

QCon
全球软件开发大会

蚂蚁大促场景下的全链路压测 体系构建与保障实践

刘凯宁

蚂蚁集团-SRE技术专家



个人简介



刘凯宁

- 蚂蚁集团 SRE 技术专家
- 多次参与蚂蚁集团超大型活动的稳定性保障，承担过大促保障队长、全链路压测负责人、全链路资源容量负责人、全链路资金安全保障负责人等角色
- QCon 2024 上海站明星讲师

目录

01 蚂蚁大型活动保障架构概览

02 蚂蚁全链路压测体系

03 蚂蚁大促全链路压测实战

04 未来已来！AI 压测初探

📍 北京

QCon

全球软件开发大会

会议时间：4月10-12日

- 大模型赋能 AIOps
- 越挫越勇的大前端
- 云上业务架构演进
- 多模态大模型及应用
- 海外 AI 应用创新实践

📍 北京

AiCon

全球人工智能开发与应用大会

会议时间：6月27-28日

- 端侧智能
- AI Agent
- 多模态大模型
- 金融+大模型

📍 上海

QCon

全球软件开发大会

会议时间：10月23-25日

- AI Agent
- 大模型训练推理
- 端侧 AI
- 搜推深度融合
- 数智金融

4月

5月

6月

8月

10月

12月

📍 上海

AiCon

全球人工智能开发与应用大会

会议时间：5月23-24日

- 通用大模型
- AI Agent
- 垂直领域应用
- 模型可解释性

📍 深圳

AiCon

全球人工智能开发与应用大会

会议时间：8月22-23日

- 模型效率与部署
- 多模态大模型
- 大模型安全
- 智能硬件

📍 北京

AiCon

全球人工智能开发与应用大会

会议时间：12月19-20日

