

Yita：基于数据流的大数据计算引擎

郑龙，Ph.D.，CTO

中兴飞流信息科技有限公司

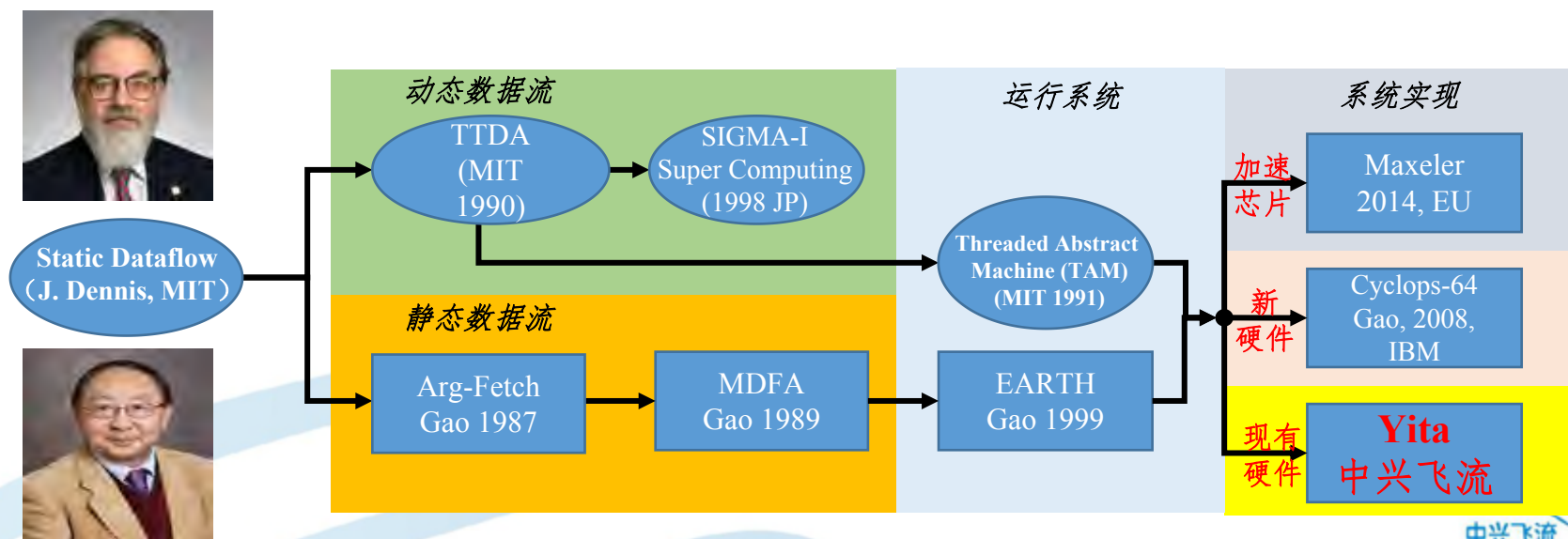
数据流起源与发展

• 数据流的起源

- 数据流基础理论由MIT的Jack Dennis教授于上世纪七十、八十年代提出；
- Jack Dennis教授是美国工程院院士，由于其在数据流理论的贡献获得2012年的IEEE冯诺依曼奖章。

• 数据流的发展

- 高光荣教授跟随Jack，坚持发展30余年，并成为数据流技术的主要代表人物；
- 高教授由于在数据流及处理器设计对中国的贡献，于2012年获得CCF海外杰出贡献奖。



控制流模型（冯诺依曼模型）

$(1 + 2) * (3 + 4) + (5 + 6) * (7 + 8)$

↑
PC

PC →

LOAD R1, 1
LOAD R2, 2
ADD R1, R2
LOAD R3, 3
LOAD R4, 4
.....

指令栈

并行控制流模型 (并行冯诺依曼模型)

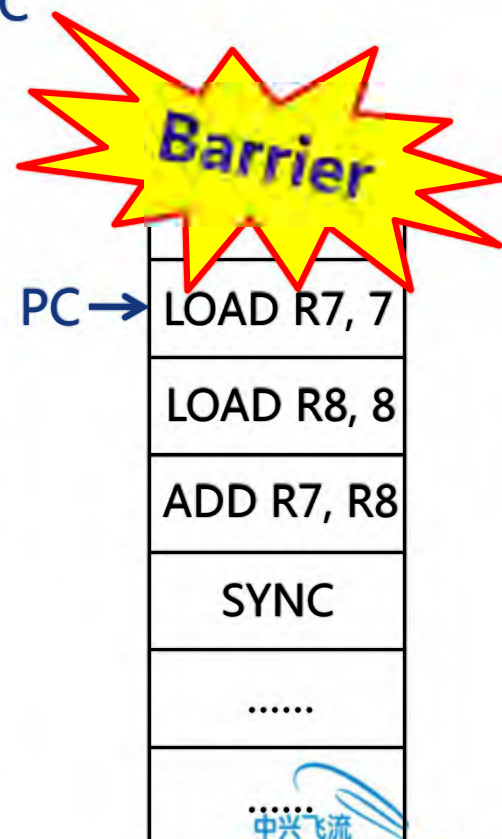
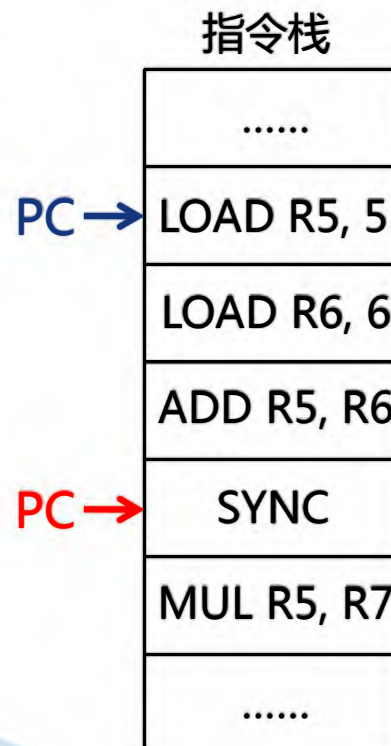
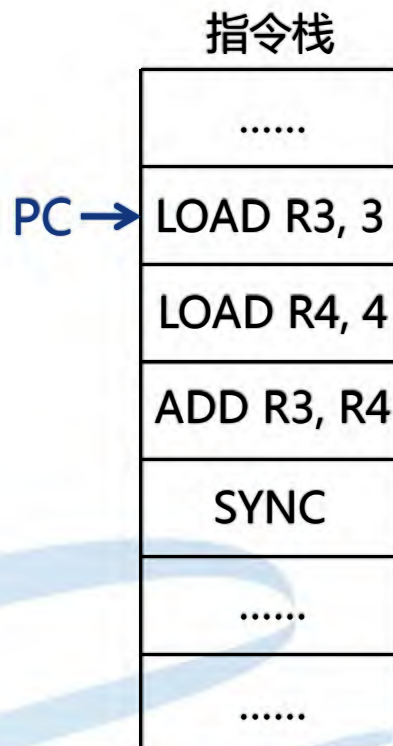
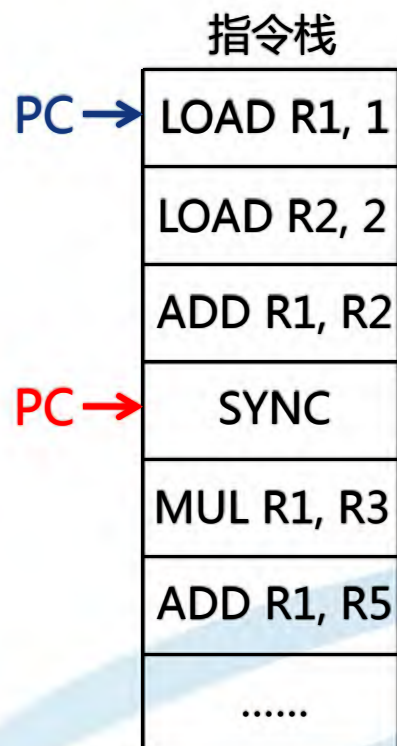
$$(1 + 2) * (3 + 4) + (5 + 6) * (7 + 8)$$

↑
PC

↑
PC

↑
PC

↑
PC



并行控制流模型（并行冯诺依曼模型）

$$(1 + 2) * (3 + 4) + (5 + 6) * (7 + 8)$$

↑
PC

↑
PC

↑
PC

↑
PC

指令栈

LOAD R1, 1
LOAD R2, 2
ADD R1, R2
SYNC
MUL R1, R3
SYNC
ADD R1, R5

PC →

指令栈

.....
LOAD R3, 3
LOAD R4, 4
ADD R3, R4
SYNC
.....
.....

