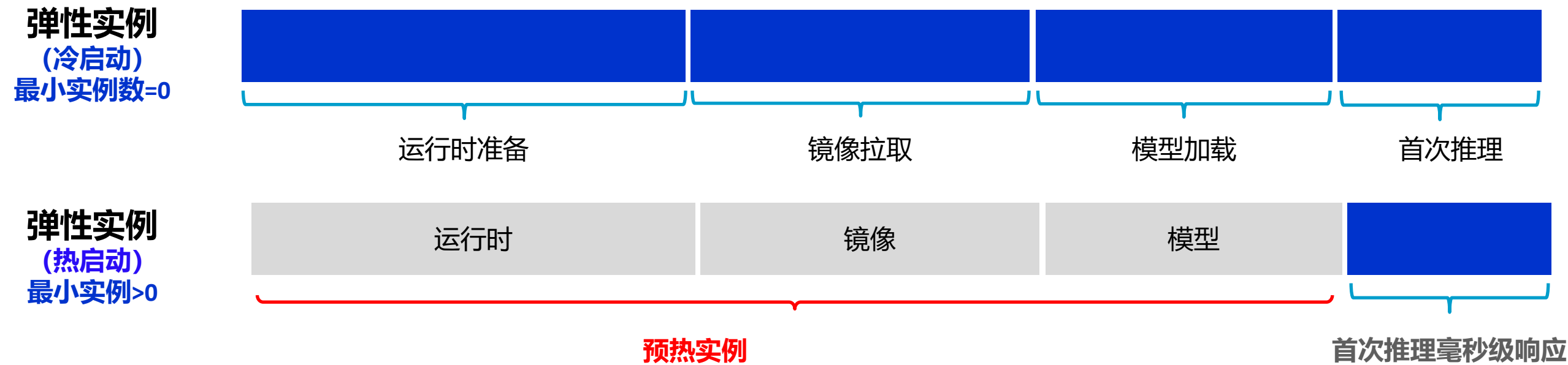
最大降本 93.75%

减少 GPU 碎片，优化稀缺资源利用率

模型运行时：GPU冷启动优化

函数计算首推 Serverless GPU 启动快照，**实现毫秒级的首次推理响应**，0->1 首包耗时对比 K8s GPU，从分钟级优化至毫秒级

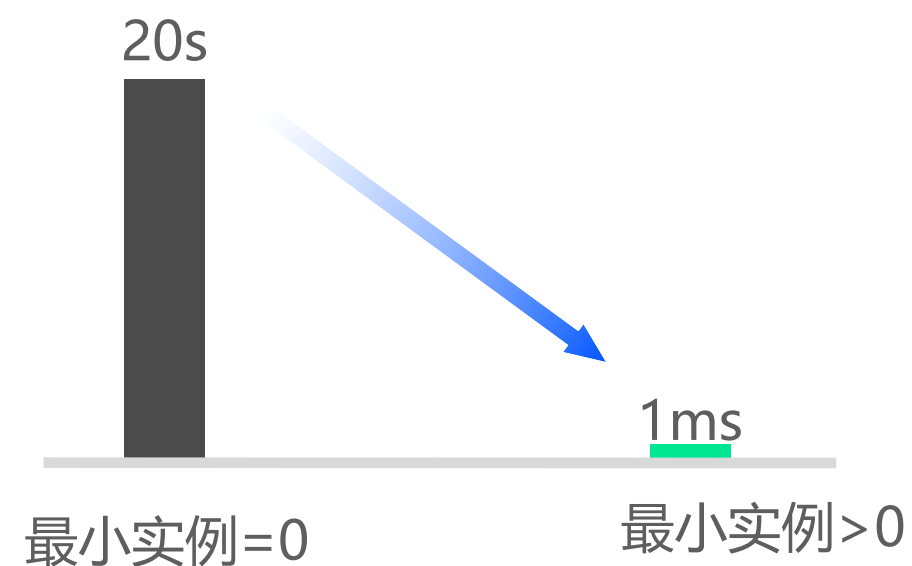
首次推理冷启动耗时分布示意图



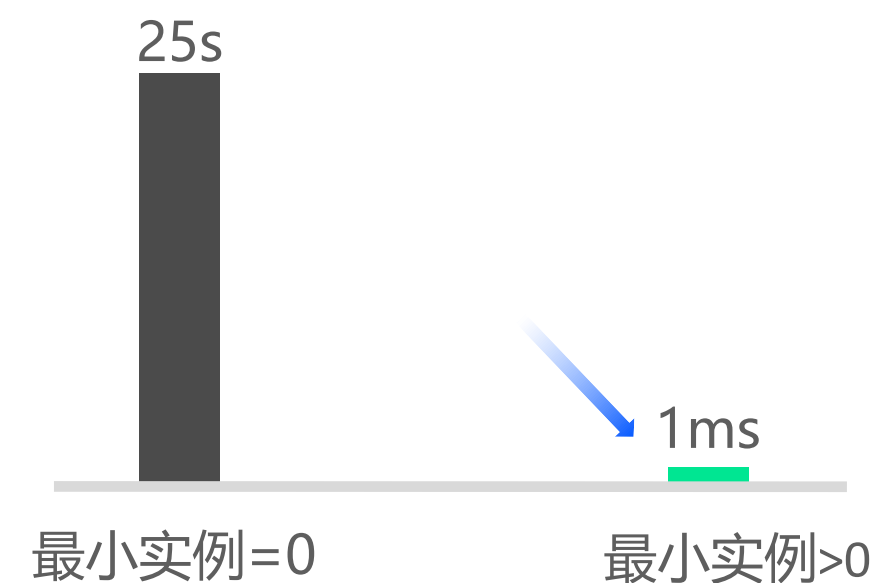
实时/准实时在线推理服务的痛点

1. 低时延：实时/准实时业务时延敏感，一般要求秒级响应，部分场景下需要毫秒级
2. 高并发：高峰期突增的吞吐量可能导致系统性能下降
3. 高成本：低峰期和小规格模型资源浪费，高峰期资源不足，成本优化难
4. 低容错：小流量推理场景单卡容灾能力差，故障率高