- 通用大模型
- 智能硬件
- LMOps
- 具身智能

极客邦科技 2025 年会议规划

促进软件开发及相关领域知识与创新的传播



参会咨询



查看会议

01

蚂蚁大型活动保障架构概览

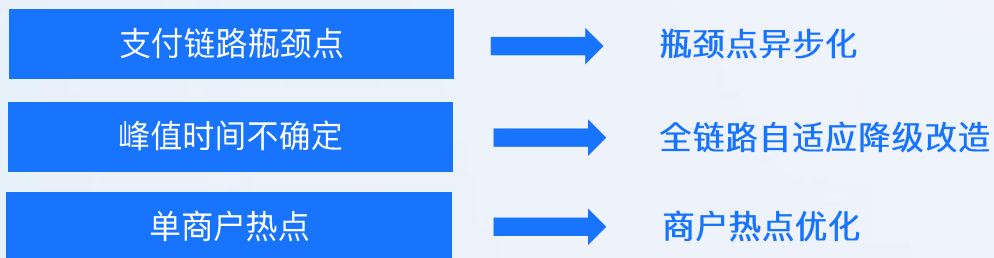
● 引子：大促活动形式及特点

① 支付峰值型大促

挑战点1：支付宝SKA商户峰值时间不确定

挑战点2：超级大促支付秒杀峰值高，链路压力巨大

挑战点3：支付商户聚集，极易引发支付链路热点

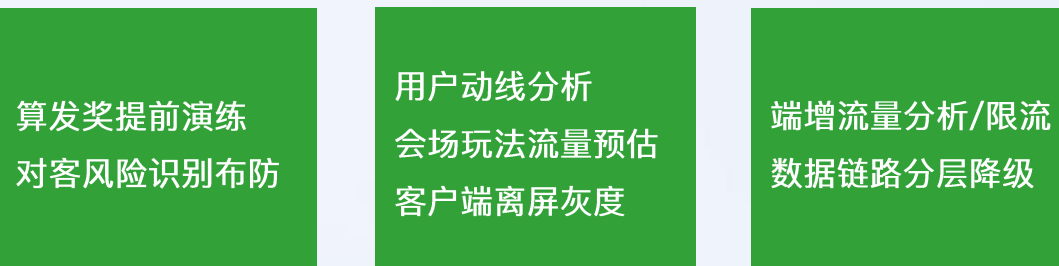


② 玩法峰值型大促

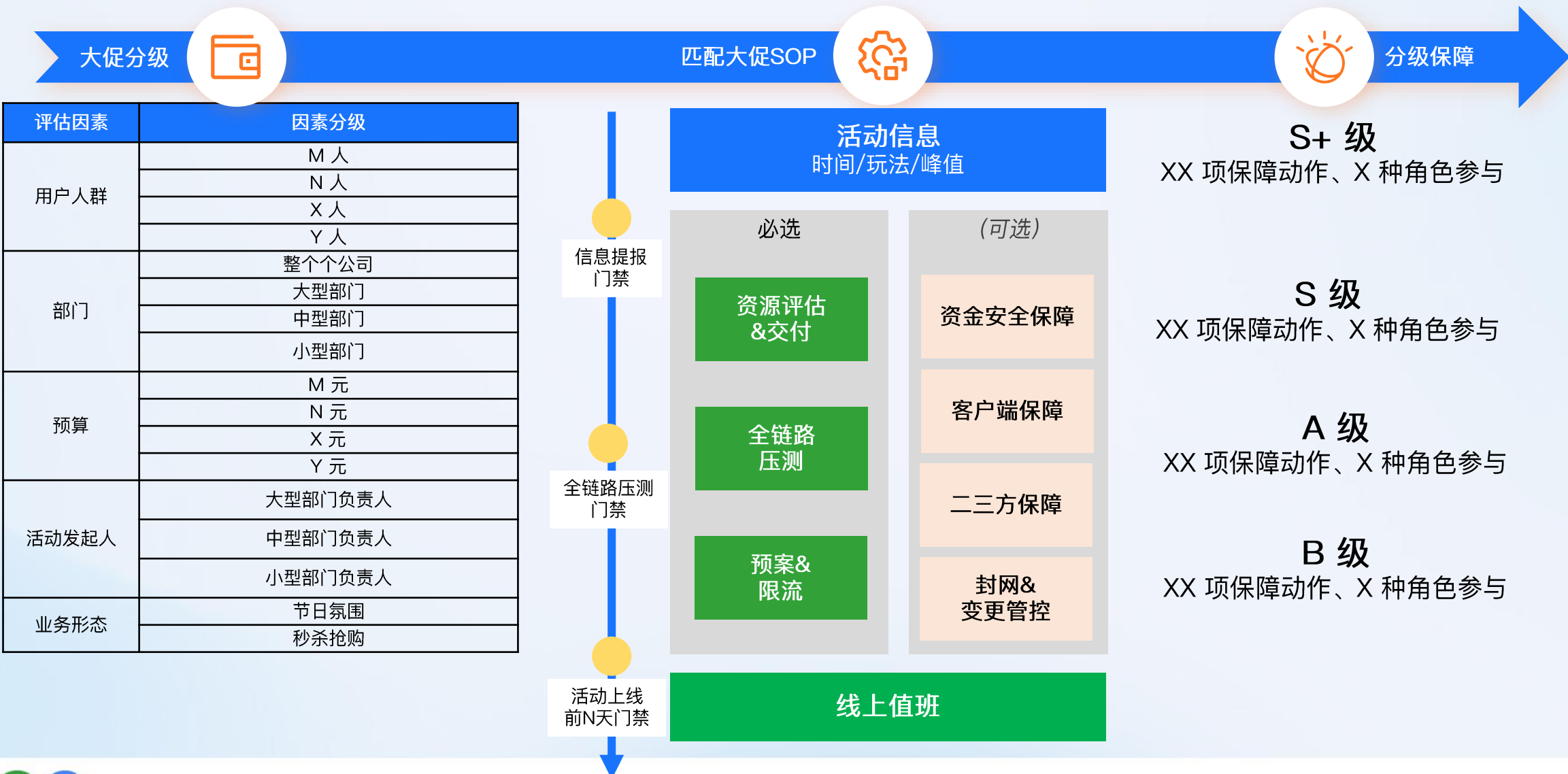
挑战点1：伴随营销秒杀抢券/红包发放，资金/客诉风险高

挑战点2：玩法多样复杂，C端用户行为难以预测准确

挑战点3：通常带来端增，整个APP以及在离线链路压力巨大

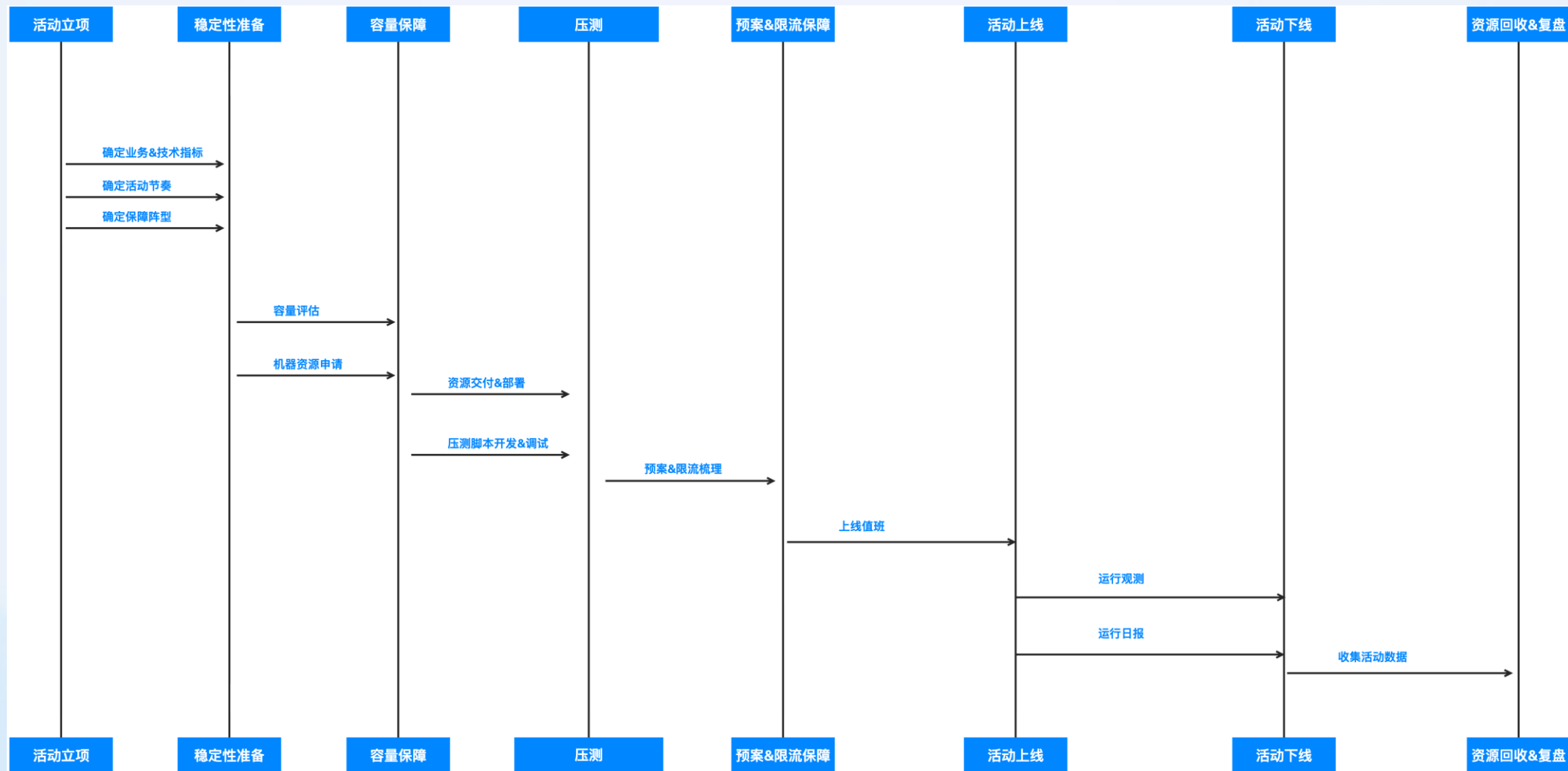


蚂蚁大促分级及活动保障 SOP





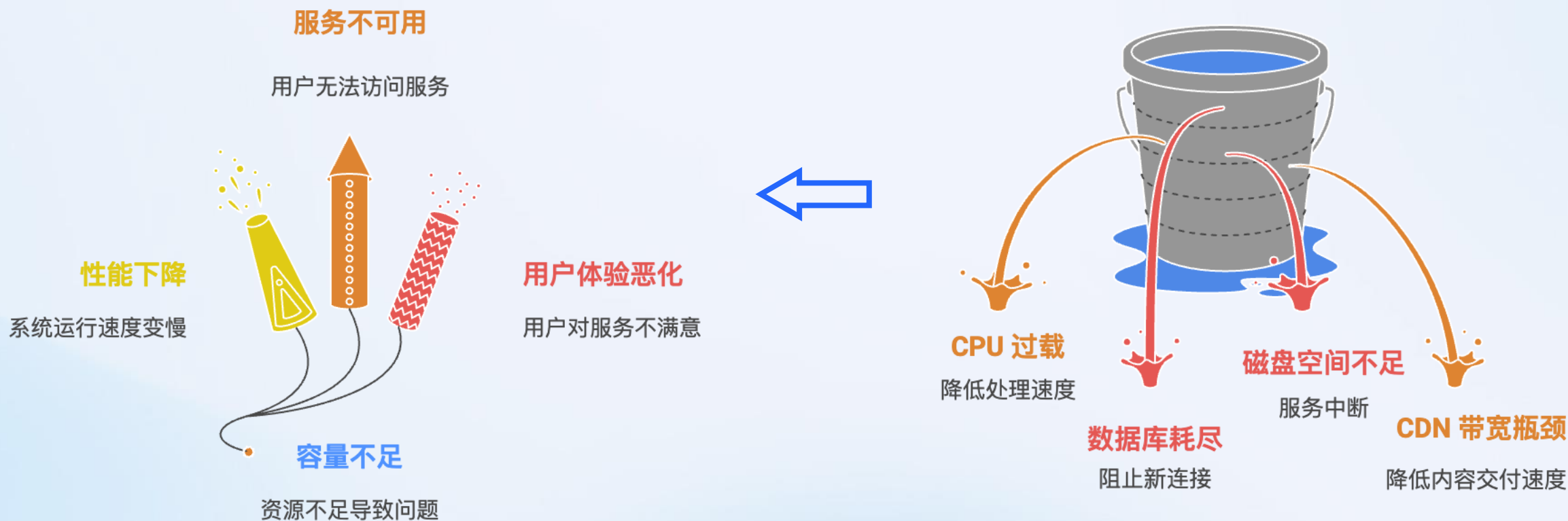
蚂蚁大促保障整体流程





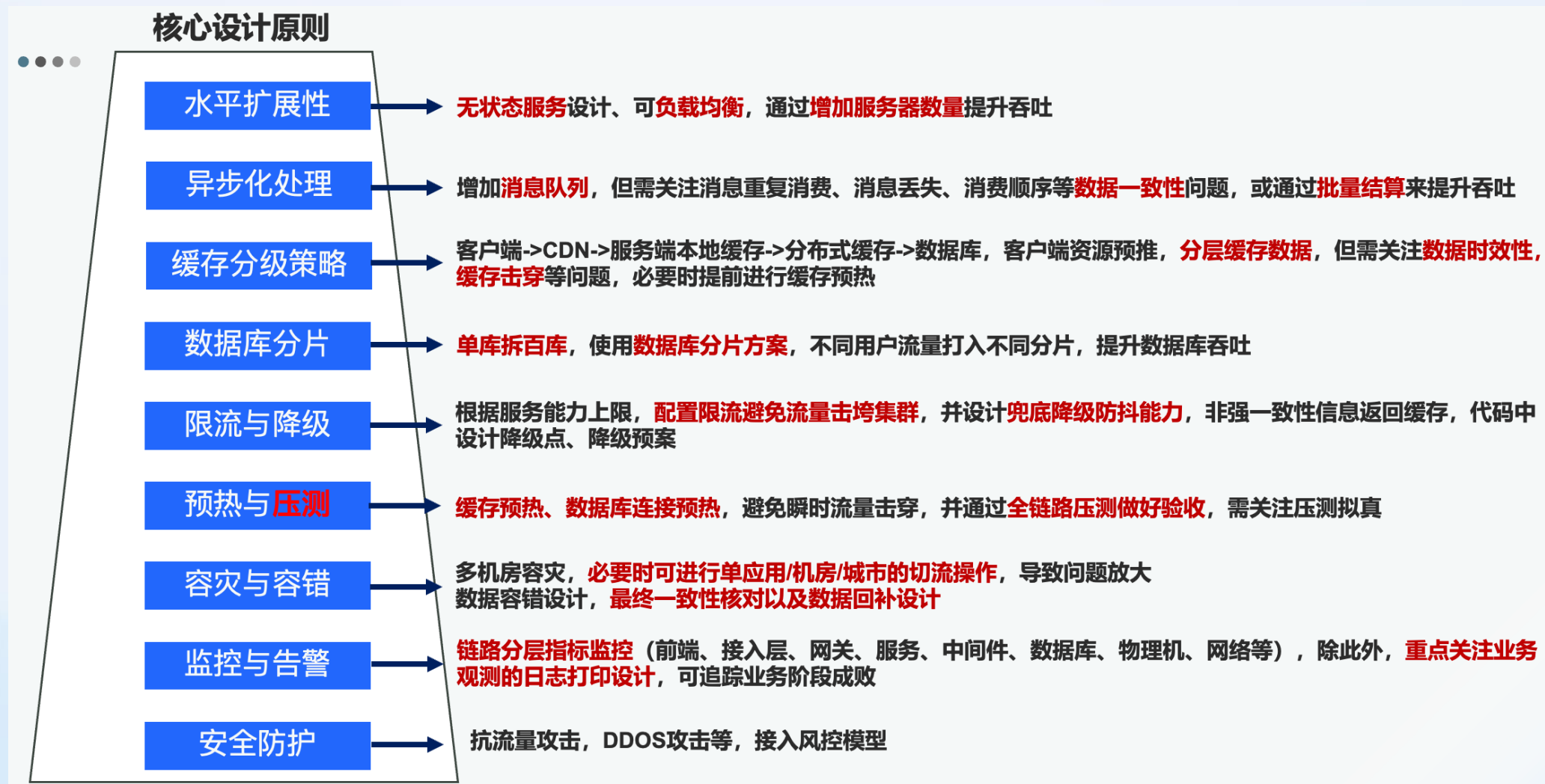
什么是容量风险？

在SRE（Site Reliability Engineering，站点可靠性工程）领域中，“容量风险”（Capacity Risk）指的是：系统当前或未来可能因资源容量不足，无法满足用户需求或业务增长，从而导致性能下降、服务不可用或用户体验恶化的潜在风险。