PC →

指令栈

.....
LOAD R5, 5
LOAD R6, 6
ADD R5, R6
SYNC
MUL R5, R7
SYNC

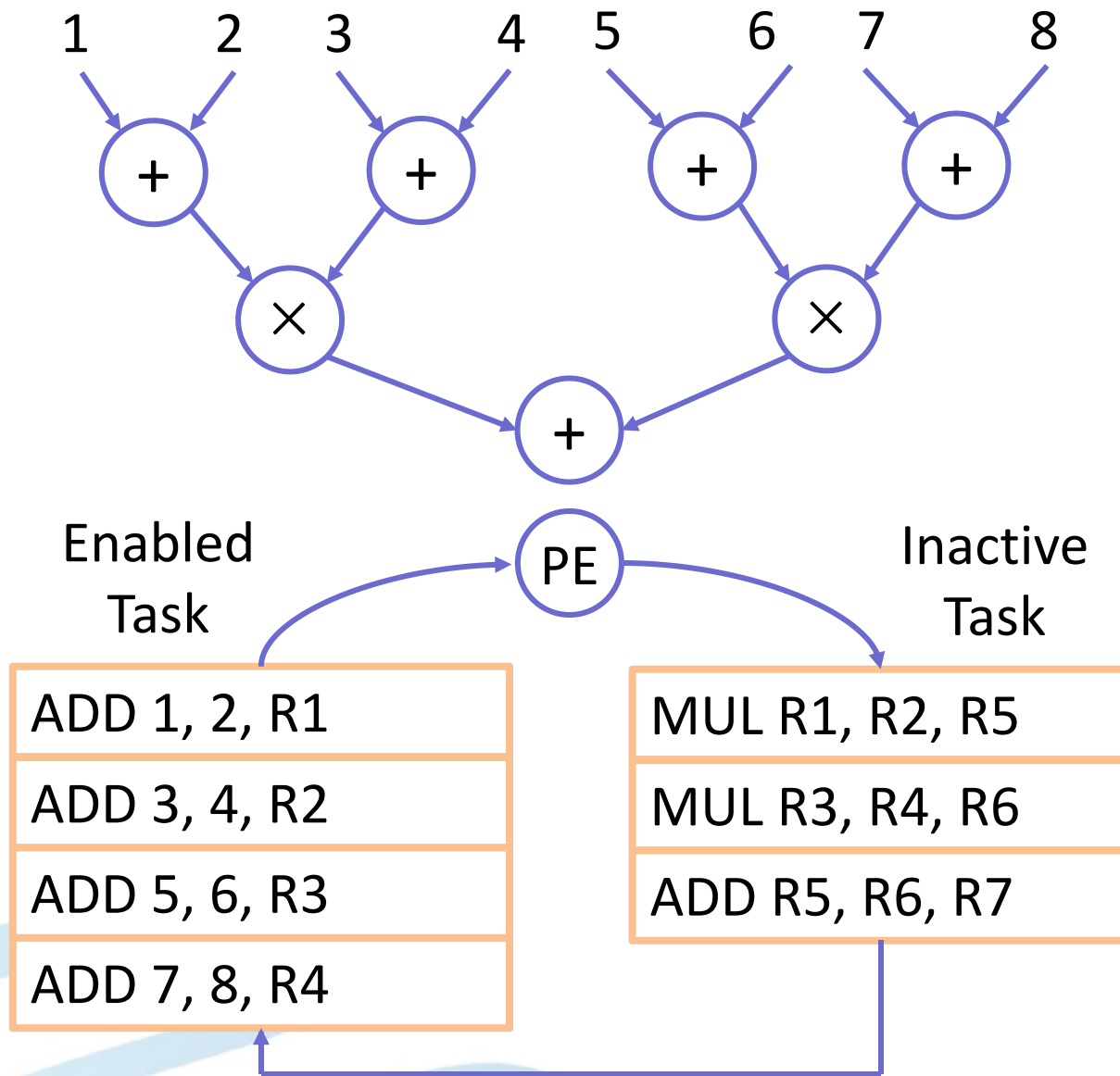
PC →

指令栈

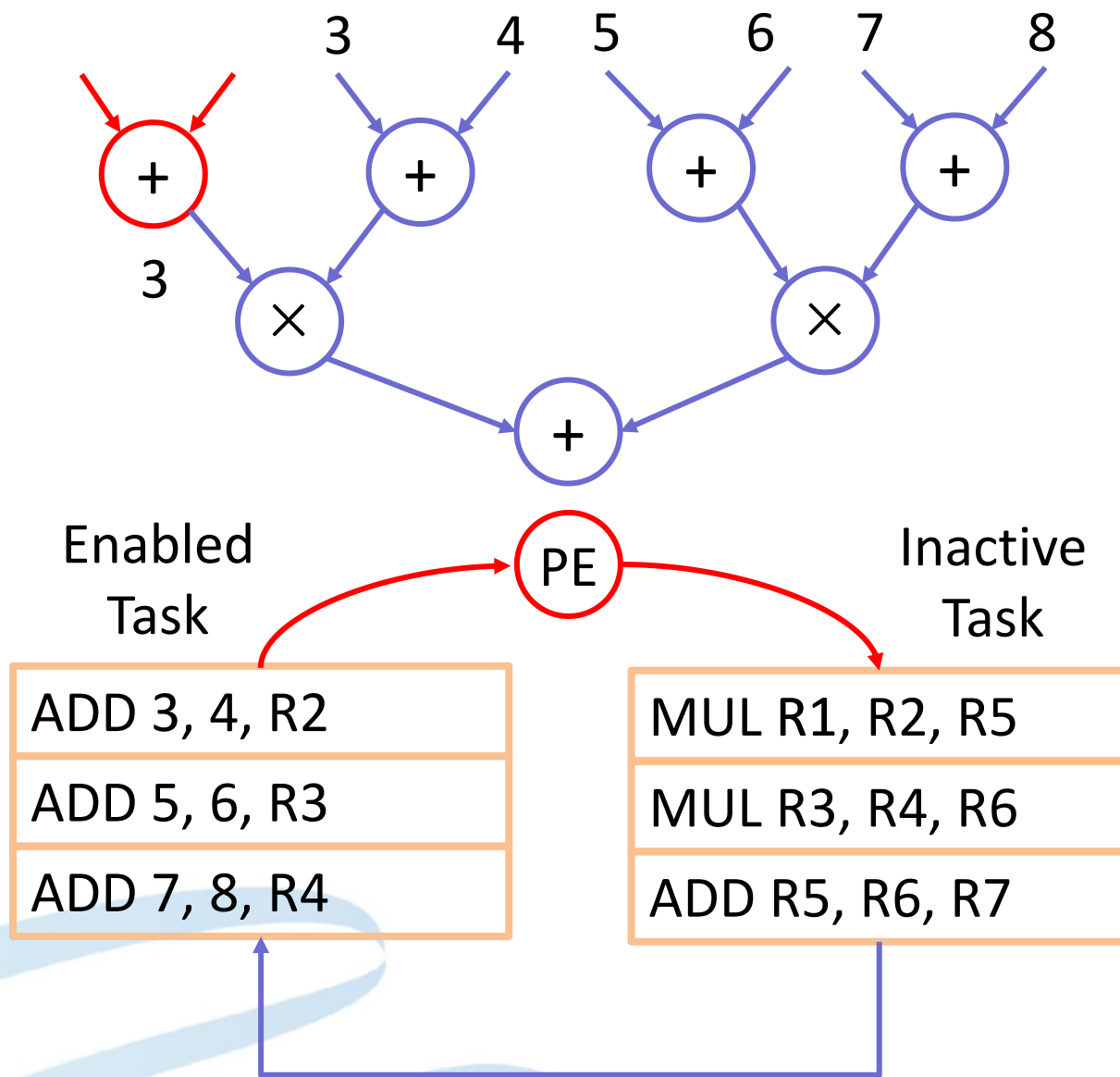
.....
LOAD R7, 7
LOAD R8, 8
ADD R7, R8
SYNC
.....
.....