SD-v1-5-inpainting (4.27GB) 0->1 TTFI



Qwen-14B-Chat-Int4 (9.01GB) 0->1 TTFI



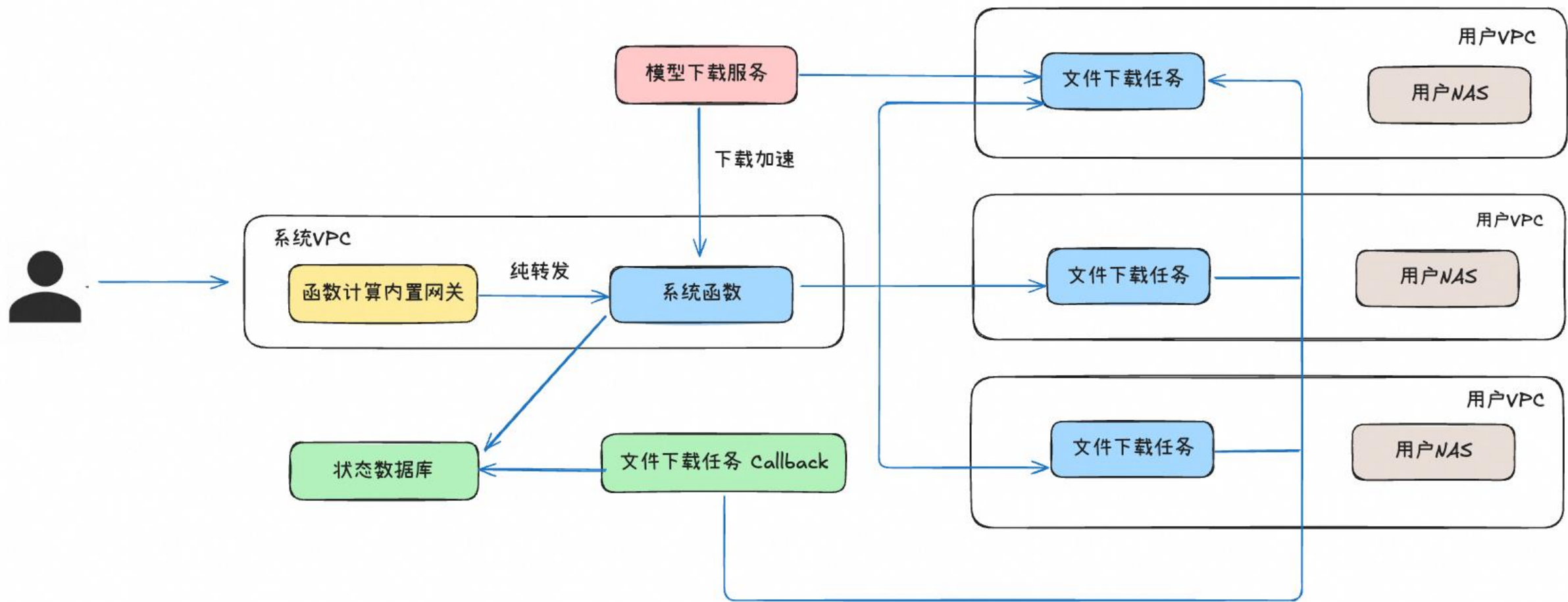
模型运行时：模型加载加速

模型随容器镜像分发

适用场景：<1GB 的传统领域模型（CV / TTS），模型变更频率比较低
模型加载加速方案：镜像加速预热 + P2P 镜像分发

模型下载加速

适用场景：模型文件放在 OSS / NAS，应用程序通过挂载点访问。对模型大小没有限制。
➢ OSS：大量实例并行加载模型、需要本地冗余，或者多地域部署的场景。访问数量较少的大文件。
➢ NAS：需要极速的启动性能。
模型加载加速方案：模型下载加速。函数计算用 OSS 缓存常用的模型，下载服务会自动判断系统是否缓存过，已缓存会走 OSS 内网下载。下载本身通过分片下载，多线程/多函数实例下载做了一些优化。



智能体/工具运行时关键技术

函数计算 FC：沙箱即服务、Session 亲和/隔离架构、毫秒级启动与按需付费

标杆客户	 ModelScope  Qwen  阿里云百炼 行业头部厂商			
沙箱即服务	Code Interpreter	Browser Use	RL Sandbox	Sim Sandbox
服务化 API，支持十万函数百万实例级别的沙箱执行				
智能体运行时	Serverless 级 Session 亲和/隔离架构			
负载感知调度，按会话弹性伸缩，支持会话亲和/会话隔离				
Serverless异构算力	CPU		GPU	
零运维，毫秒级启动，最大支持2w实例/分钟极速交付，免费提供 3AZ 自动容灾				