容量风险的本质：如何承接高并发流量？



02

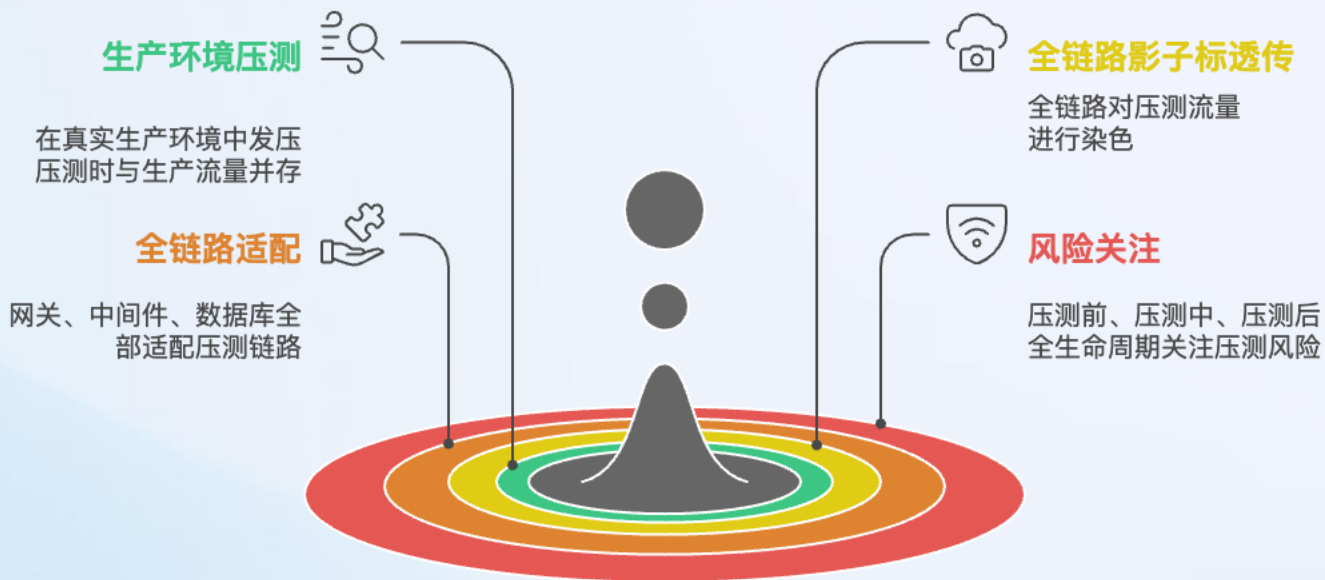
蚂蚁全链路压测体系

● 全链路压测定义

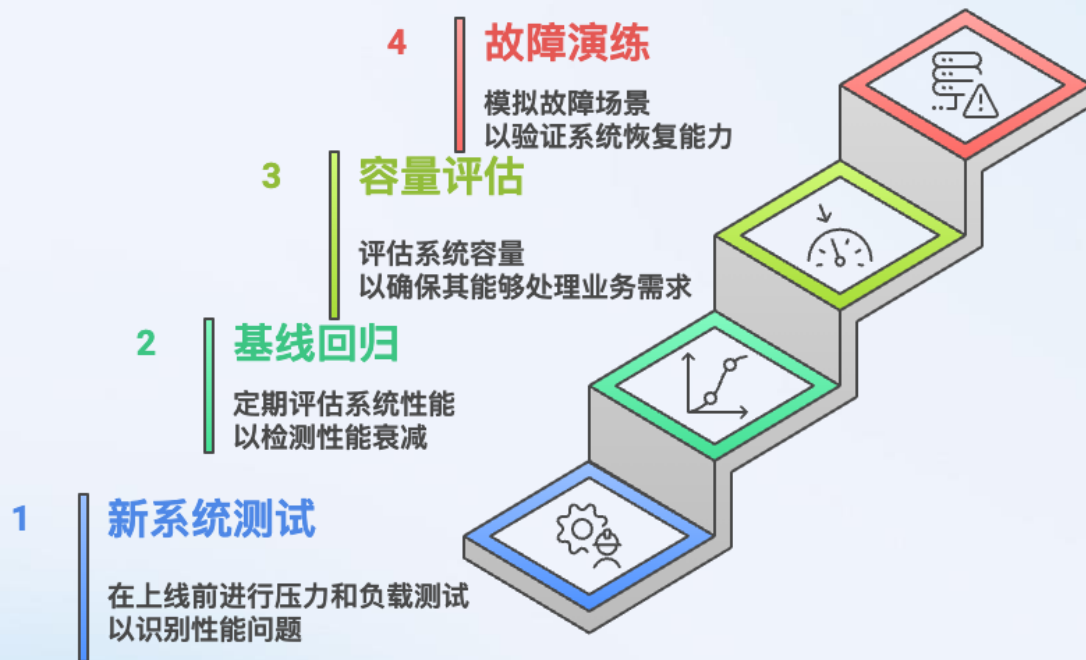
全链路压测是模拟真实业务场景流量，对完整系统链路（前端→后端→数据库等）进行高并发压力测试，发现性能瓶颈与单点故障，确保系统在峰值（如大促）下稳定运行。通过流量染色、影子表等技术实现零业务影响

—— From [Ling-IT](#)