PC →

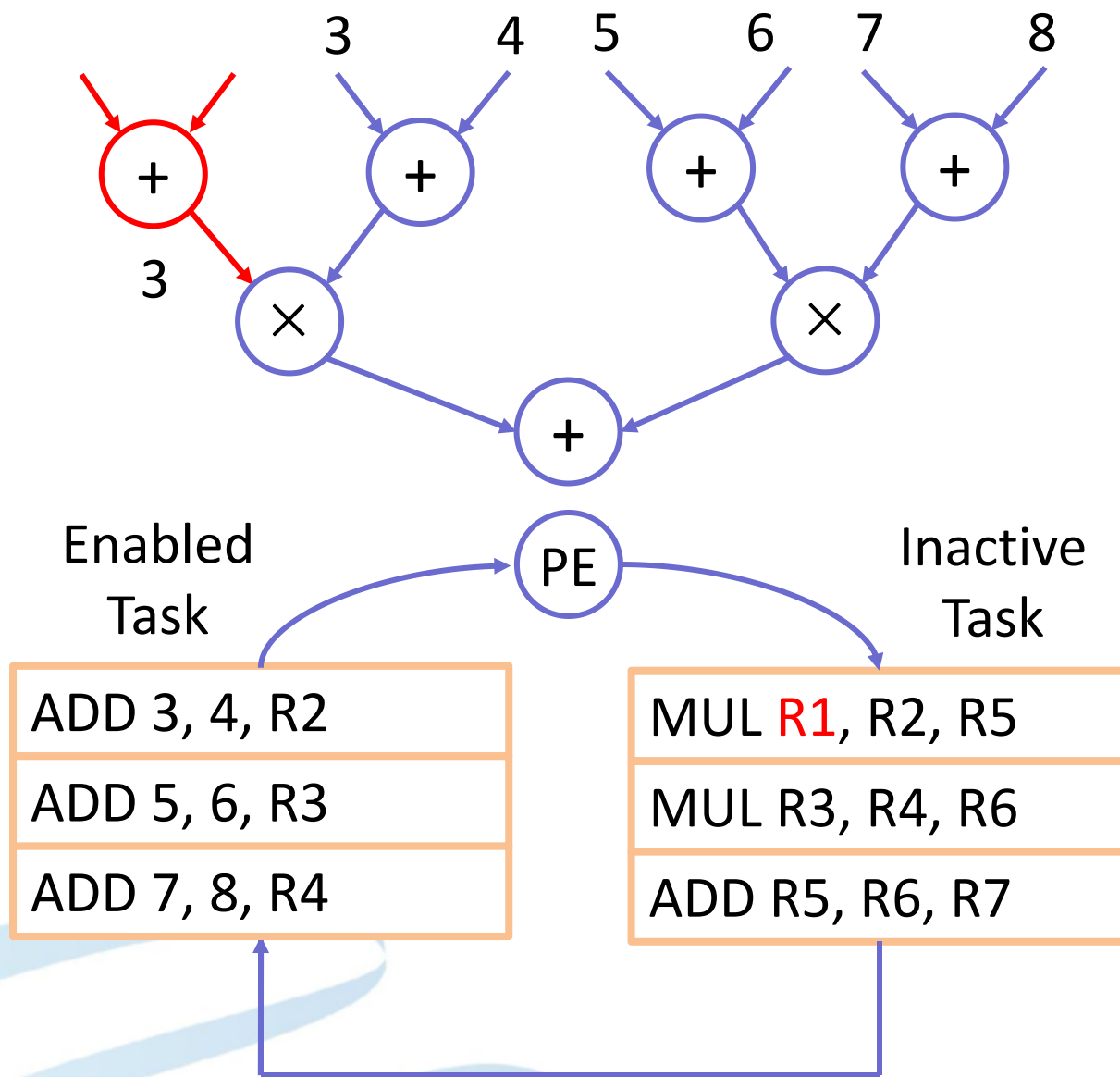
数据流模型



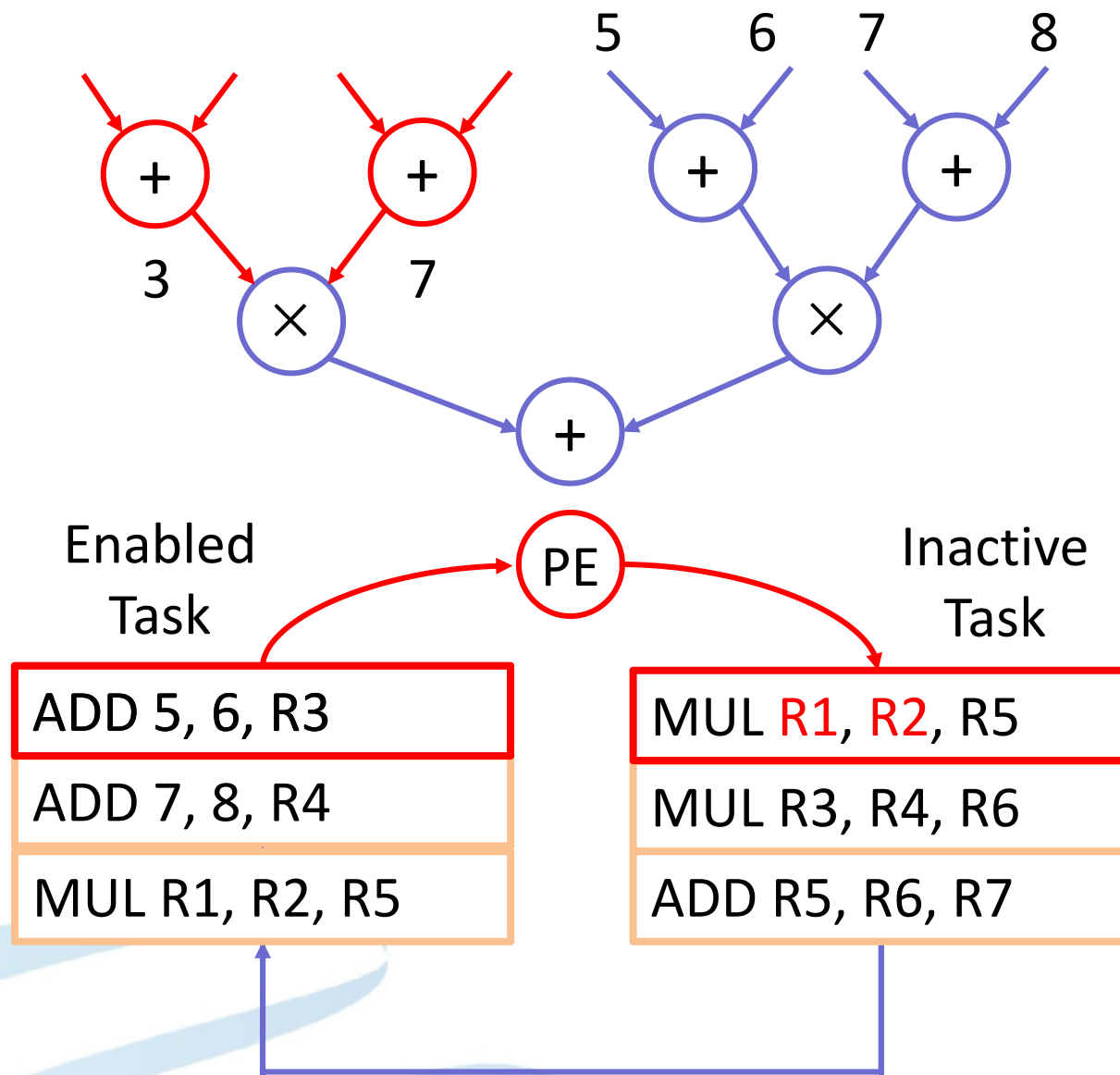
数据流模型



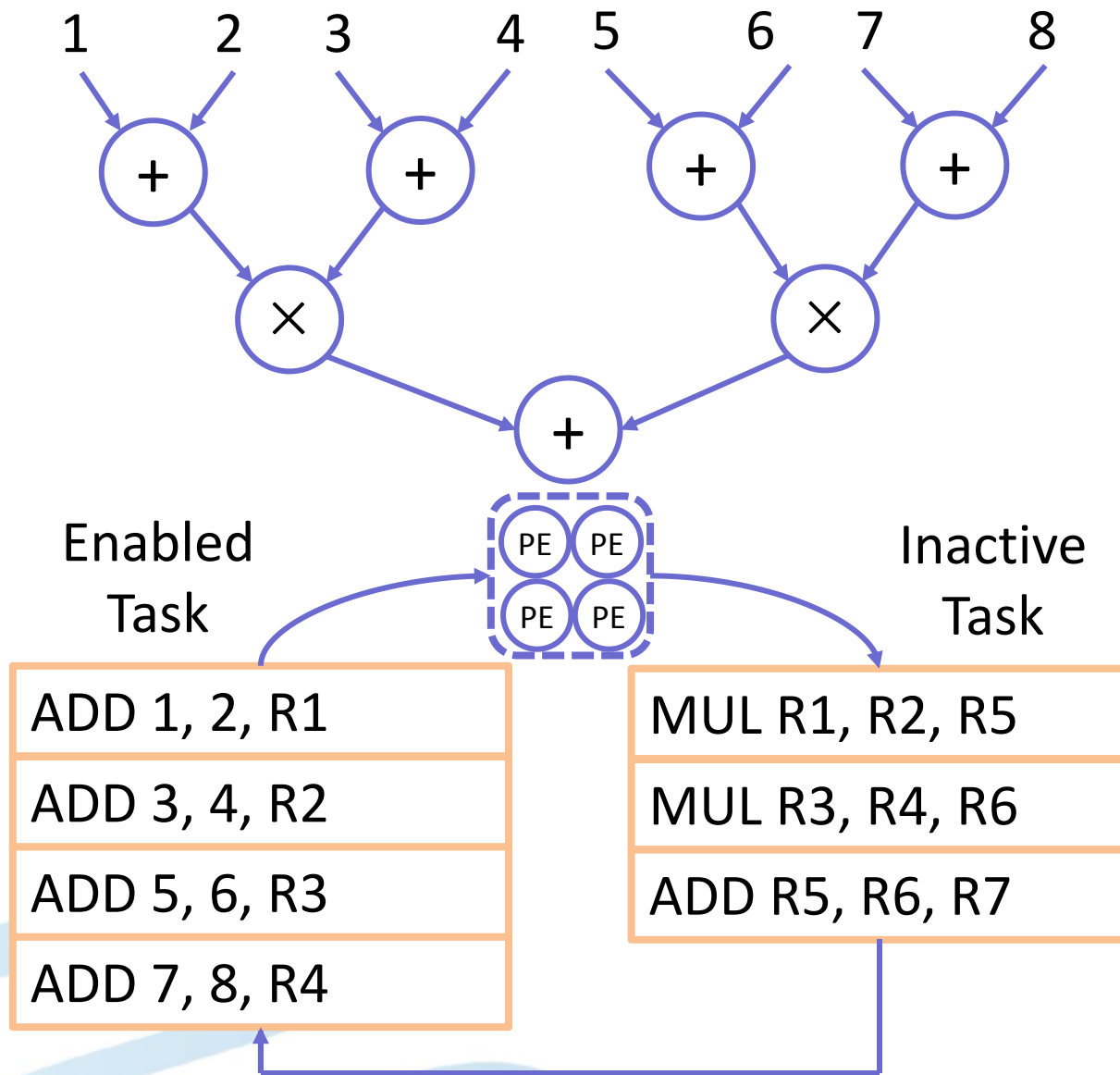
数据流模型



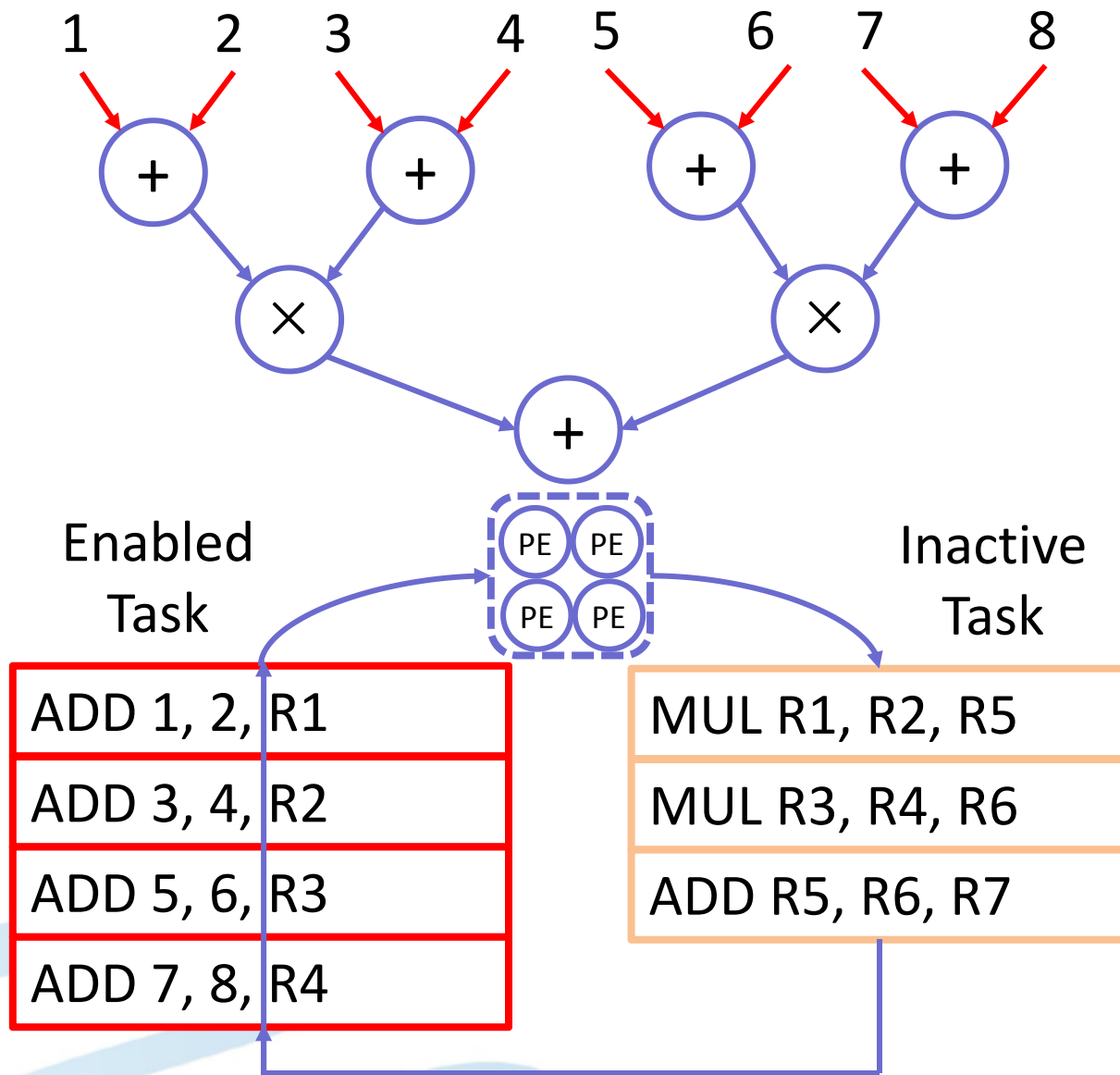