业界领先的开箱即用、多语言代码安全执行引擎

沙箱即服务

服务化：提供 Code Interpreter API、Browser API

内置开发环境：Python/Node.js/Java/PHP/ Shell/.NET 等 50+ 多语言环境，支持 OCI 标准镜像和自定义运行时灵活扩展

业界首创 Serverless 级 Session 亲和/隔离架构

智能体运行时

开源开放：与 AgentScope、LangChain、LlamaIndex 等主流开发框架集成

毫秒级启动与按需付费：强隔离、突破性上下文保持，启动效率领先传统容器方案 100 倍，按需使用，按量付费，低峰缩 0 成本最优

AI 时代计费演进 —— 从请求驱动到价值驱动

阶段一：按资源租用计费

虚拟机 / 容器 计费模式

问题：为实例的持续运行付费。无请求时不能缩0仍计费，资源空转成本高。

阶段二：按请求计费

传统 FaaS 计费模式

问题：为代码运行时刻付费，无请求时 0 成本。但长连接保活场景（如 MCP Server/WebSocket）因低负载存活仍计费，成本高。

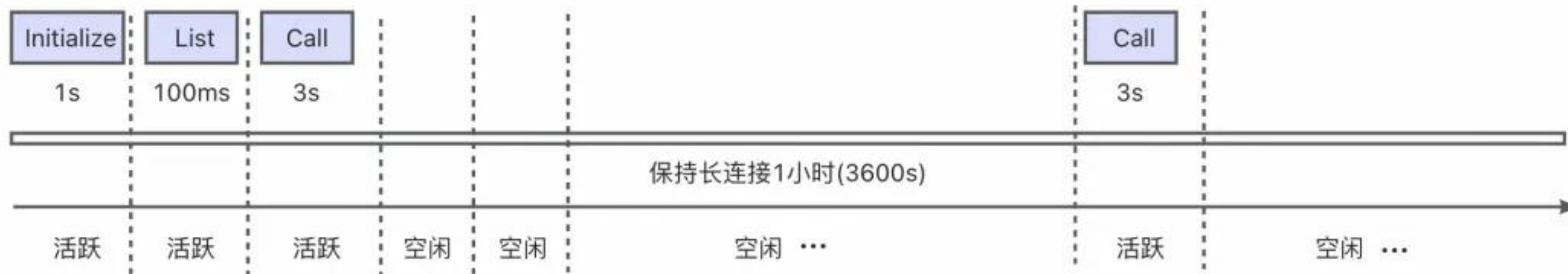
阶段三：按实际资源消耗计费

Serverless AI 计费模式

价值：按实际资源消耗，精准区分忙闲时计费。

消除长会话/低负载保活冗余成本，无缝支持 AI 强交互场景。

MCP Server 基于Serverless AI 的计费方案，长连接闲置计费最高降低 **87%**



Sandbox 实例动态挂载 — 从计算隔离到存储隔离延伸

传统共享存储问题（虚机/容器/FaaS架构）

Agent Code Sandbox 多租户数据共享，有安全问题。无法满足同一个函数的每个实例路径不同的需求。挂载存储路径是变化不确定的。

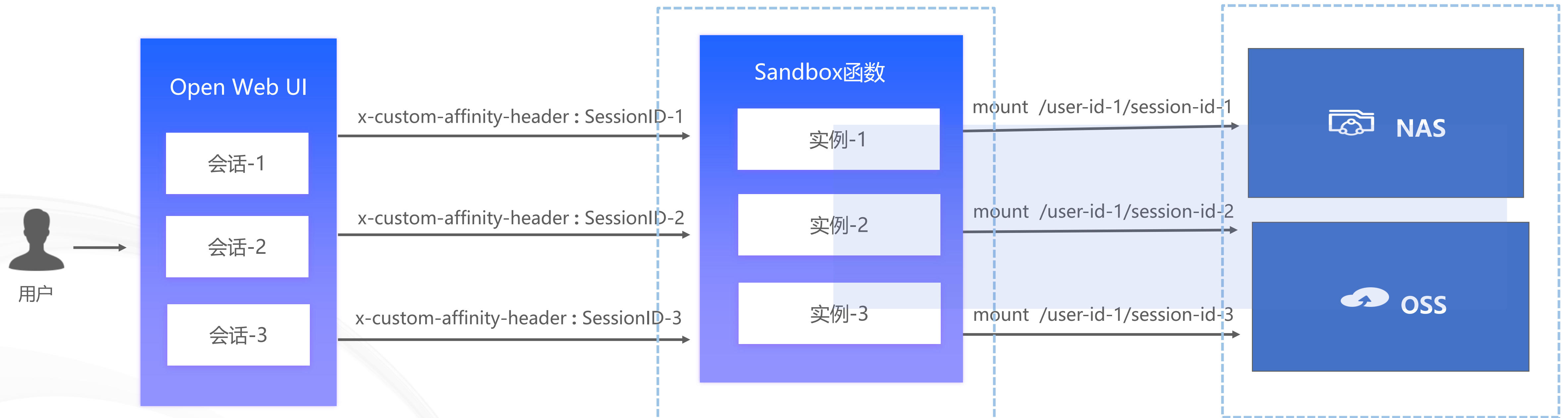
Serverless AI 解决方案

- 引入会话粒度度存储粘性，将会话和一个持久化的，归属特定租户的存储子目录进行强绑定。
- 平台基于POSIX 标准多租存储安全实践框架，落地层次化纵深防御体系



函数计算 FC

持久化存储





阿里内部案例 — 智能体/工具运行时最佳实践

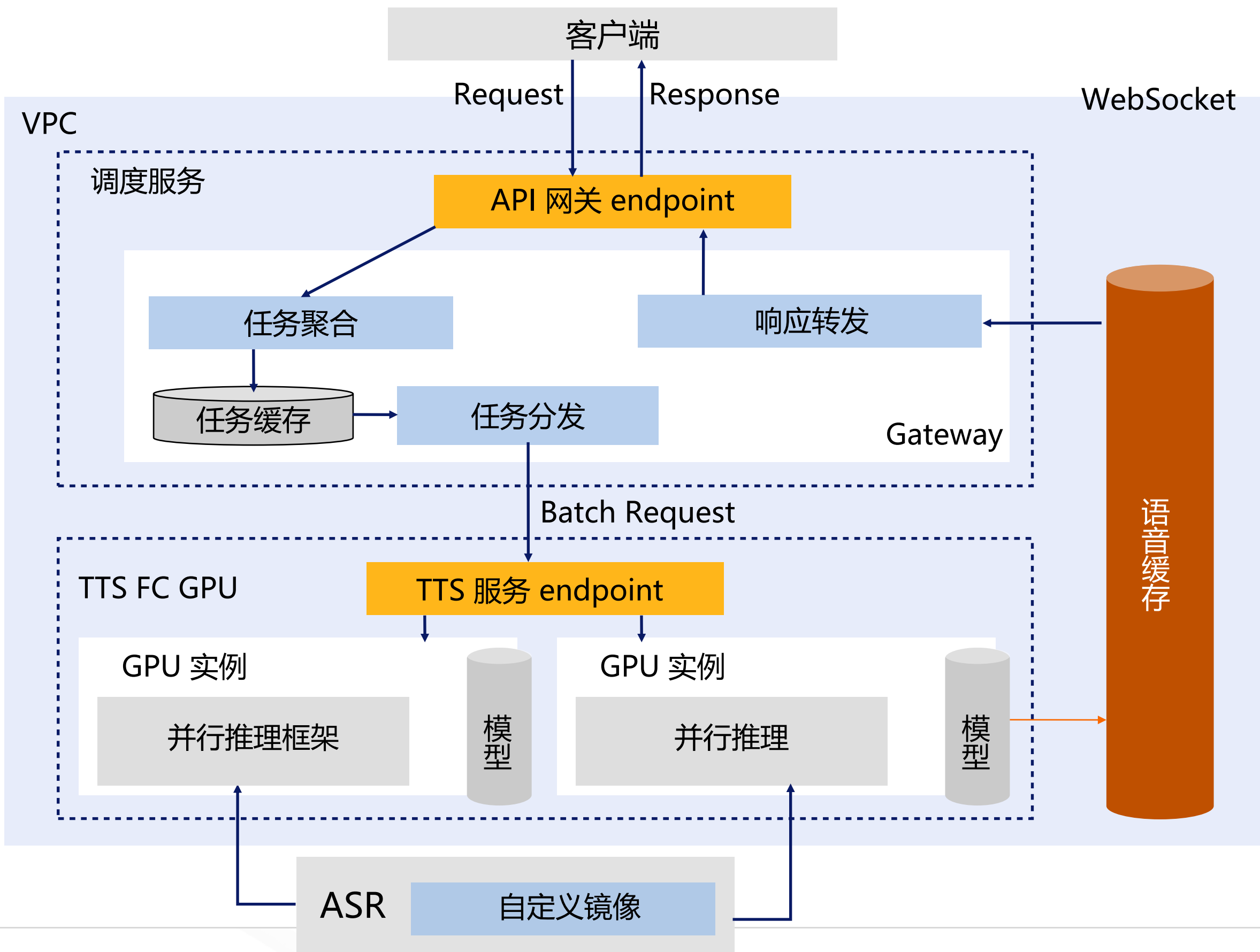
魔搭社区、Qwen、百炼，大规模使用函数计算 FC 提供的 Serverless 运行时构建模型、智能体和 AI 工具