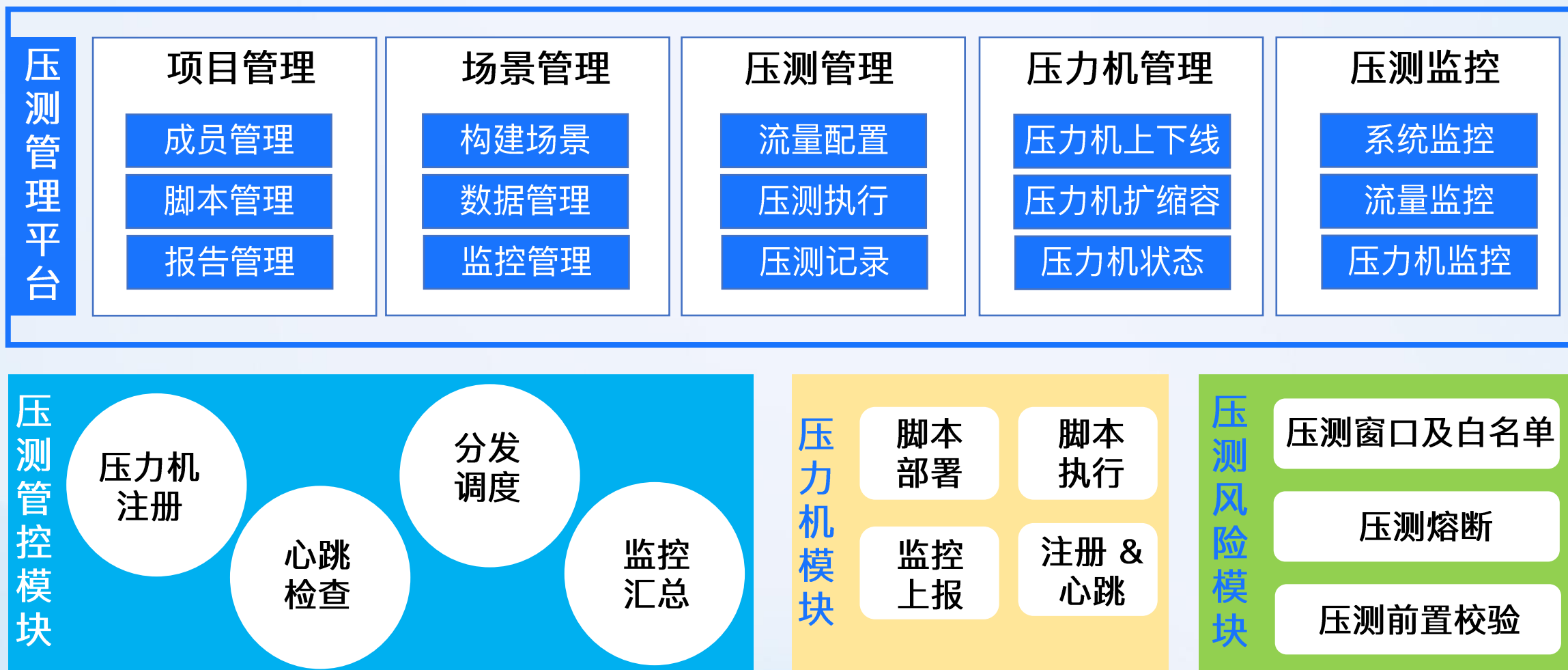
蚂蚁全链路压测特点



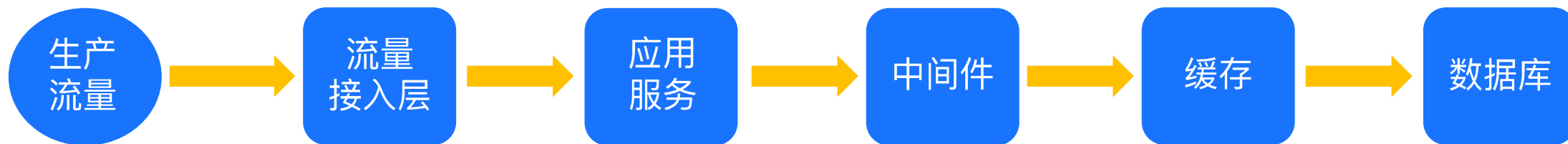
典型使用场景



蚂蚁全链路压测平台架构介绍



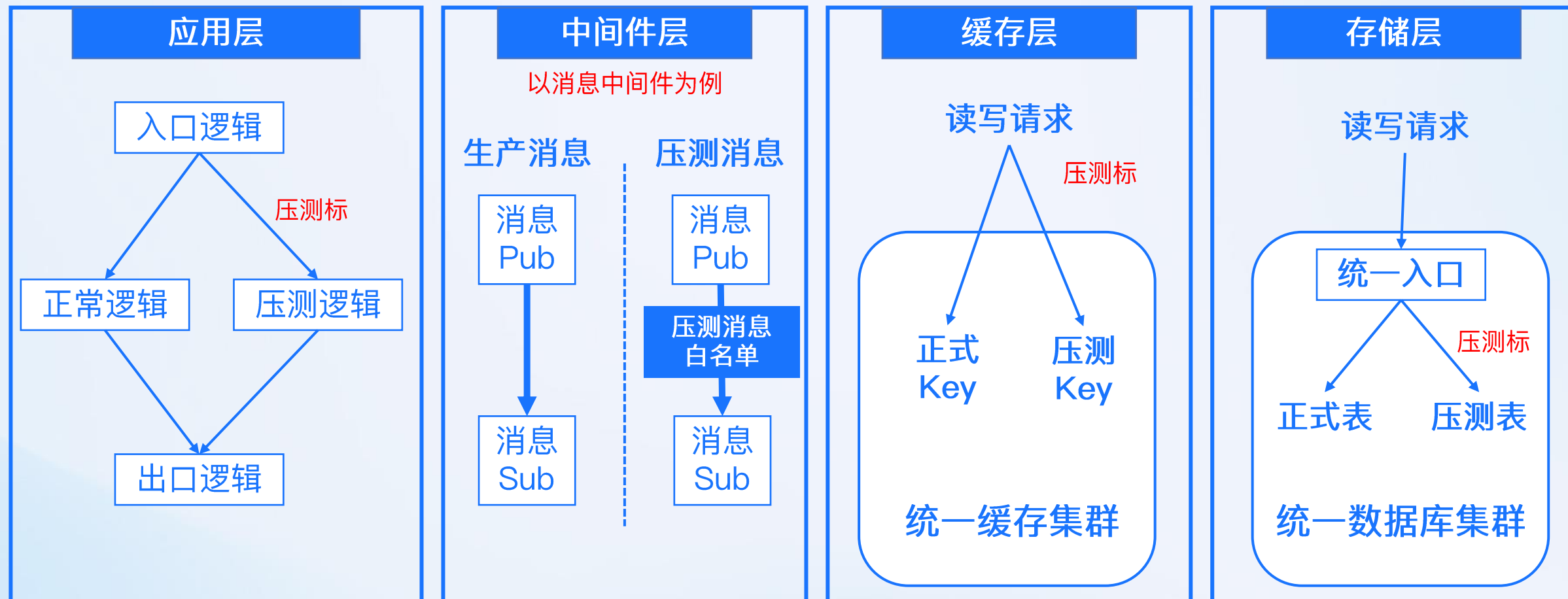
● 蚂蚁全链路压测平台核心技术介绍_全链路染色打标



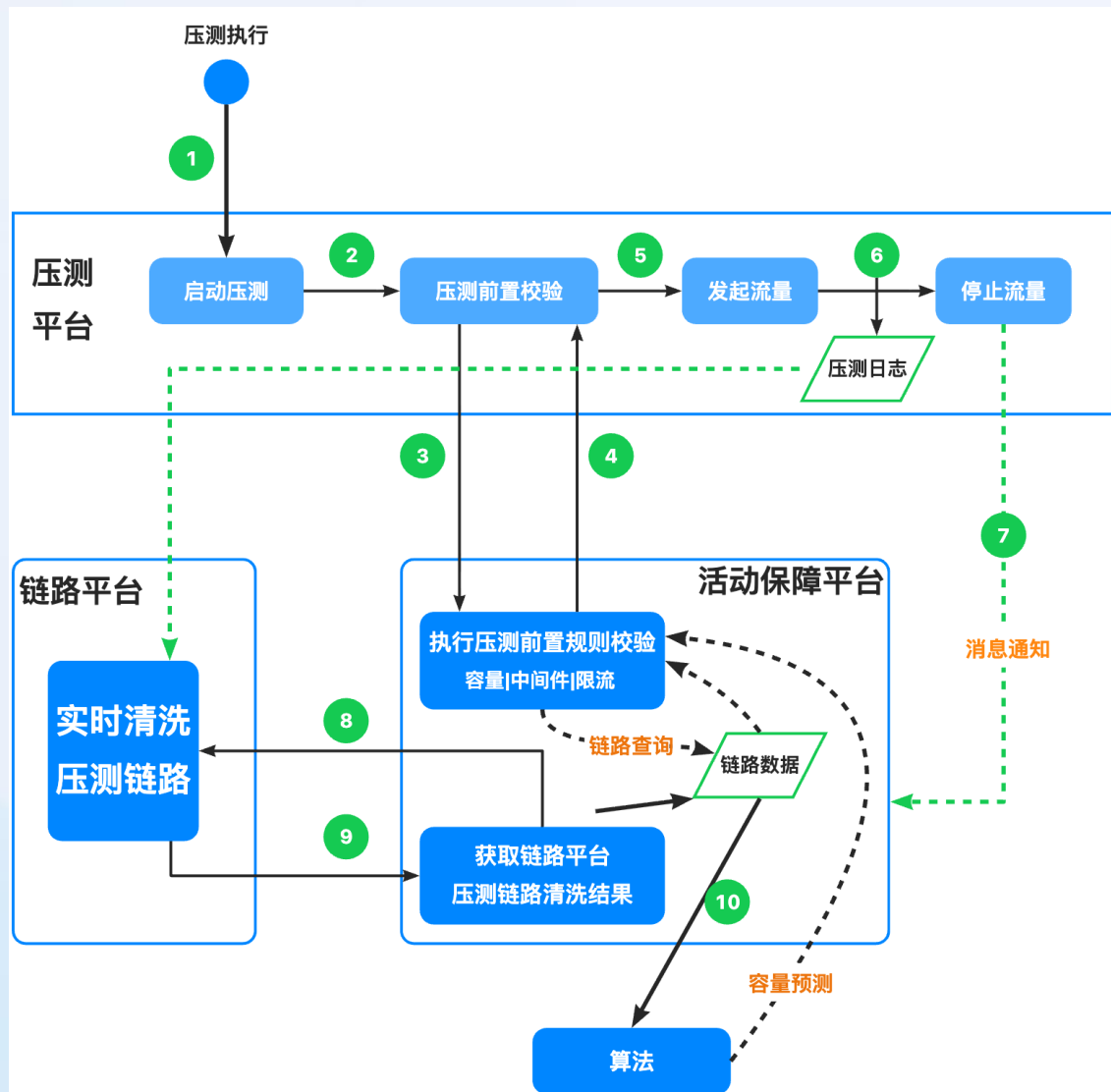
基于 Trace 实现全链路上下文传递压测标



蚂蚁全链路压测平台核心技术介绍_压测隔离

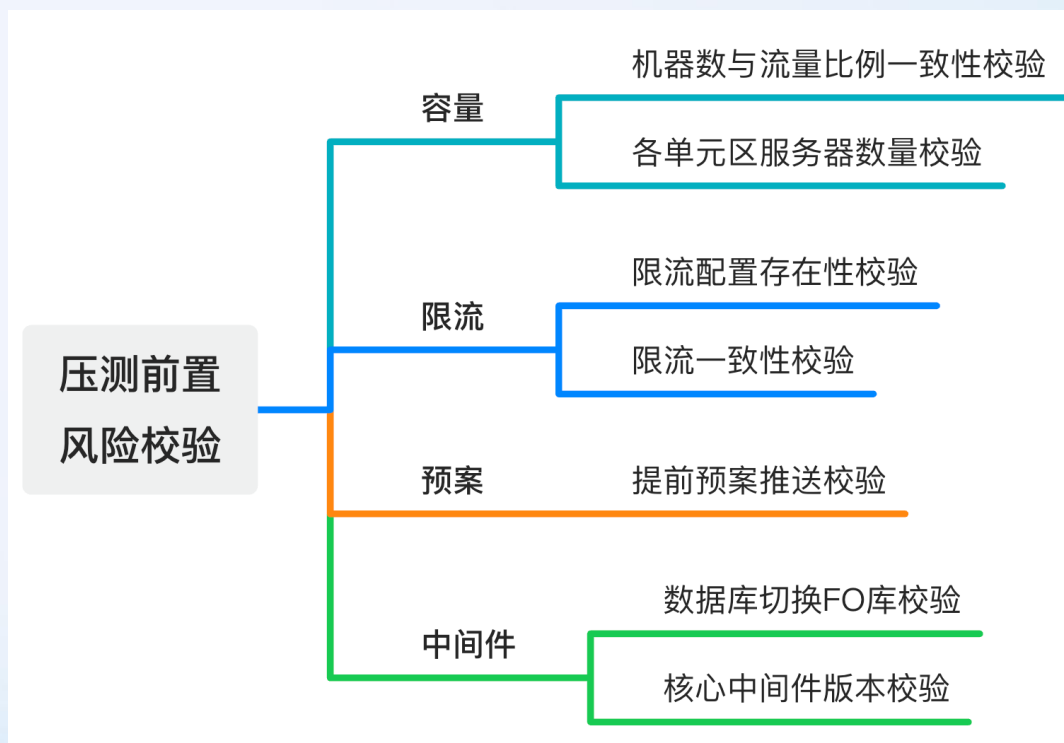


蚂蚁全链路压测平台核心技术介绍_压测前置风险校验



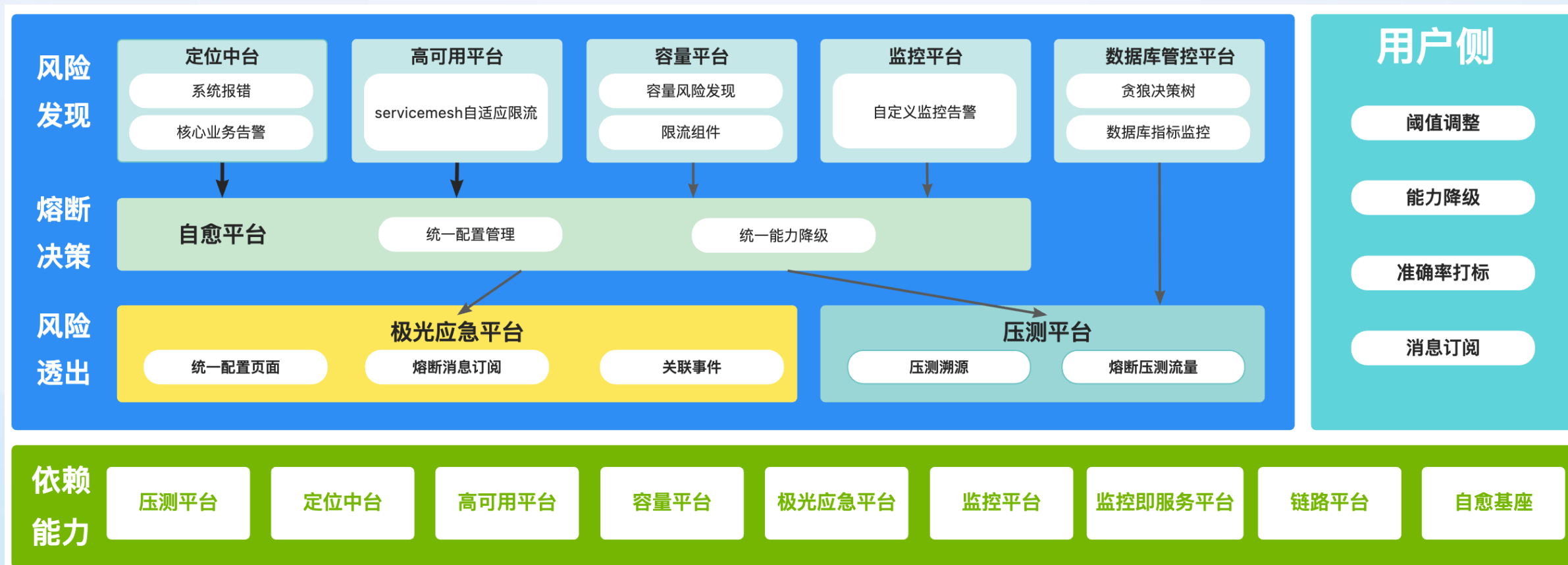
压测前置校验

压测起量之前，主动进行容量、中间件、限流、预案等风险的校验，提前感知并解除压测风险，避免压测过程中出现低级问题