数据流模型



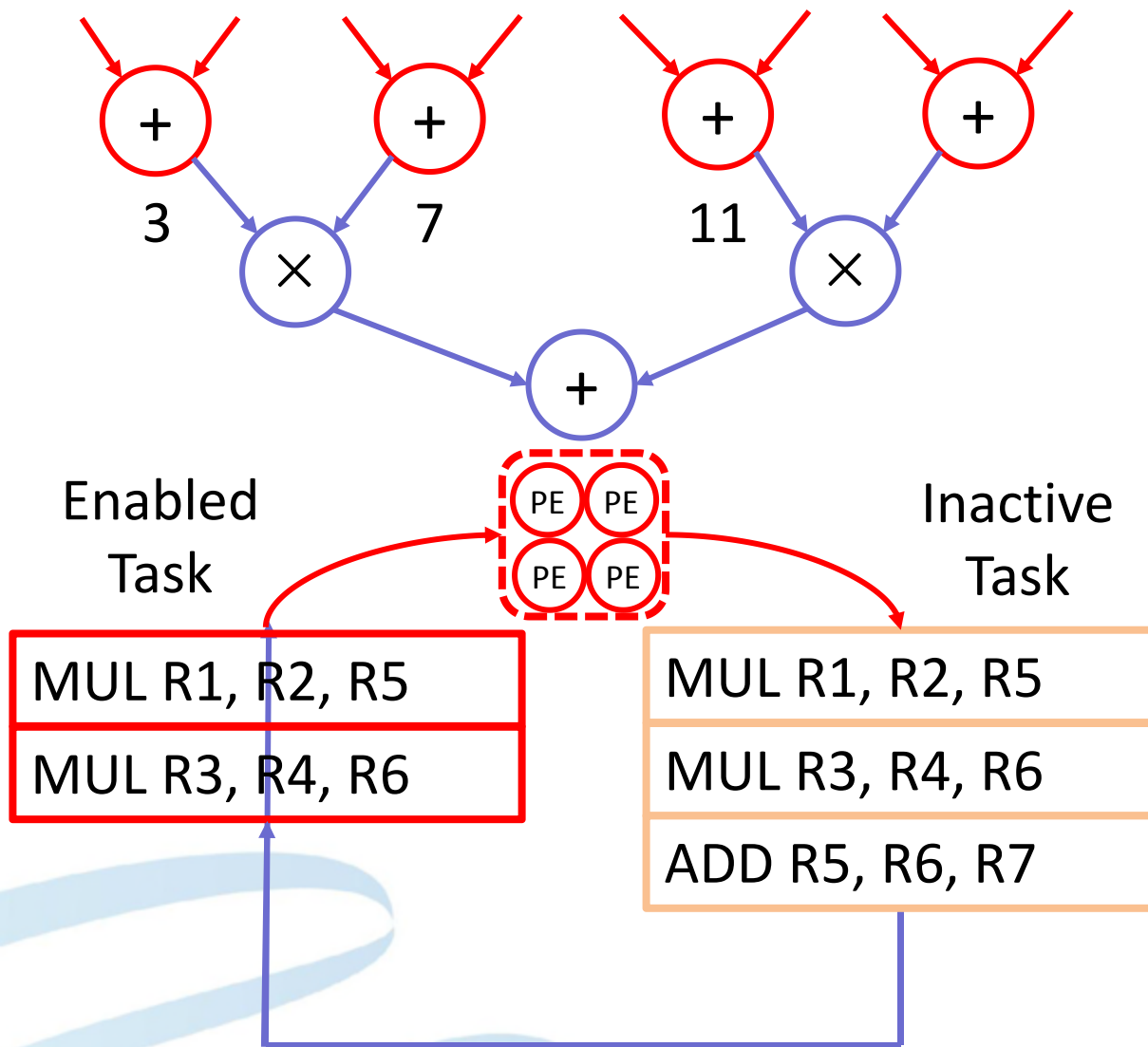
并行数据流模型



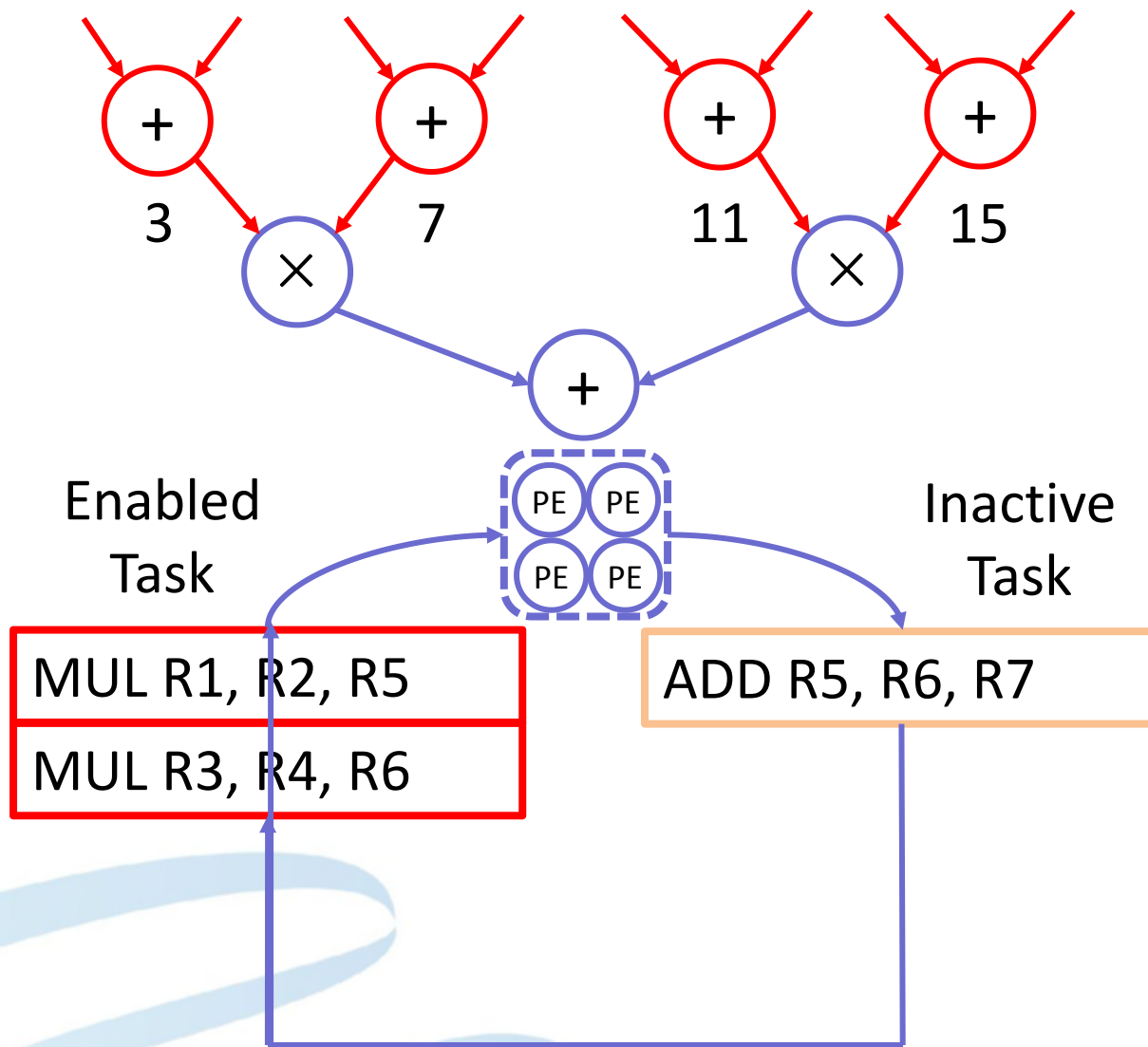
并行数据流模型



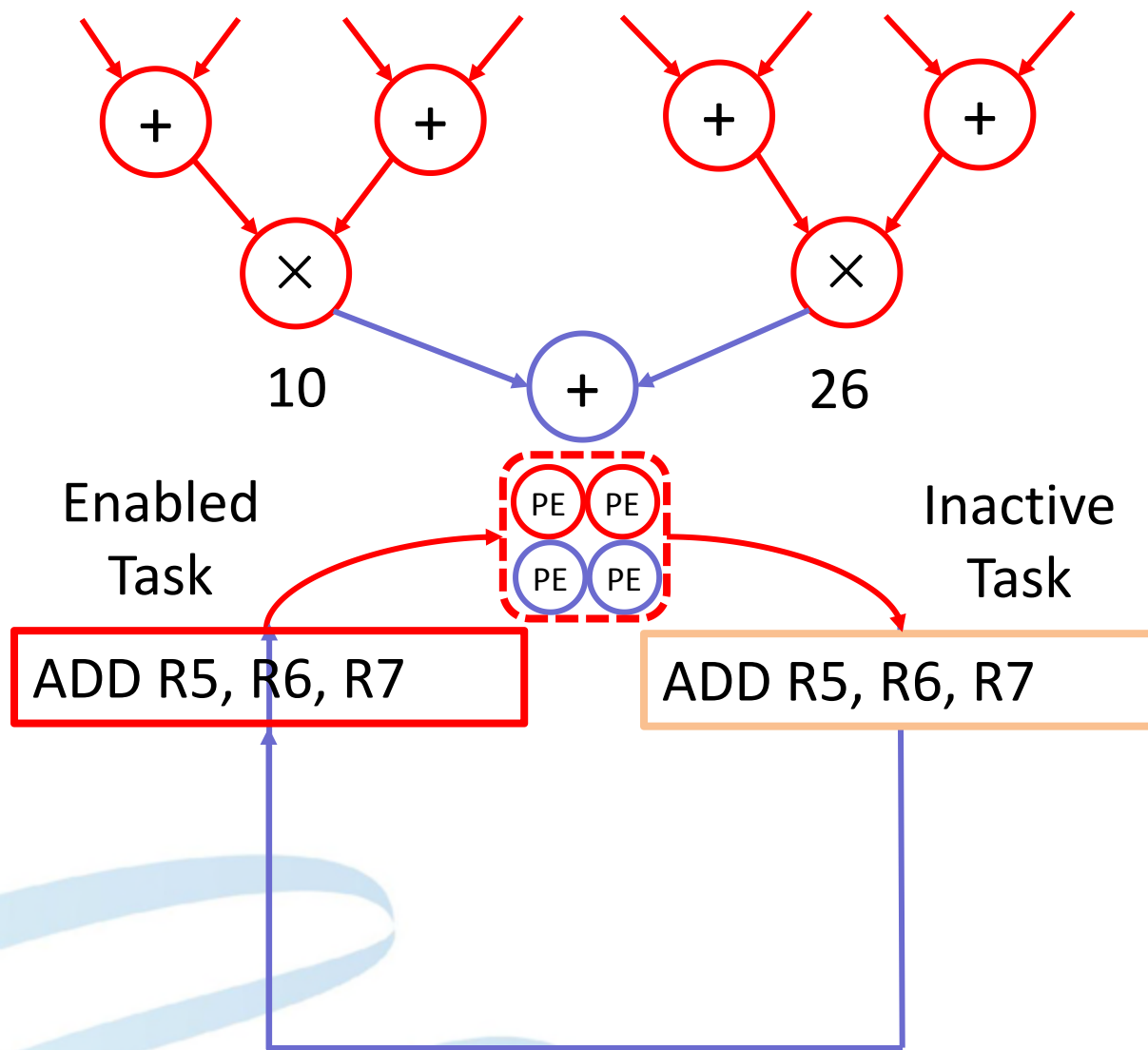