Serverless 运行时已经成为阿里云 AI 原生应用的核心载体

实时/准实时推理场景— Serverless GPU 解决方案

函数计算给吉利 AI 座舱的交互和娱乐功能提供大规模推理服务，共同打造**大规模、高可用、高性能**的推理引擎。
场景覆盖：意图解析、文生图、情感TTS 等。



痛点 & 挑战

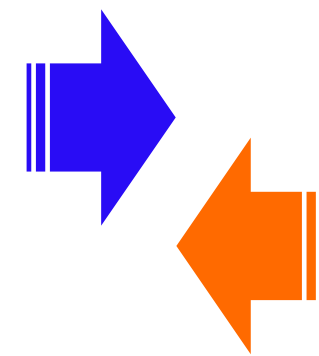
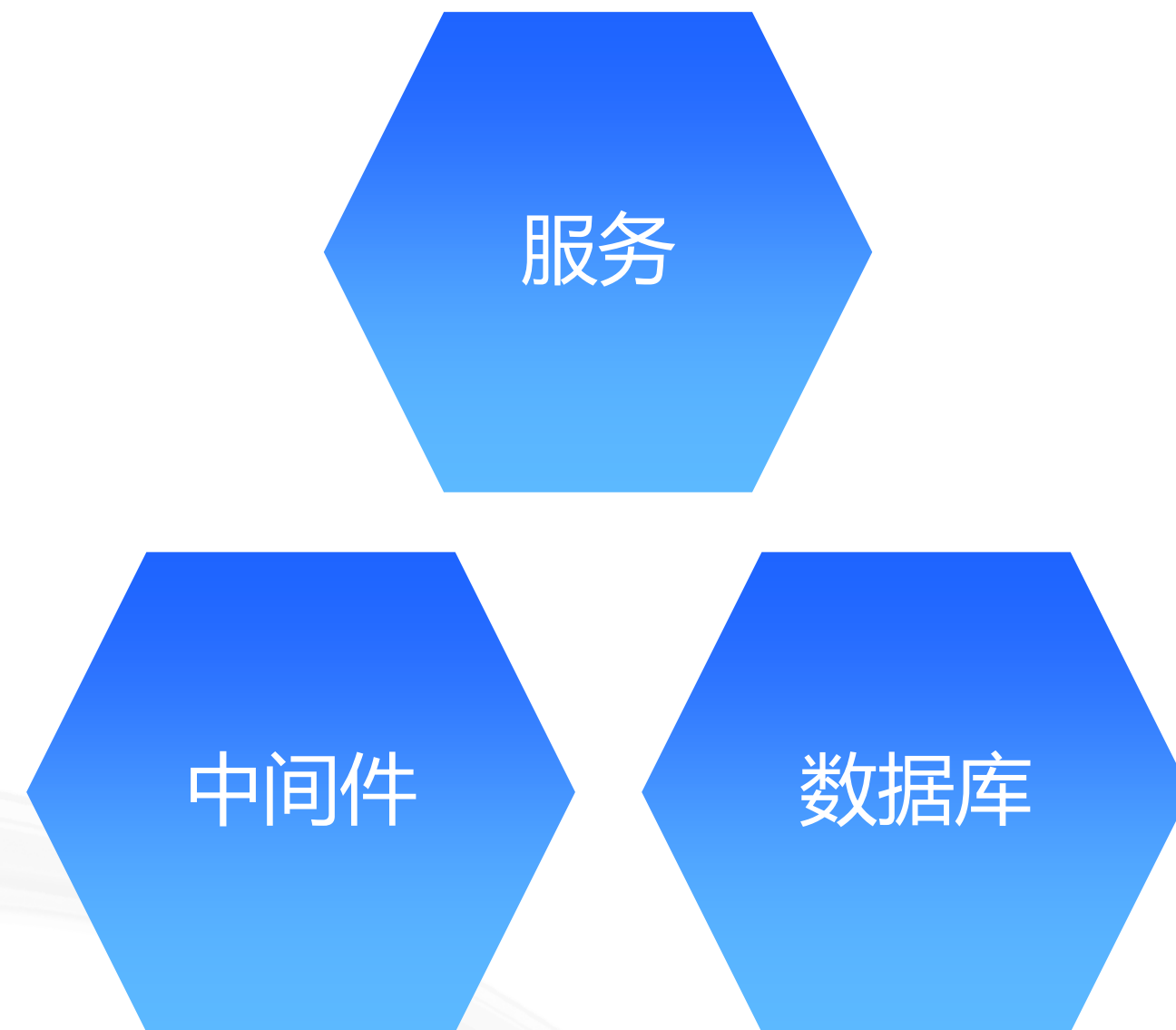
- 高性能：冷启动低延迟 & 模型预热、推理请求批量执行等。
- 低成本：提升 GPU 资源利用率
- 高可用：模型请求高可靠接入、推理服务高可靠、故障恢复。
- 故障恢复策略：快速定位和恢复。

解决方案

THANKS

AI 应用组成核心抽象发生变化

传统应用开发



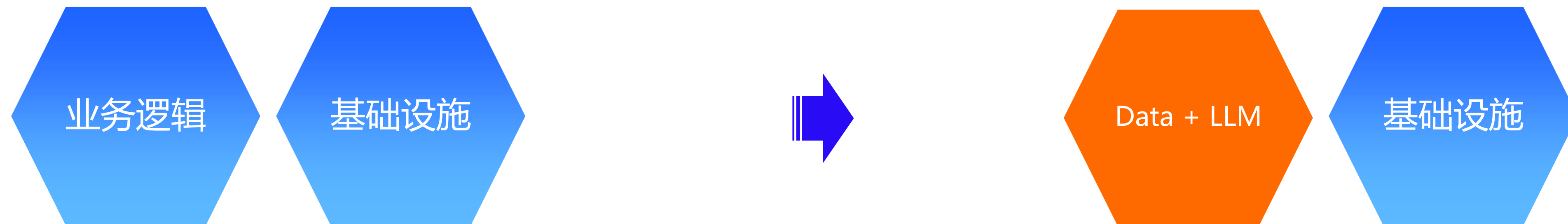
AI 应用开发



AI 应用研发的关注点发生变化

传统应用开发：如何确保业务逻辑正确稳定运行

AI 应用开发：如何最大可靠的发挥LLM价值



AI 业务优先要求基础设施更加敏捷高效

Serverless AI 原生架构新范式

Building a New Paradigm of AI-Native Architecture on Serverless



AI应用可观测：云监控2.0针对AI应用提供全栈智能可观测能力

全栈Serverless

全栈高可用

双层安全



Function AI：从算力到应用，AI 全栈升级

魔搭社区、Qwen、百炼、PAI、Qoder，大规模使用函数计算 FC 提供的 Serverless AI 运行时构建模型、智能体和工具



FunctionAI：让 AI 应用开发更简单

FunctionAI 一键创建应用

Function AI / 探索

返回函数计算控制台 钉钉答疑群 帮助

Serverless + AI，让应用开发更简单！

Function AI 是基于函数计算构建的 Serverless AI 应用开发平台，内置极致的弹性算力、探索丰富的 AI 应用模板、开箱即用的 AI 工程能力与应用的全生命周期管理。让技术不再是束缚，让AI化作创新的画笔，即刻体验，开启您的 Serverless AI 之旅！