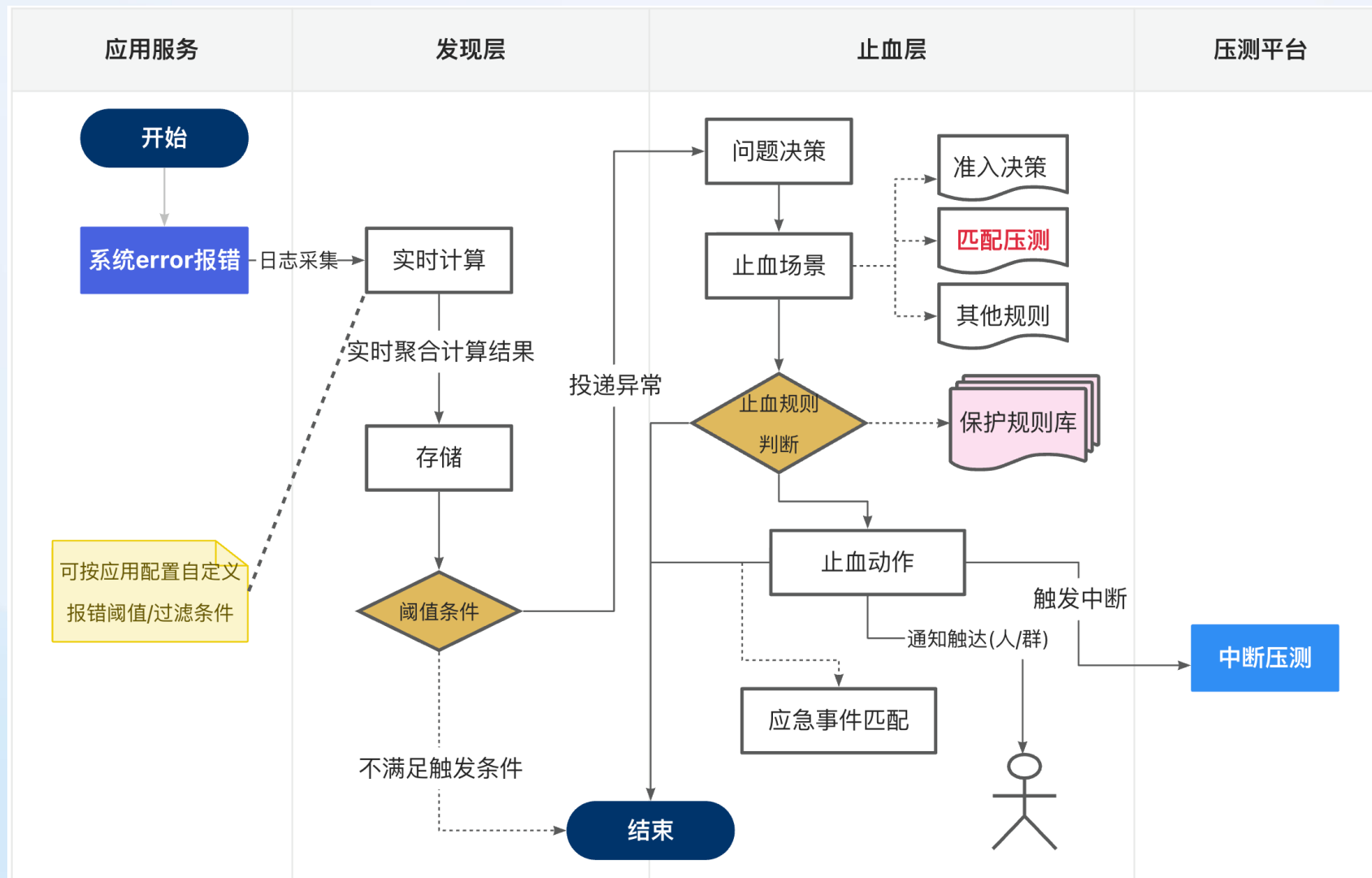


蚂蚁全链路压测平台核心技术介绍_压测熔断

压测熔断：及时发现压测链路中的各种异常，快速、有效地阻隔压测风险



蚂蚁全链路压测平台核心技术介绍_压测熔断_举例



系统error熔断：监测应用的系统报错日志量级，超过特定阈值则会触发熔断

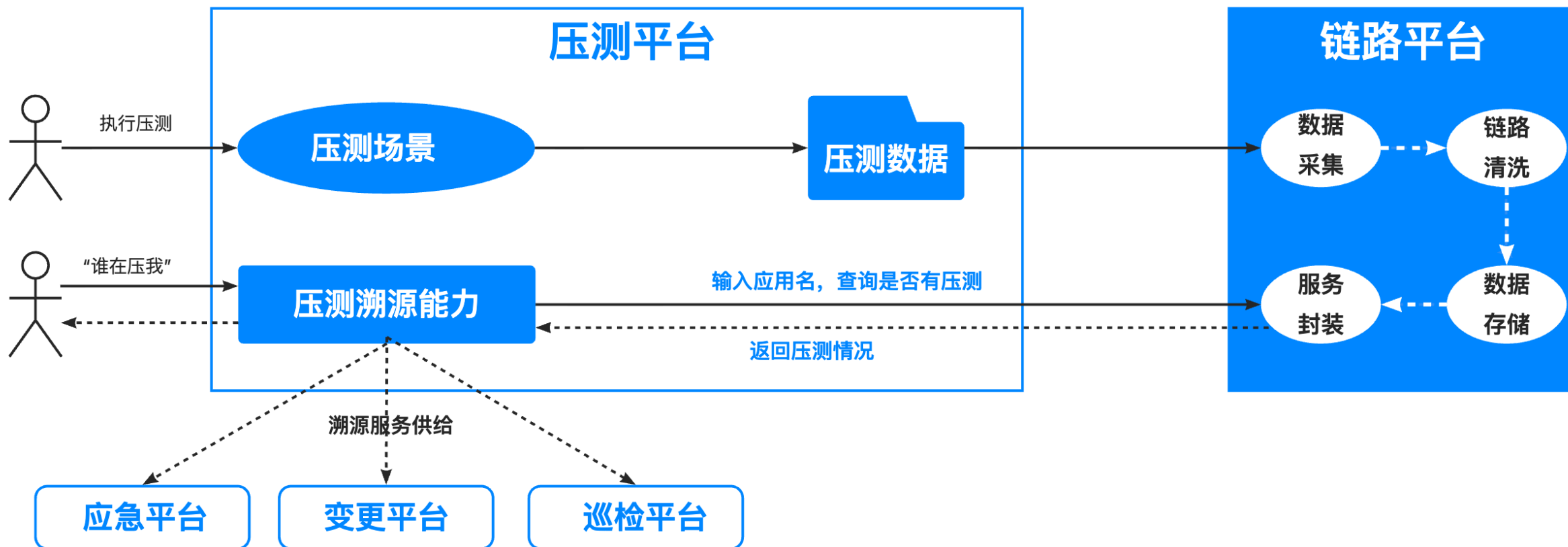
熔断场景：被压测应用的系统报错量级上涨

熔断时效：秒级，20秒之内

熔断要求：应用需要接入标准日志框架，并配置熔断报错阈值

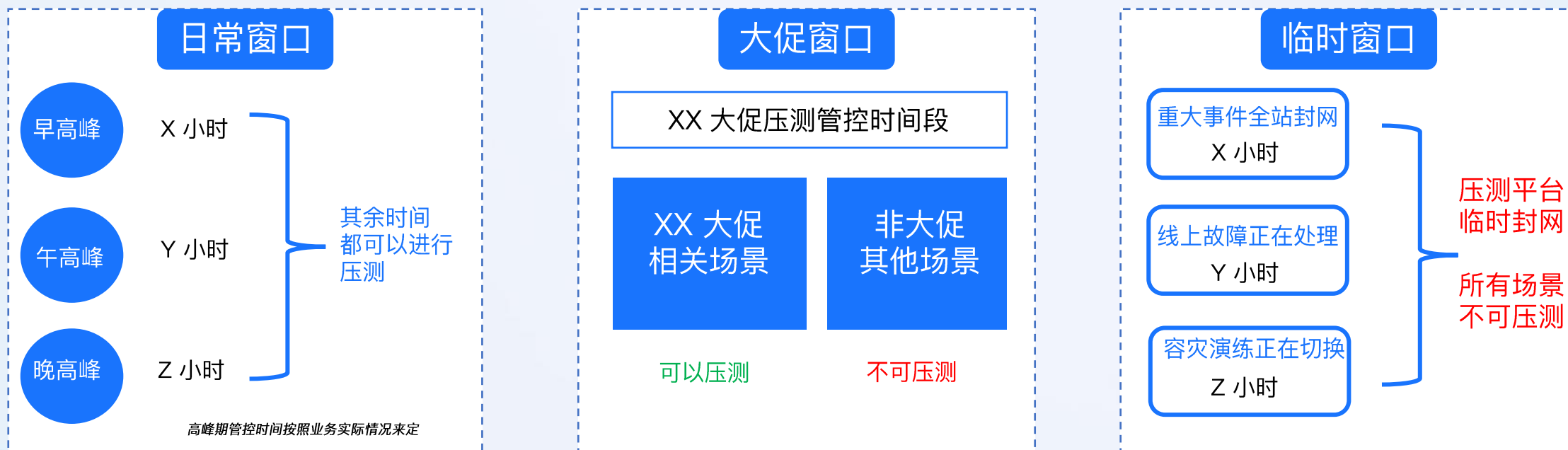
蚂蚁全链路压测平台核心技术介绍_压测溯源定位

压测溯源：通过应用 appName 查询该应用当前的压测情况，快速获取某一应用被哪个场景压测，便于进行压测来源判断，异常情况下能够及时应急处理