并行数据流模型



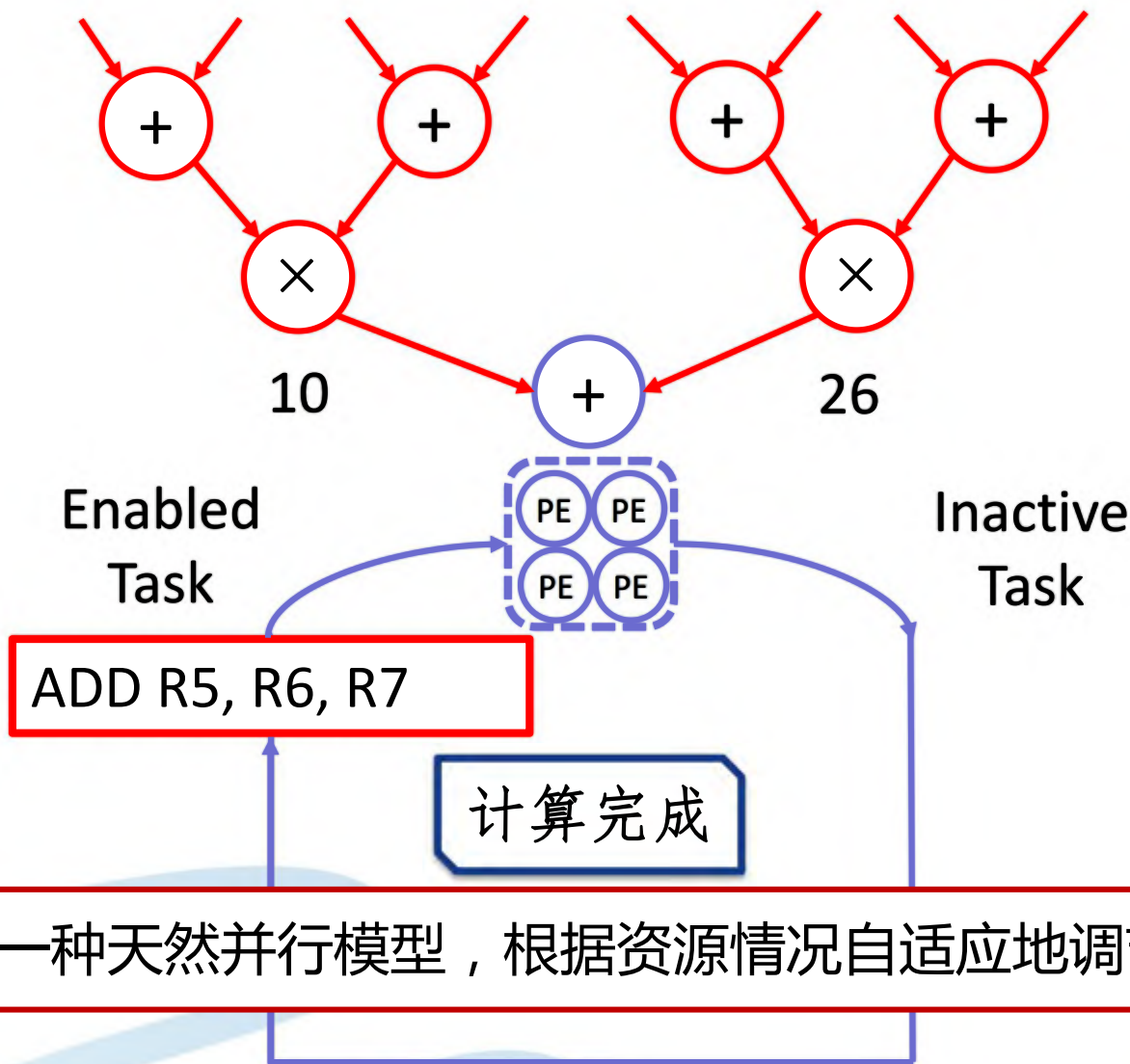
并行数据流模型



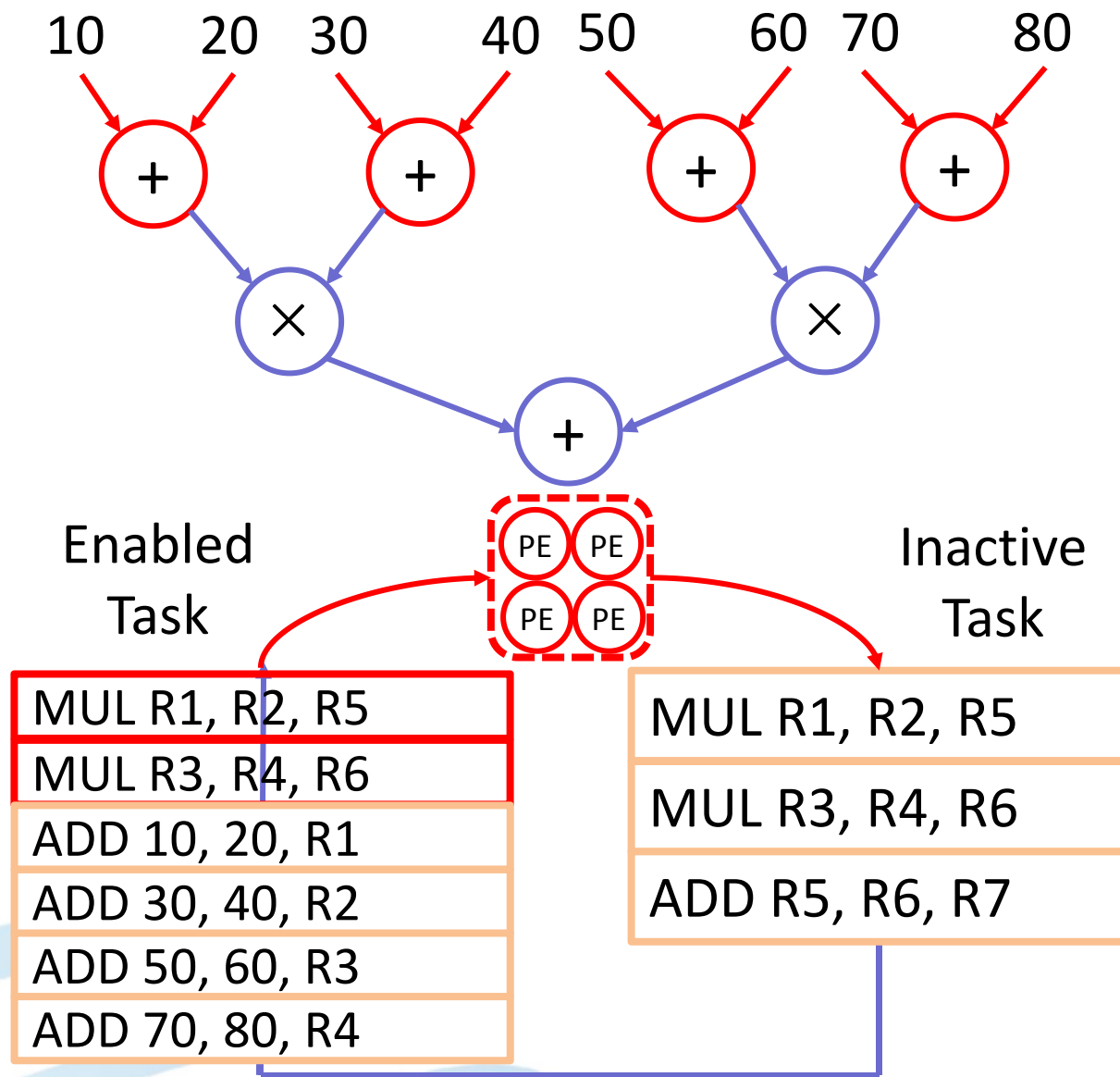
并行数据流模型



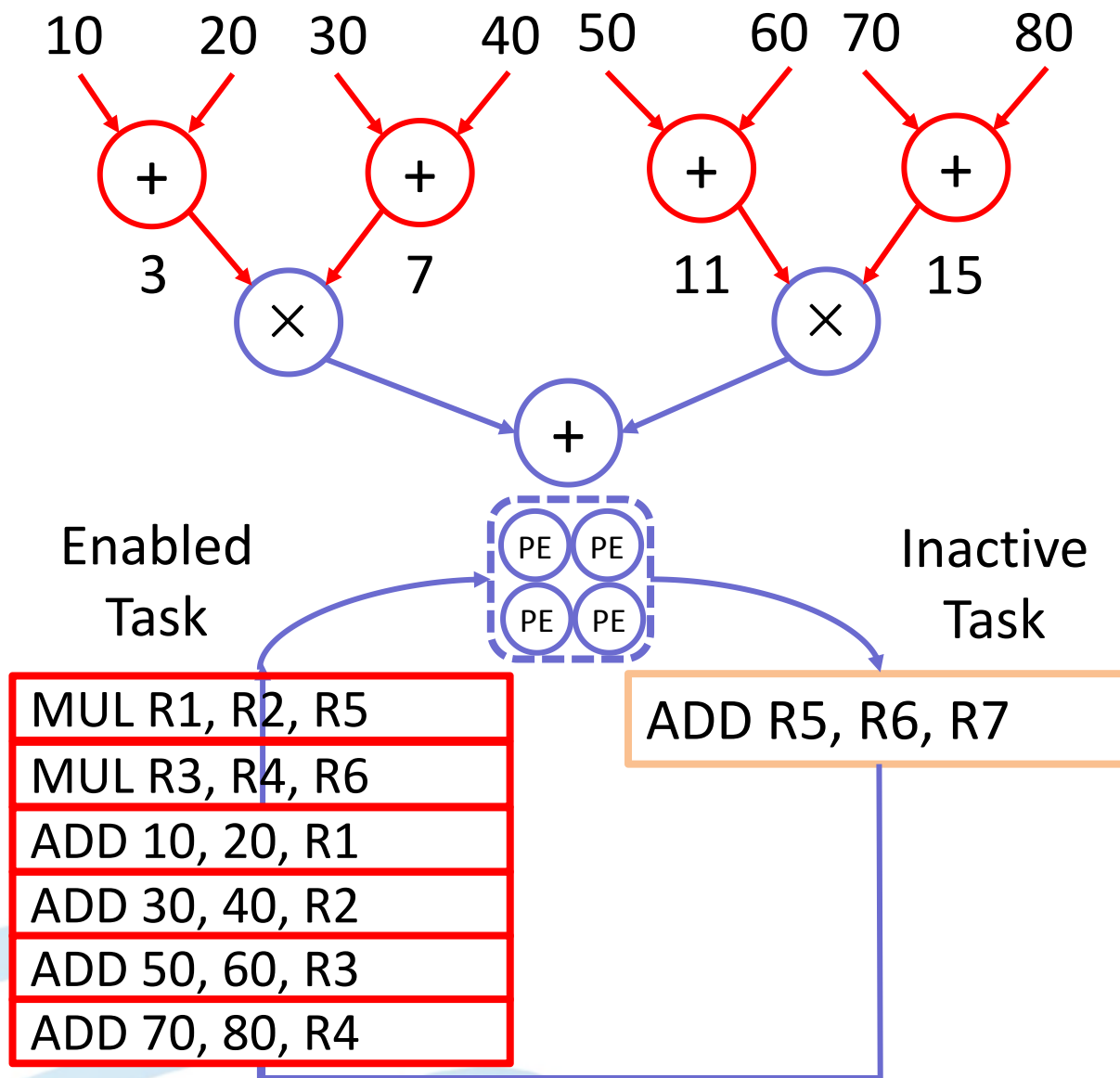
并行数据流模型



并行数据流模型-流水线并行