筛选应用模板

☐ 热门模板

☐ 人工智能

☐ MCP Server

☐ 音视频处理

☐ Web 开发框架

☐ 文件处理

☐ Web 应用

☐ 游戏

☐ 流程式开发

搜索应用模板

基于 Qwen-QwQ 推理模型构建 AI 聊天助手

官方 AI Ollama Open-WebUI

使用 Qwen-QwQ 系列模型，部署 Ollama 与 Open-WebUI 到阿里云应用开发平台

2025.03.06 更新 272

基于 DeepSeek-R1 构建 AI 聊天助手

社区 大模型 AI Ollama Open-WebUI

部署 Ollama 与 Open-WebUI 到阿里云应用开发平台, 使用 DeepSeek-R1-Distill-Qwen 系列模型

2025.03.04 更新 1055

基于 DeepSeek-R1 构建 RAG 应用

社区 大模型 AI 检索增强生成 模板

CAP rag with deeepseekr1

2025.03.16 更新 173

AI 剧本生成与动画创作

官方 AI 大模型 阿里云技术内容

AI 剧本生成与动画创作示例工程

2025.03.04 更新 707

组装式开发，弹性开放，按需选择

模型托管
FunModel

图像生成
FunArt

Agent 开发
AgentRun

Agent 低代码
AIStudio

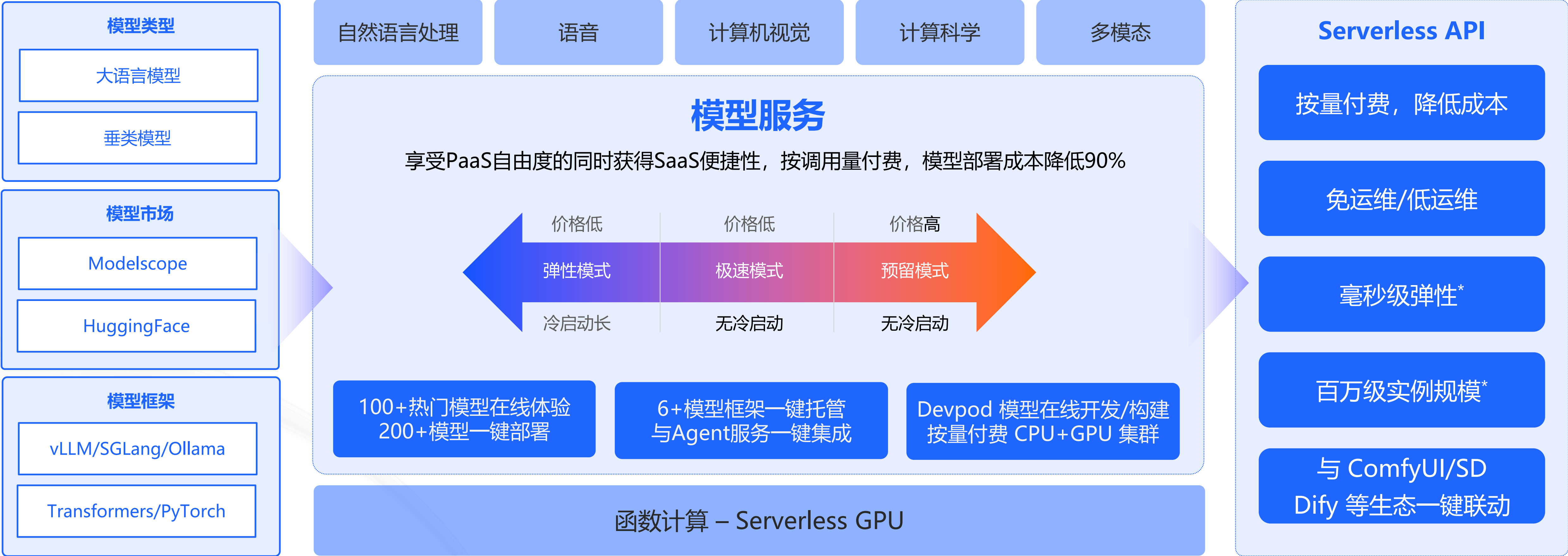
无缝升级 AI 应用开发范式

应用开发

模型服务：AI 模型一键转化为 Serverless API

Model Service: One-Click AI Model to Serverless API Transformation

开源模型一键部署，AI模型一键 Serverless 化，云端模型开发部署零门槛



* 毫秒级弹性和百万规模集群紧针对部分模型的测试结果，并不代表全部模型都可以具备该能力

AgentRun: Agentic AI 应用基础设施

AgentRun: Agentic AI Application Infrastructure

函数计算/Function AI 赋能企业 AI 应用高效开发与稳定运行

开发

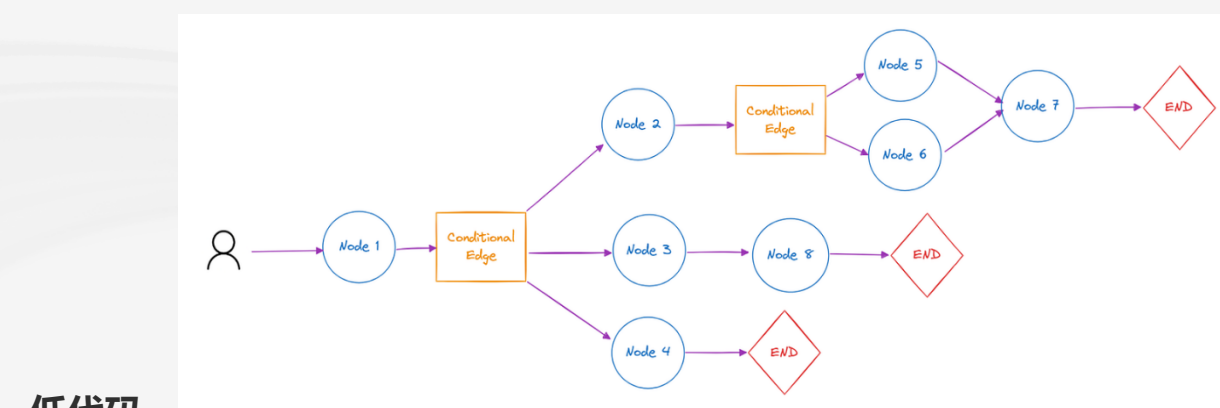
部署

运维