蚂蚁全链路压测平台核心技术介绍_压测窗口及白名单

为统一管理压测流量，同时尽可能降低压测对线上产生影响，特设置多维度的压测管控窗口



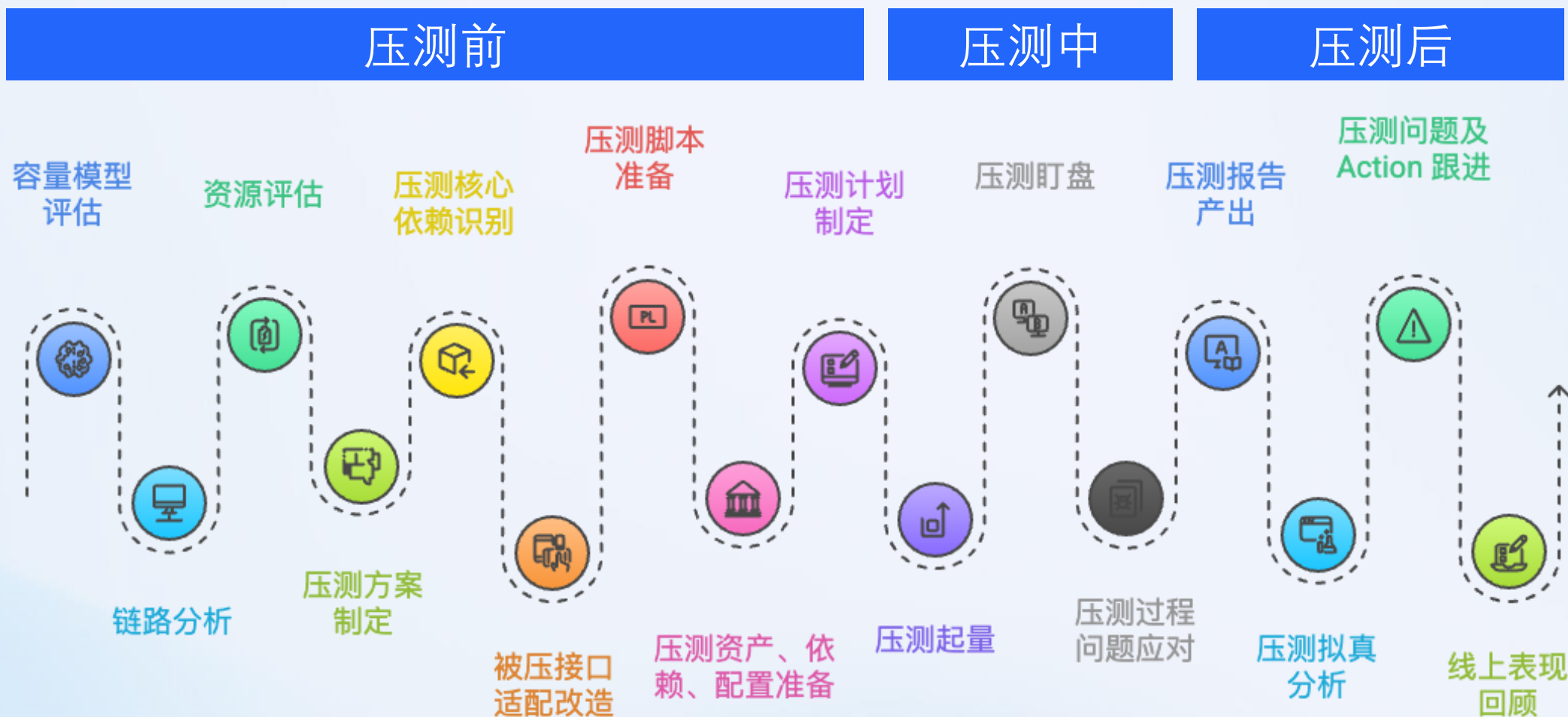
压测白名单机制

- 原则上压测管控窗口期间需要严格按照封网要求进行压测变更管控，但在落地执行的时候实际情况可能较为复杂，本着实事求是的考虑，压测平台有白名单机制，允许某压测场景在一定时间之内可以不受压测窗口管控
- 当前申请压测白名单需要走审批，审批人为申请人主管+压测平台管理员

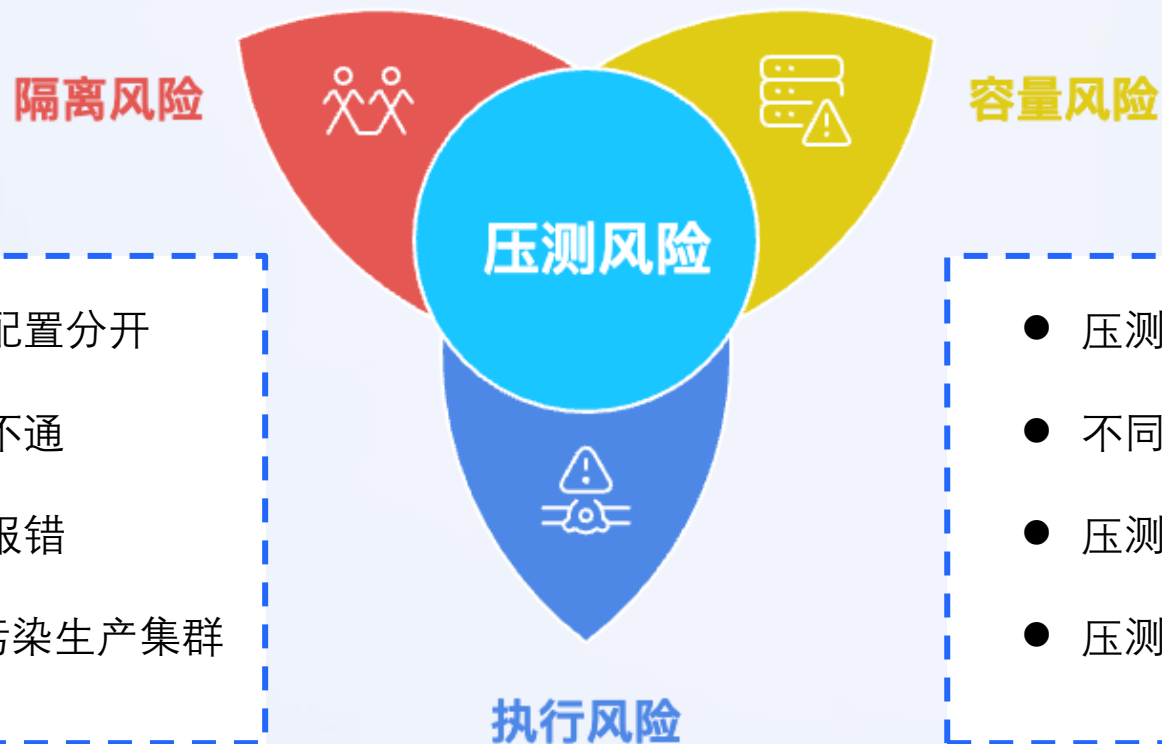
03

蚂蚁大促全链路压测实战

全链路压测基本流程



● 压测风险



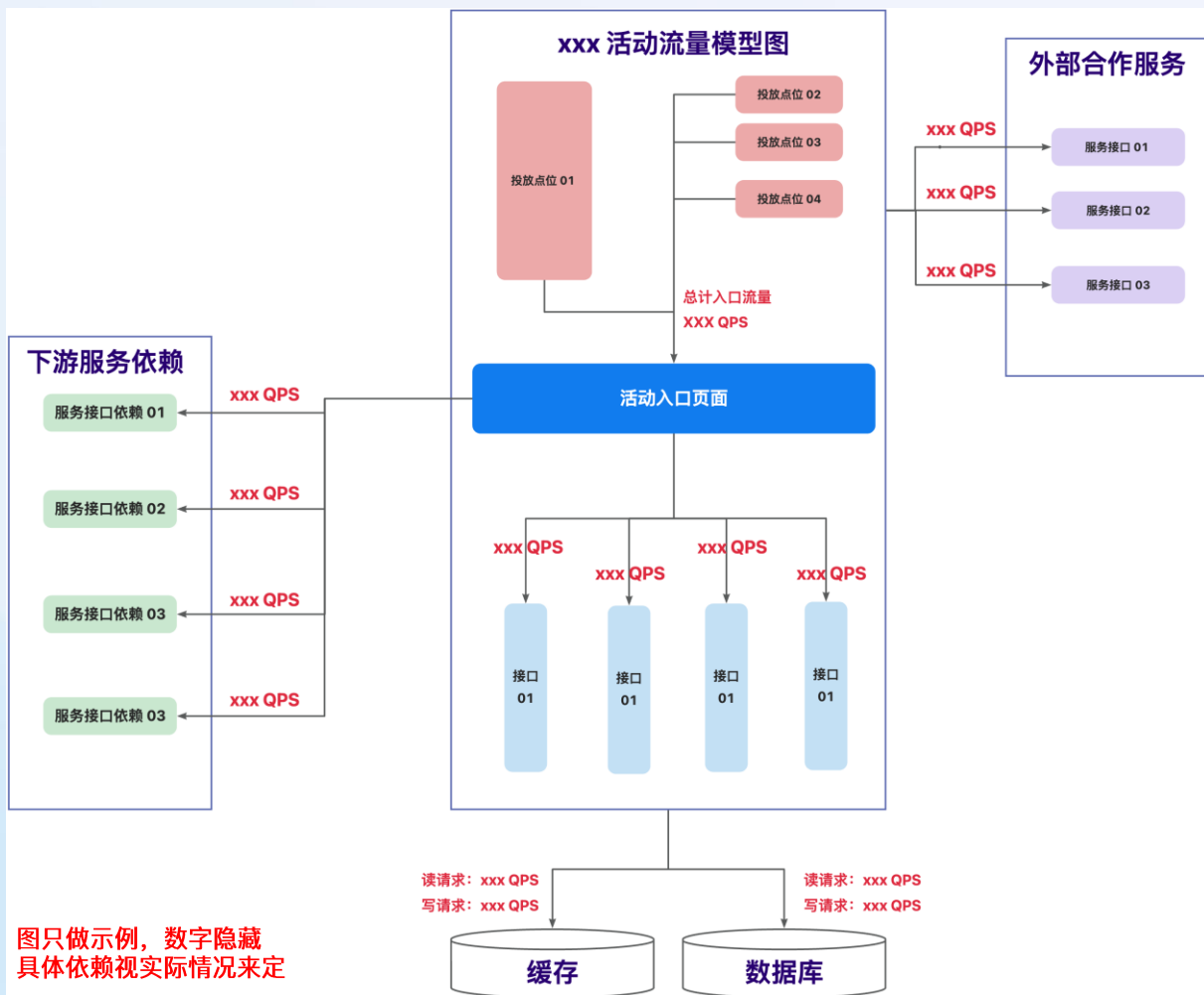
- 压测配置未完全与生产配置分开
- 压测标记未透传，链路不通
- 压测表未创建，数据库报错
- 缓存未使用压测 key，污染生产集群