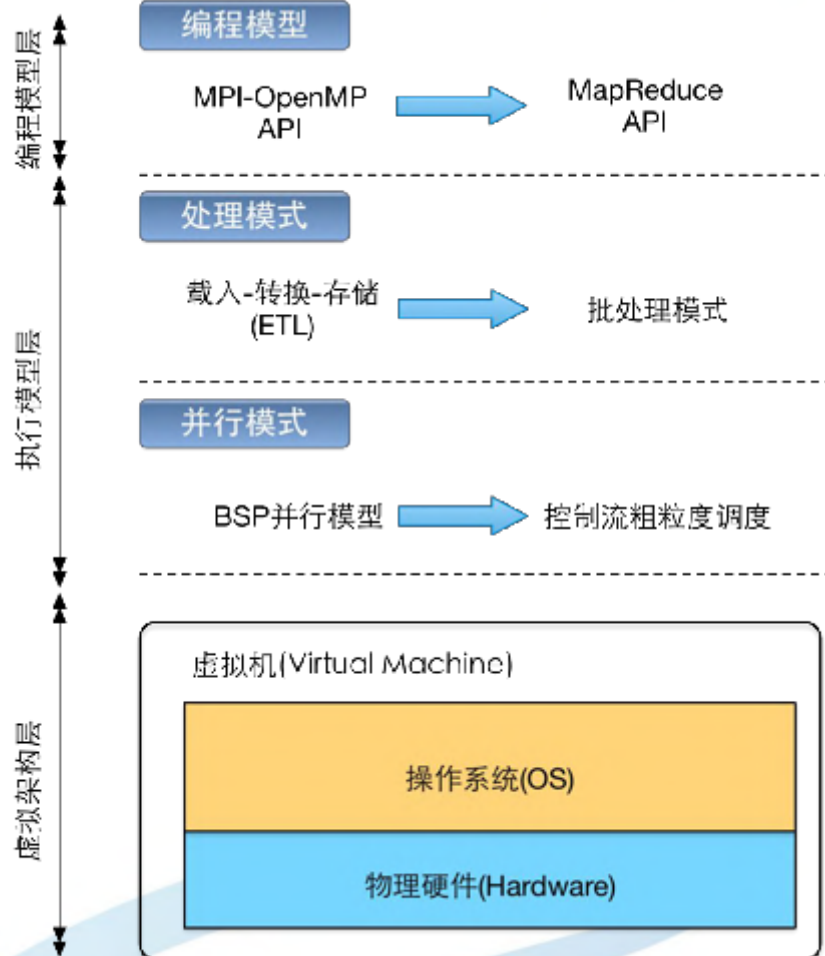


并行数据流模型-流水线并行

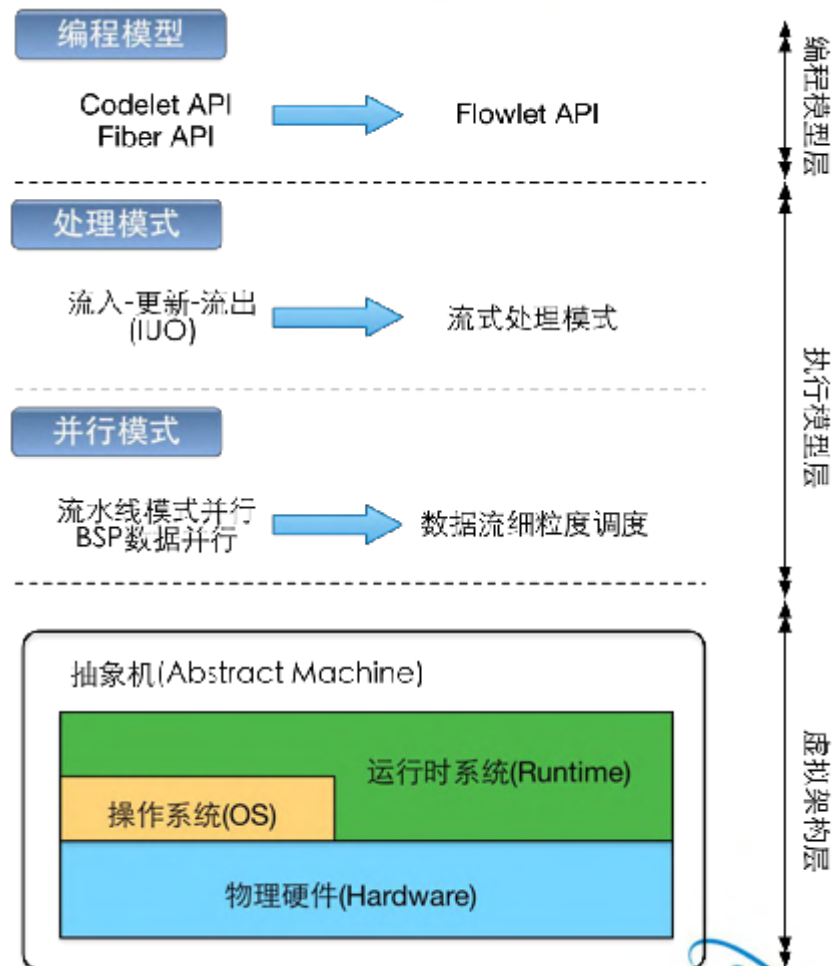


Yita与传统控制流大数据对比

基于控制流的大数据系统



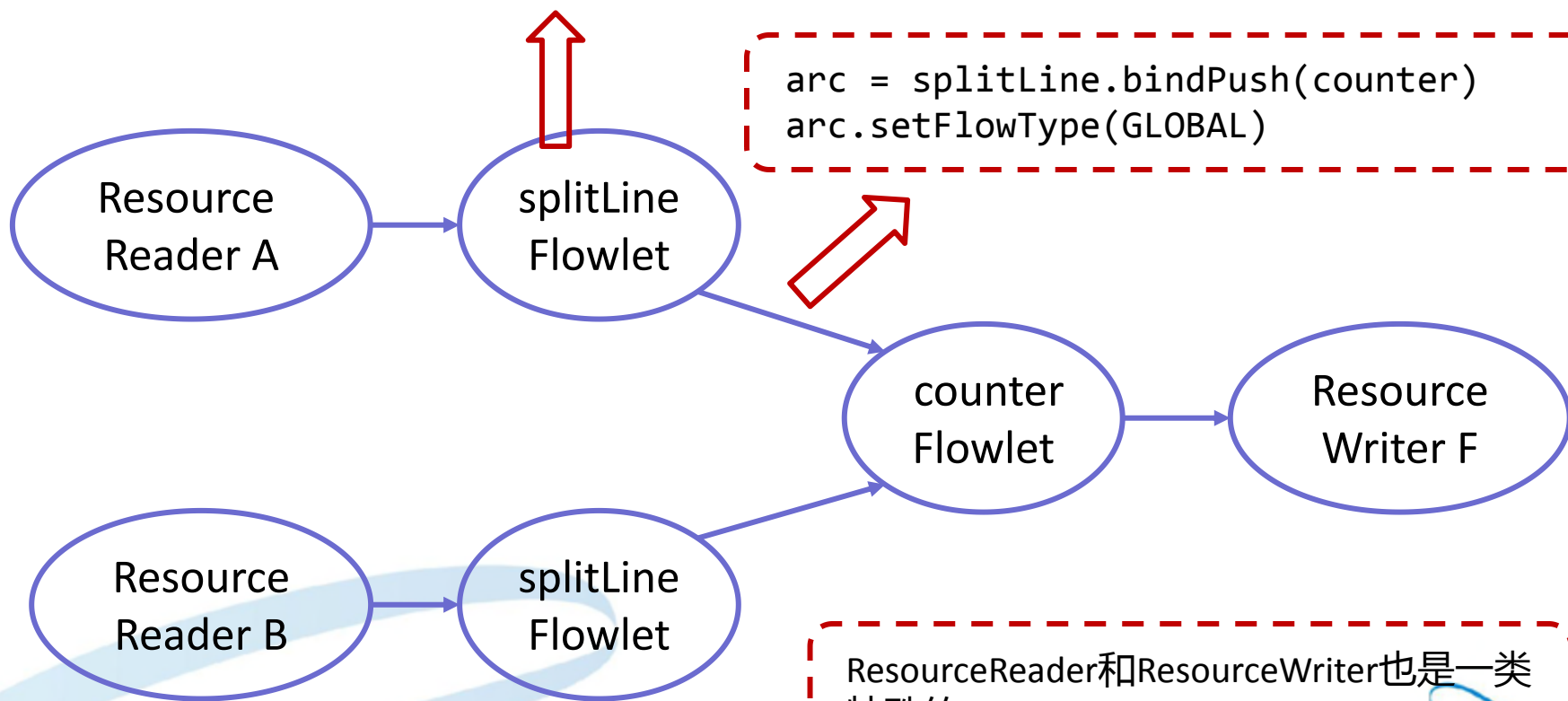