高代码深度定制，低代码快速搭建，全面提升 AI 应用开发效率

轻量、安全隔离、极致弹性的 Serverless AI 运行时

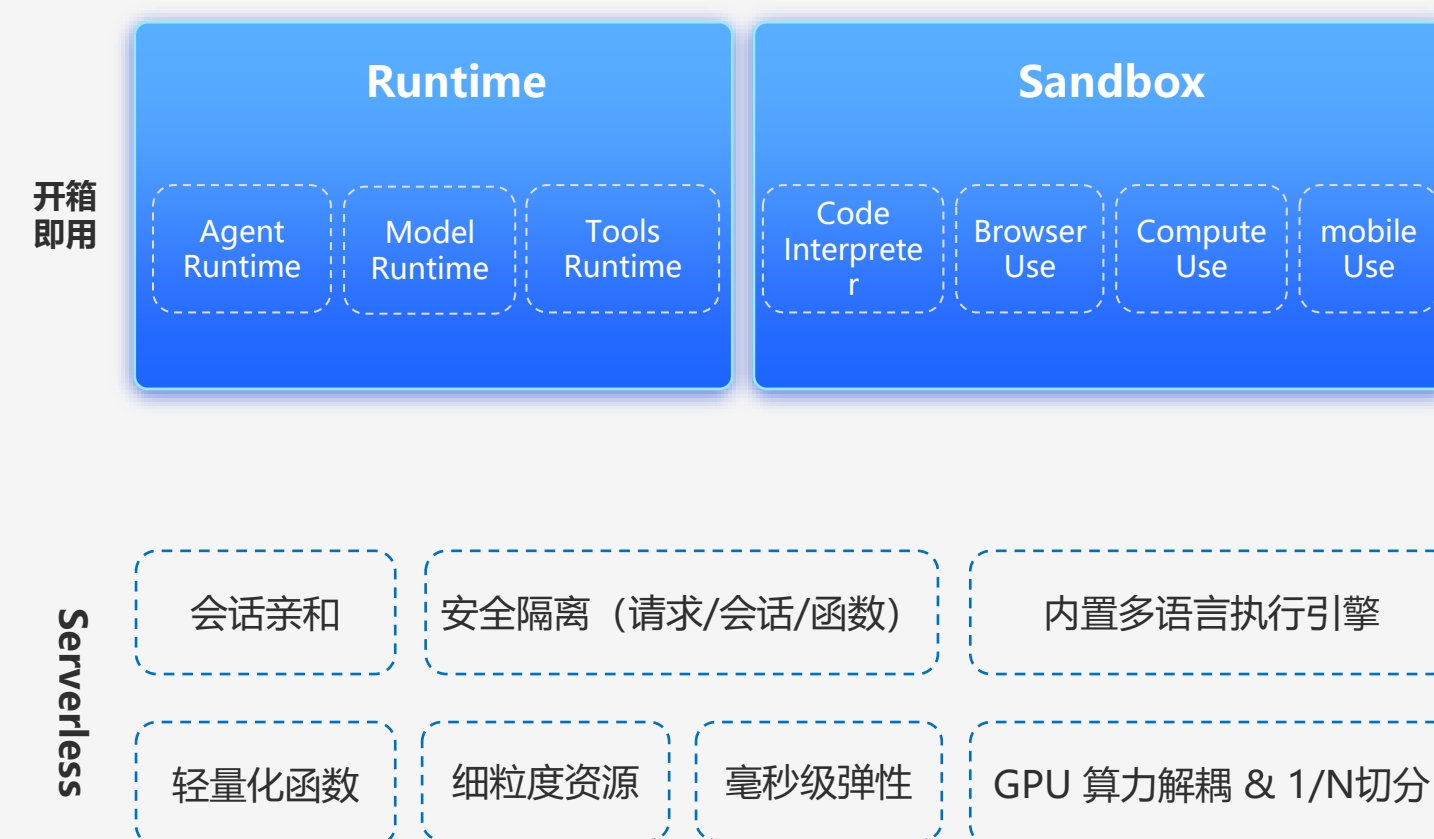
服务治理与可观测，为 AI 应用保驾护航



快速创建

自然语言 AI 生成

流程编排



AIStudio: 高性能低代码 Agent 开发平台



与光同尘 AIGC 案例：业务特点及挑战

业务流量在 平峰 与 高峰期 有十几倍的 Gap



成本难以控制

平常资源消耗相对平稳，传统包月方案资源配置不够灵活，资源闲置率较高



弹性能力要求极高

AIGC 课程高峰期需要几分钟内弹出几百张卡的 GPU 资源，弹性能力要求很高



运维管理投入大

业务发展迅速，人员需要聚焦在业务本身，传统方案运维程序繁琐，大大增加运维的工作量

技术选型

Technology Solution Comparison

一般有 2 种技术选型：模型服务商（如 OpenAI、百炼等）、开源自建（Qwen、DeepSeek、ComfyUI、SD 等）

维度	SaaS（模型服务商）	PaaS（函数计算 Function Compute）	IaaS（VM/容器自建）
安全	✗ 数据风险高： 数据在第三方，合规不可控	✓ 数据风险低： 数据在客户私网，厂商基础安全保障	✓ 数据风险低 数据在客户私网，厂商基础安全保障
效率	✓ 开发效率最高： 开箱即用，零配置 ✗ 几乎不可定制： 无法修改底层框架，可选模型少	✓ 开发效率高： 开箱即用，无需管服务器及其环境依赖 ✓ 定制效率高： 框架/模型自由，开源选择多	✗ 开发效率低： 需手动配置集群、网络、依赖 ✗ 定制效率低： 小规模效率高，大规模效率低
可靠	✓ 可靠性最高： 服务商提供 SLA，自动容灾 ✗ 完全黑盒： 故障依赖服务商修复	✓ 可靠性高： 3AZ 高可用，自动容灾 ✓ 黑盒+白盒： 自带监控日志链路追踪等工具	✗ 可靠性低： 自主实现高可用架构 ✓ 完全白盒： 需要自建监控日志，代价高
弹性	✓ 有限弹性： 按请求弹性，配额受限则不可再弹 ✗ 成本不可控： 按请求单价高，突发流量费用激增	✓ 极致弹性： 按请求弹性，毫秒/秒级供给资源 ✓ 成本可控： 按资源单价低，利用率高浪费少	✗ 普通弹性： 手动/自动扩展 VM 或 Pod，分钟级 ✗ 成本可控： 按资源单价低，利用率低浪费多