- 压测之前应用未扩容到位
- 不同 Zone 的机器不均匀
- 压测限流未配置
- 压测熔断配置没开

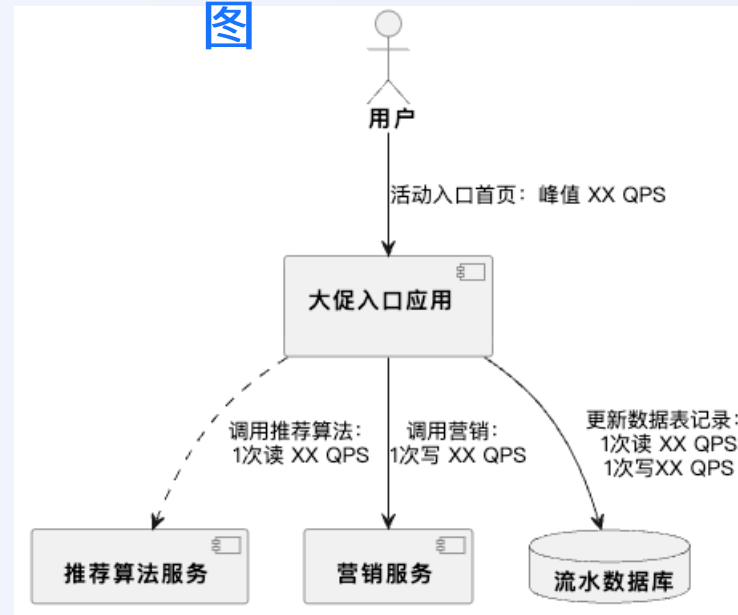
- 未配置压测监控，压测期间未正确盯盘
- 起量过快，系统某些节点无法承载瞬间高压流量
- 线上有异常未及时感知并停压

容量评估与资源准备

整体流量模型图



单接口读写模型图



资源需求汇总表

应用名	总计 QPS	资源需求	具体机房分配
主应用	xxx QPS	xxx cores	城市A: xxx1 cores 城市B: xxx1 cores
下游应用 01	xxx QPS	xxx cores	城市A: xxx1 cores 城市B: xxx1 cores
下游应用 02	xxx * 2 QPS	xxx cores	城市A: xxx1 cores

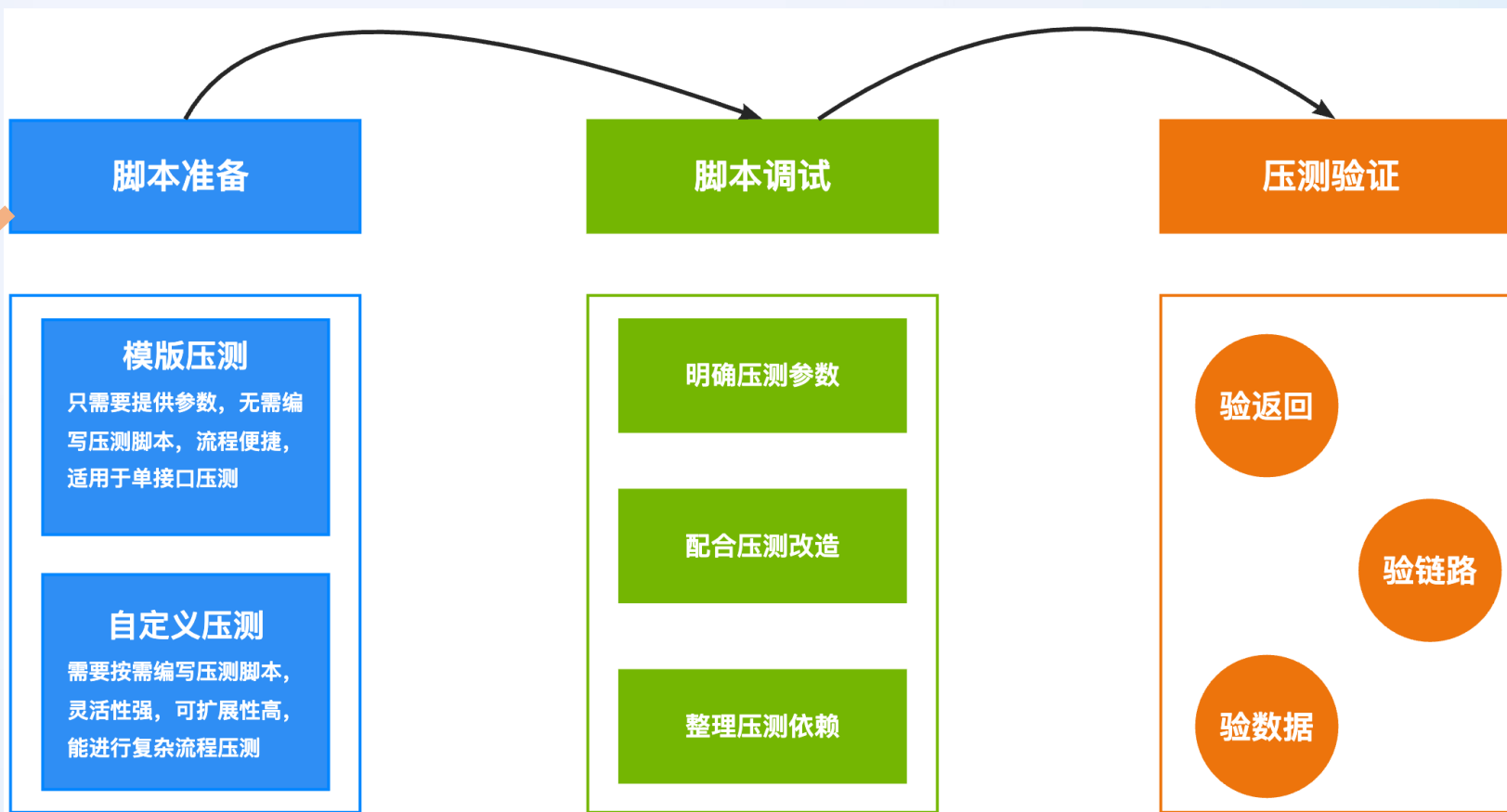
● 压测脚本编写、调试、验证

压测基本概念



压测前期准备

- 按照流量模型确认被压接口
- 为每个被压接口准备起量压测脚本
- 调试压测脚本并进行压测验证

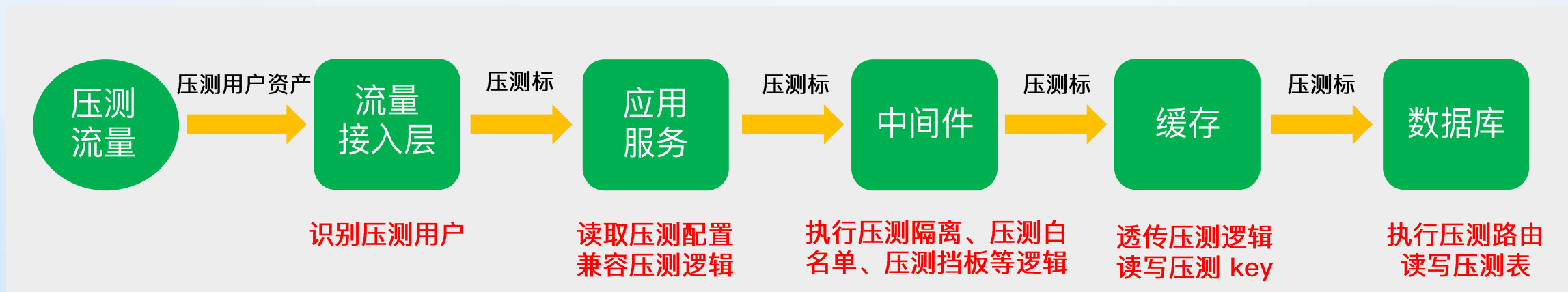


● 压测配置与资产准备

压测配置：一般分为全局配置和业务自定义配置。全局配置一般用于通用内容配置，例如访问地址、环境信息等内容；业务自定义配置一般实现在业务接口中，在压测期间会推送对应的压测开关、压测特定配置等，来保证全链路压测能走通