基于数据流的Yita系统



Yita基于Flowlet的编程模型

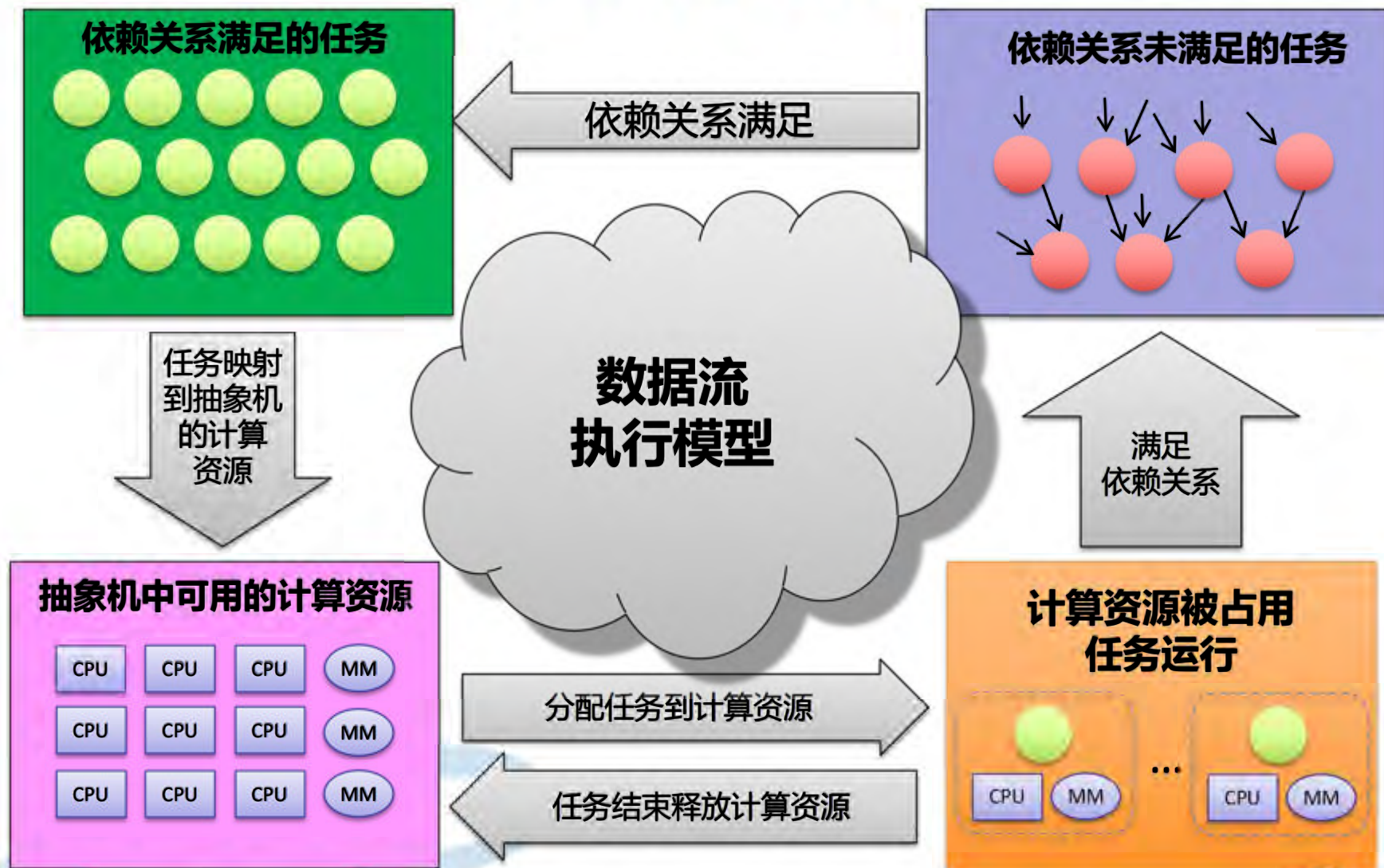
```
def splitLine(lineNum, line, flow):  
    for word in line.split(" "):  
        flow.push([word, 1])