基于上页提到的挑战，我们可以发现IaaS基于成本、效率、稳定性、弹性上均不满足我们要求。

与光作为创企，在 PoC 阶段快速验证效果，SaaS / PaaS 的开箱即用无疑是最简单的选择。

随着规模扩大，从百卡 GPU 需求增长到千卡级别，核心业务的自建和创新PaaS无疑是 ROI 最高的选择。

等到业务成熟，有专业团队，PaaS / IaaS 自建



FunctionAI 客户交流群

<https://functionai.console.aliyun.com/welcome>

稳：从架构到防护，全程保障线上稳定运行

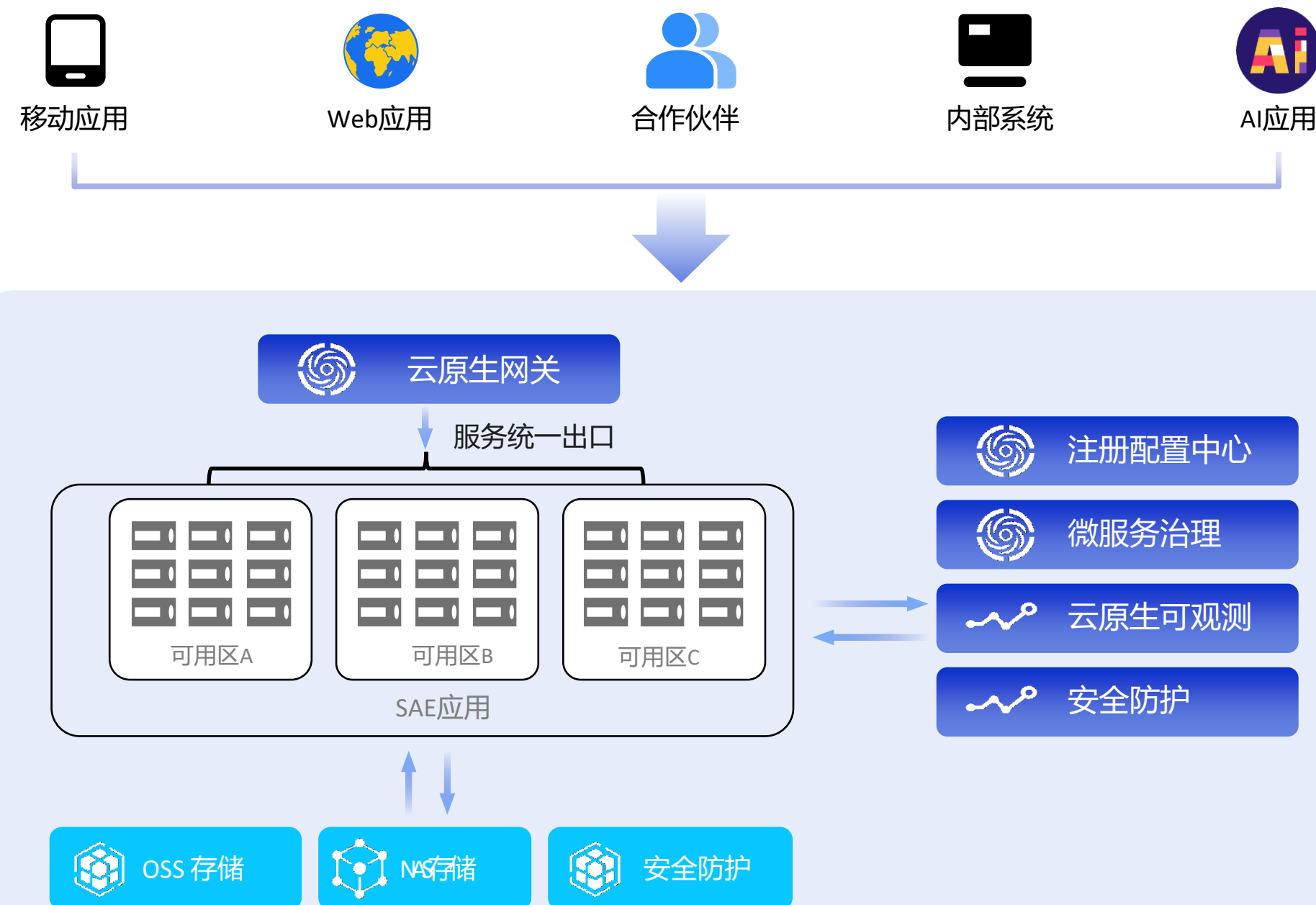
多可用区容灾

多可用区容灾对于SAE是默认的能力：

[一键开启](#)

SAE默认应用实例分散部署在多个可用区，实现跨机房容灾。单个可用区故障时，流量自动切换至其他可用区，保障业务连续性。

SAE内部架构示意图



多可用区优势

跨可用区容灾 (Multi-AZ)

- 应用实例自动分发至多可用区 (AZ)，单区故障秒级流量切换，可用性达99.95%+。
- 秒级自动切换，RTO≈0，RPO≈0

智能流量调度

- 同可用区优先路由：优先访问同AZ实例，跨区延迟降低80% (1ms→0.2ms)。
- 全局负载均衡 (SLB)：故障时自动切换至健康AZ，业务零中断。

全托管运维

- 无需维护资源，自动维护多AZ资源池，无需手动配置，运维成本降低70%。
- 按需跨AZ弹性伸缩，资源利用率提升50%

规模化落地 AI 应用的痛点

开源平台性能差

- 各组件（如：Worker、Plugin、数据库等）参数非最优配置
- 管控面与数据链路耦合，高并发无法保证稳定性
- 数据源存储格式单一，推理服务需要大量的计算资源，资源分配不均会导致性能瓶颈