压测数据资产：用于业务流或业务单元脚本中需要用到的压测数据，一般是压测用户资产；在执行压测时，压测中心会将压测数据分发到各台压力机中，用于压测任务的执行

全链路压测必须搭配压测配置及压测数据资产



压测计划制定

由于全链路压测涉及到的范围广、链路长、人员众多，压测负责人必须制定详细的压测计划，以便压测过程顺利、压测问题及时解决，助力全链路压测尽快达成既定目标

压测范围

- 压什么场景？
- 压哪些接口？
- 涉及到哪些应用？
- 可能会影响哪些业务？

压测模型

- 每个场景要压多少量级？
- 核心应用总计要承接多少流量？
- 哪些链路节点需要降级？
- 机器资源是否扩容完毕？

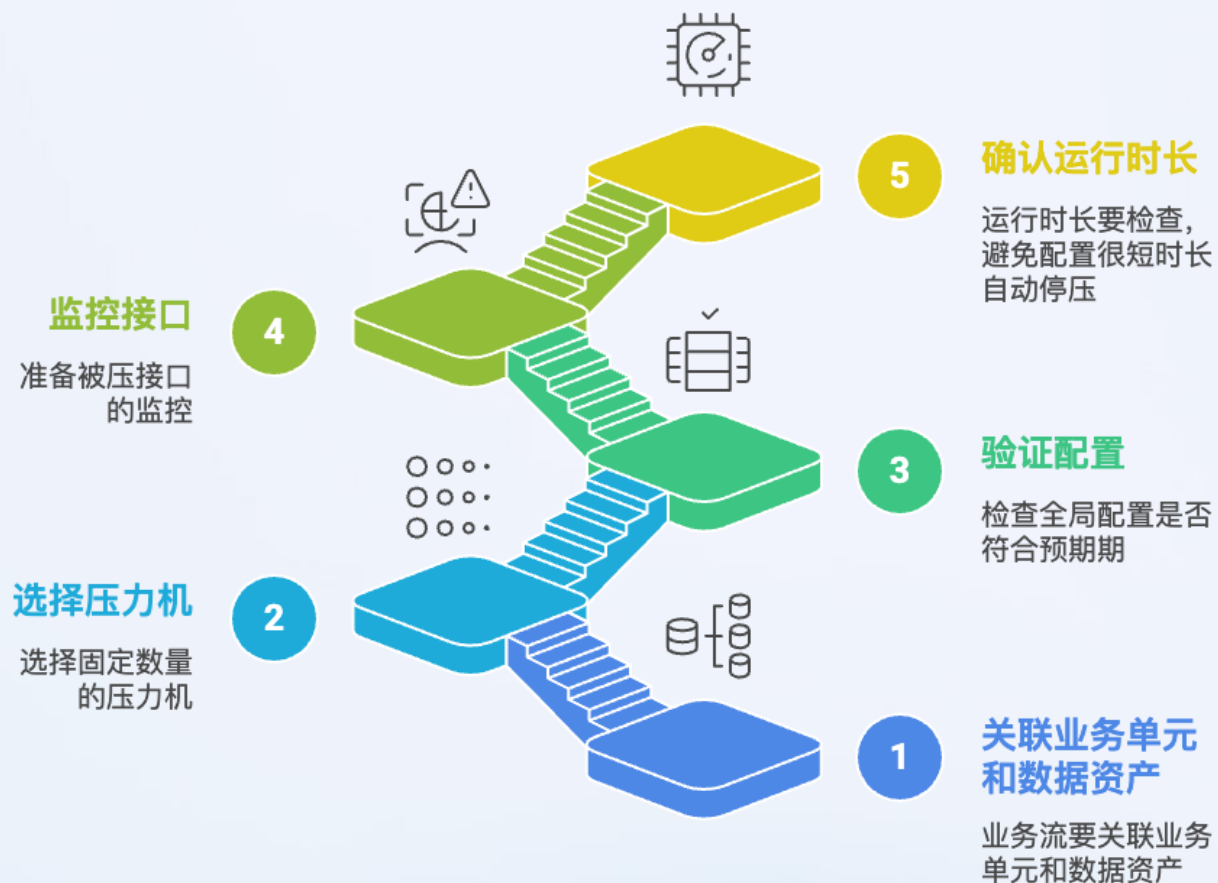
压测组织

- 全链路压测需要哪些人员必须到场配合？
- 整体压测形式是现场压测还是线上压测？
- 压测要组织几个场次？

压测执行

- 每场压测具体起量计划是怎样的？
- 压测需要重点关注哪些监控大盘？
- 压测问题如何记录？
- 压测复盘如何进行？

压测执行及盯盘



压测准出

整体压测结论：压测是否通过？

压测到量情况

- 各个接口目标量级、压测量级、到量百分比
- 各个下游服务依赖的实际到量情况

系统负载情况

- 业务指标：峰值成功量、峰值成功率、峰值耗时
- 系统指标：CPU水位、内存水位、JVM 情况、线程池情况、DB 负载、缓存负载.....

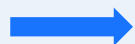
压测问题情况

- 已解决的问题：问题数量、问题原因、问题当前状态
- 未解决的问题：问题原因、问题跟进人、问题解决要求时间、问题当前状态

当前风险情况

- 当前压测还遗留哪些风险？是否需要上升解决？

压测产出



压测报告

系统峰值
水位截图

压测到量
监控截图

压测执行
记录链接

压测问题
记录列表

04

未来已来：AI 压测初探

AI 链路的特点

传统业务关注的核心指标

总量

成功量

平均
耗时

CPU
水位

峰值
QPS

网络
延迟

对比传统业务链路，AI 业务链路最关注的技术风险目标：

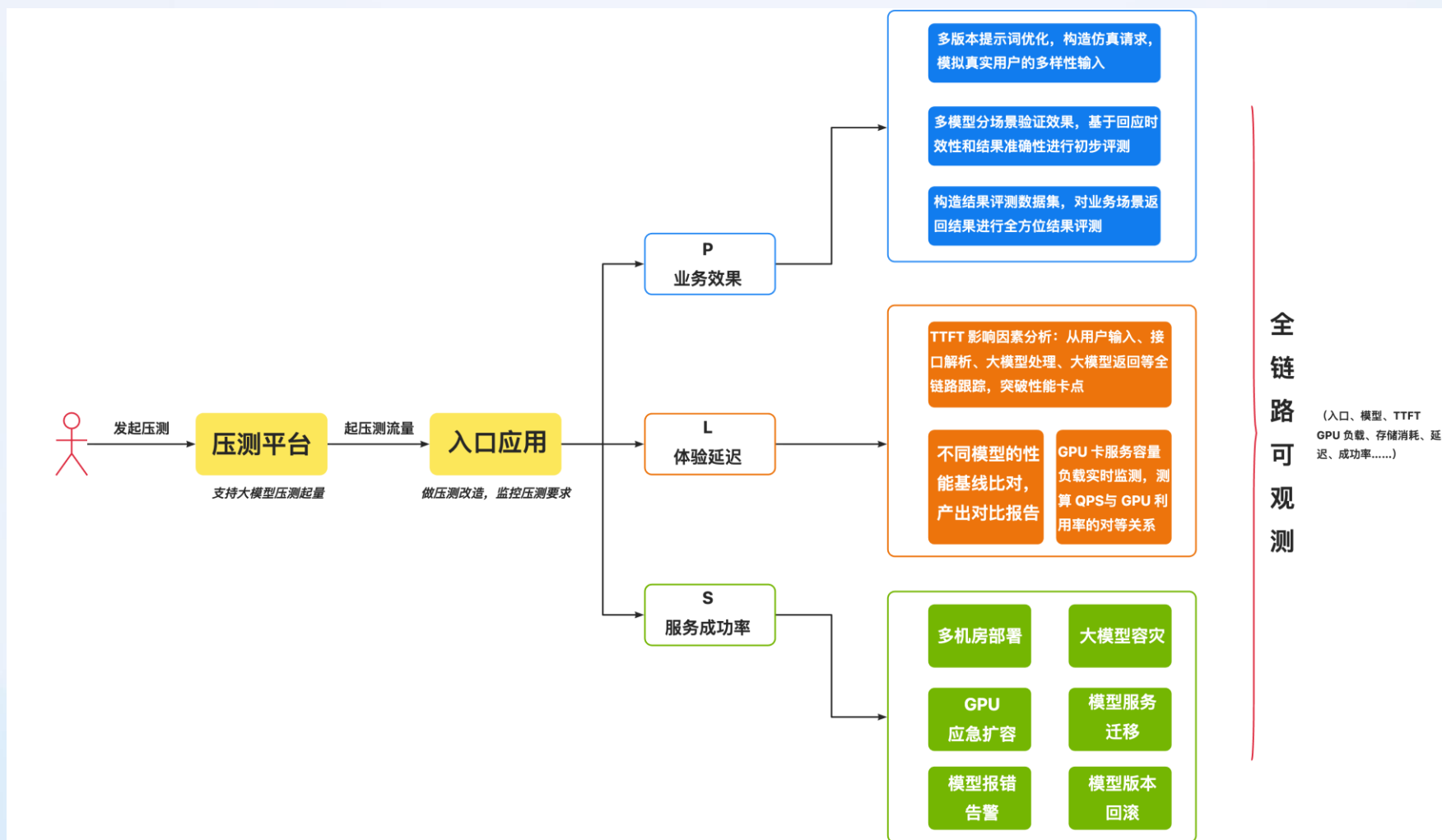


- 重点关注端到端 TTFT 指标，持续优化 AI 业务链路各阶段耗时情况
- 从基模选择、AI 服务多机房部署、GPU 容量供给、全链路延迟可观测等方面提升用户体验
- 为 AI 业务效果交付极速服务的同时提升 AI 保障效率

- 重点关注同一个业务场景在不同模型下的业务表现差异
- 从提示词调优、全链路可观测、输出结果评测等多方面进行业务表现评估
- 为用户交付优质结果的同时加强大模型风险管理

- 重点关注 AI 服务层、AI 基础设施层的稳定性
- 从模型部署成功率、AI 可观测、超大 GPU 集群调度等方面加强技术风险能力建设
- 为 AI 服务保驾护航的同时兼顾 AI 成本

AI 链路压测实践



📍 北京

QCon

全球软件开发大会

会议时间：4月10-12日

- 大模型赋能 AIOps
- 越挫越勇的大前端
- 云上业务架构演进
- 多模态大模型及应用
- 海外 AI 应用创新实践

📍 北京

AiCon

全球人工智能开发与应用大会

会议时间：6月27-28日

- 端侧智能
- AI Agent
- 多模态大模型
- 金融+大模型

📍 上海

QCon

全球软件开发大会

会议时间：10月23-25日

- AI Agent
- 大模型训练推理
- 端侧 AI
- 搜推深度融合
- 数智金融

4月

5月

6月

8月

10月

12月

📍 上海

AiCon

全球人工智能开发与应用大会

会议时间：5月23-24日

- 通用大模型
- AI Agent
- 垂直领域应用
- 模型可解释性

📍 深圳

AiCon

全球人工智能开发与应用大会

会议时间：8月22-23日

- 模型效率与部署
- 多模态大模型
- 大模型安全
- 智能硬件

📍 北京

AiCon

全球人工智能开发与应用大会

会议时间：12月19-20日

- 通用大模型
- 智能硬件
- LMOps
- 具身智能

极客邦科技 2025 年会议规划

促进软件开发及相关领域知识与创新的传播



参会咨询



查看会议

THANKS

大模型正在重新定义软件

Large Language Model Is Redefining The Software