运维复杂度高

- 本地部署复杂且维护成本高，需要频繁升级版本
- 需要自己管理应用的版本发布
- 周边配套不完善：没有配套的治理、可观测体系，事前事后无法及时发现并定位问题

成本不可控

- 资源错配，要么业务低峰期闲置烧钱，要么业务高峰期瞬间被打满，影响业务
- 人力维护投入大

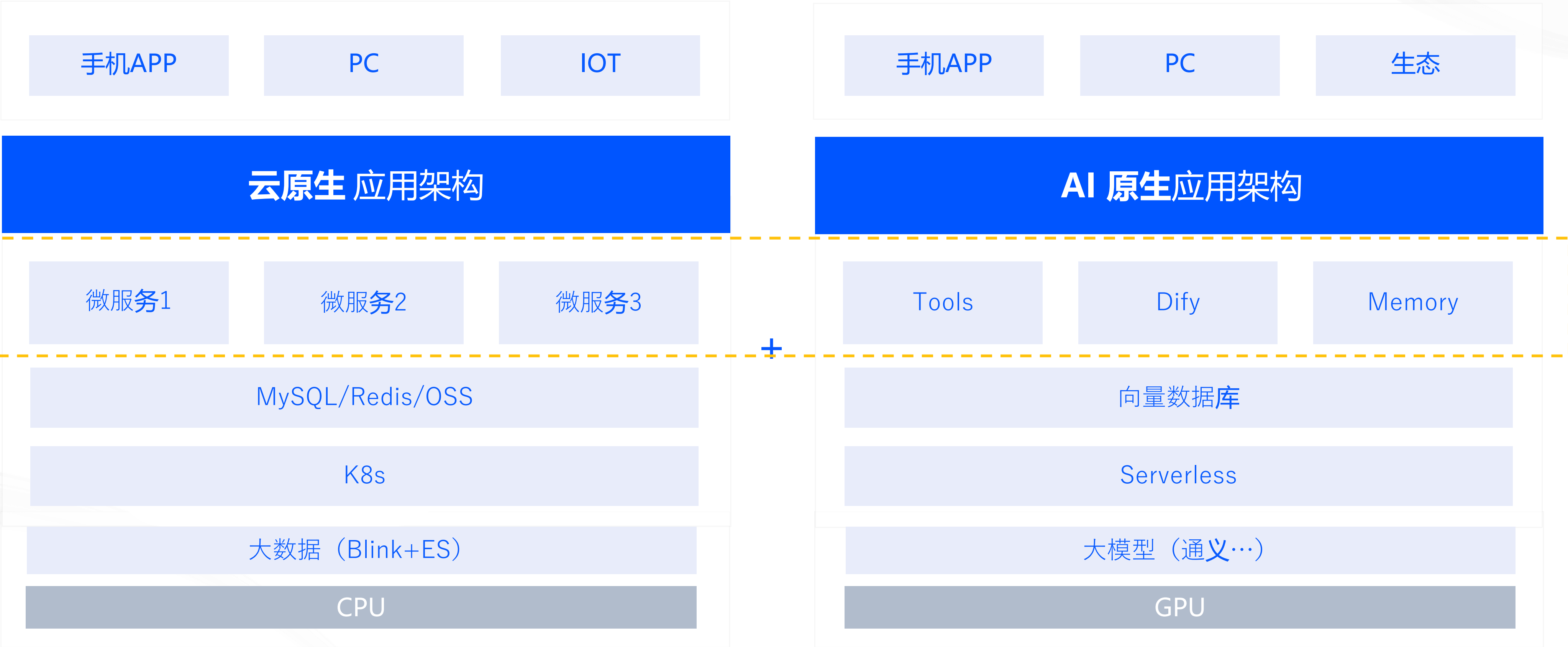
安全合规风险

- 流量防护弱，很容易被穿透
- 数据隐私与合规性管理困难

企业真正需要的是：开箱即用的开发体验 + 生产级的性能、稳定性及安全保障

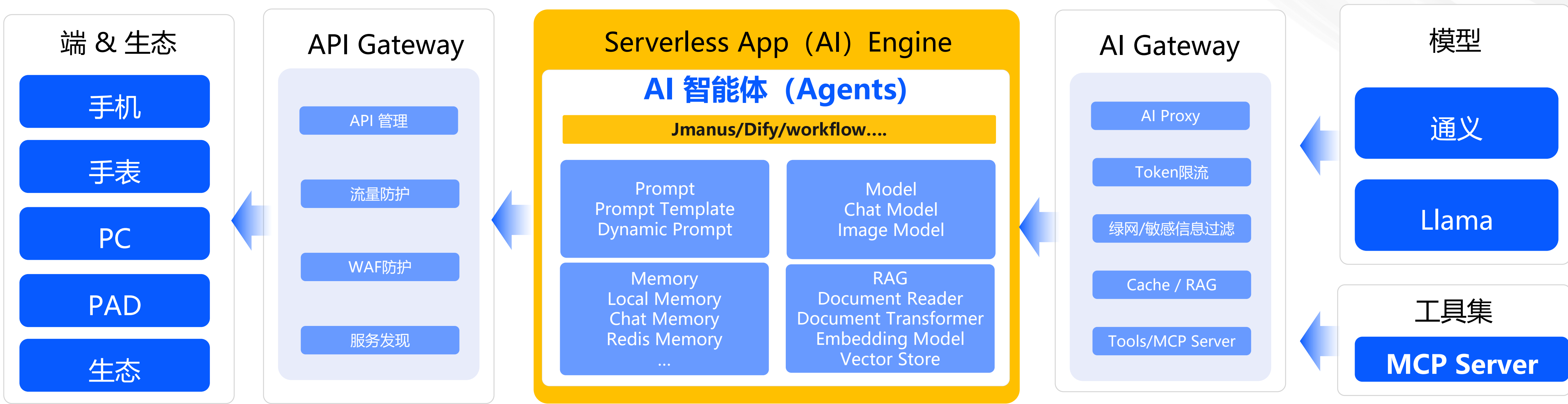
SAE 在 AI 原生应用领域的定位

不做开发平台的替代者，而是做它们的“护航舰”



SAE 致力于托管主流开源 AI 智能体应用开发平台（深度适配+全局赋能）

SAE 全托管 AI 智能体解决方案



Serverless (AI) 应用引擎托管 AI Agents 方案优势

简单易用

- 一分钟创建 AI 应用，无需任何额外配置
- 默认集成全链路监控，保证系统稳定性
- 无需关系底层资源，按需弹缩资源

稳定高可用

- 配置化，支持三 AZ 部署，默认支持智能化可用区，实例粒度的自动化迁移
- 默认支持负载均衡与健康检查联动保证无损上下线

低成本

- 按需按量付费，潮汐流量弹性使用，无需冗余保证资源
- 支持多种规格资源，并提供闲时计量资源类型，提供更低成本的算力

安全保障

<

运维配套 – 自定义弹性伸缩

作为 SAE 的核心竞争力，相对传统 ECS 的弹性，SAE 更精准更降本；相对 K8s 弹性，SAE 的指标和策略更丰富，上手门槛更低。

1. 指标弹性（CPU、Mem、QPS、RT等）

适用于有突发流量、典型脉冲的应用场景，多用于互娱 / 游戏 / 社交平台 / 电商等行业。



优势：比开源K8s HPA 指标丰富，且可以自定义指标。

2. 定时弹性

适用于资源画像存在周期性的应用场景，多用于餐饮 / 出行 / 证券 / 医疗政府等行业。



优势：操作简单，易用。

3. 混合弹性（定时弹性 & 指标弹性混用）

适用于固定时段内有突发流量、典型脉冲，常稳时段内流量波动不均的应用场景，多用于 媒体报社 / 在线教育 / 语音识别合成等行业。



优势：比开源K8s HPA 指标丰富，且可以自定义指标。

一键起停开发测试环境

中大型企业多套环境，内部环境长期资源浪费，使用 SAE 一键启停，可以节省一部分资源成本。