```

```
arc = splitLine.bindPush(counter)  
arc.setFlowType(GLOBAL)
```



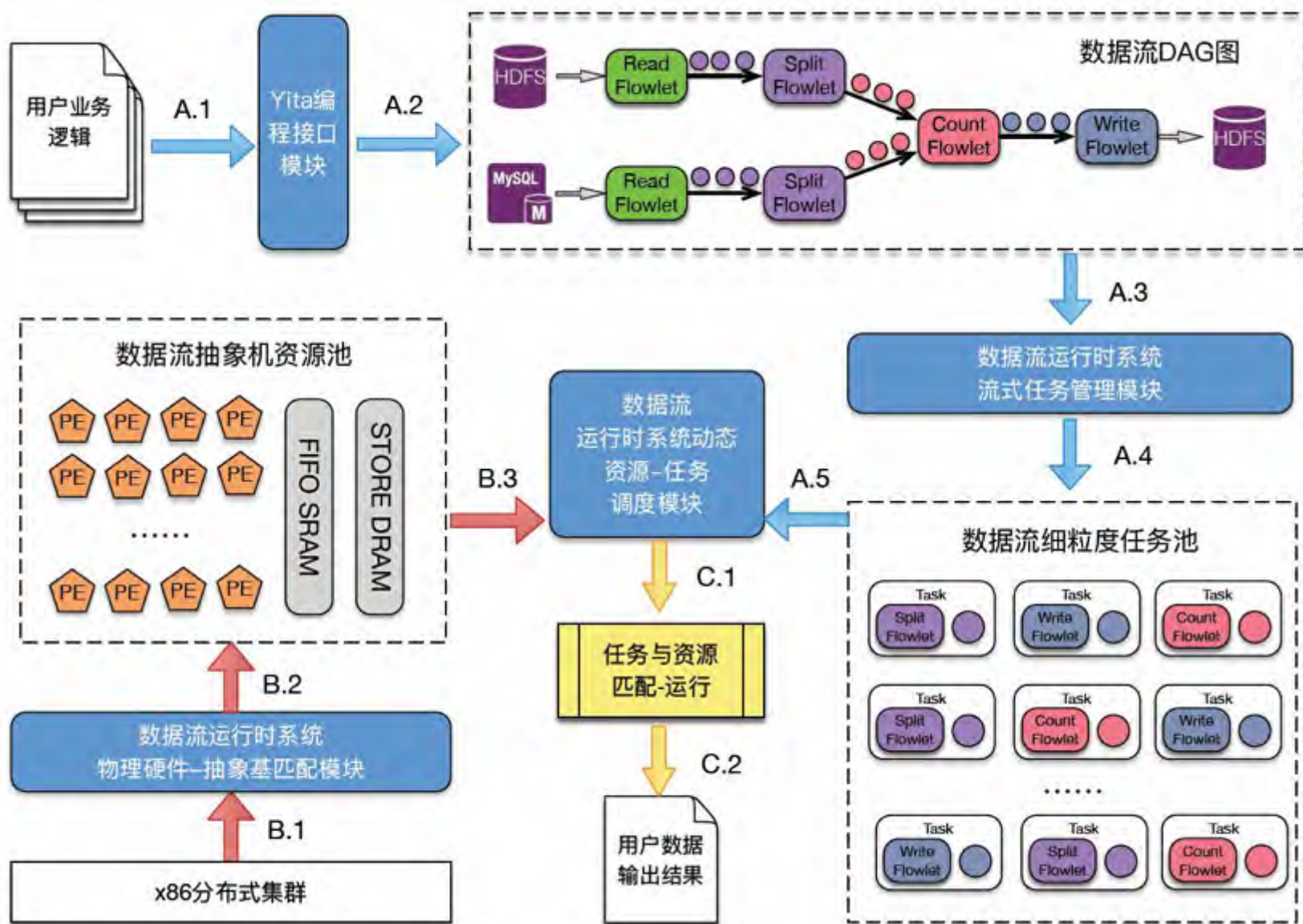
ResourceReader和ResourceWriter也是一类特殊的Flowlet。

Yita执行模型

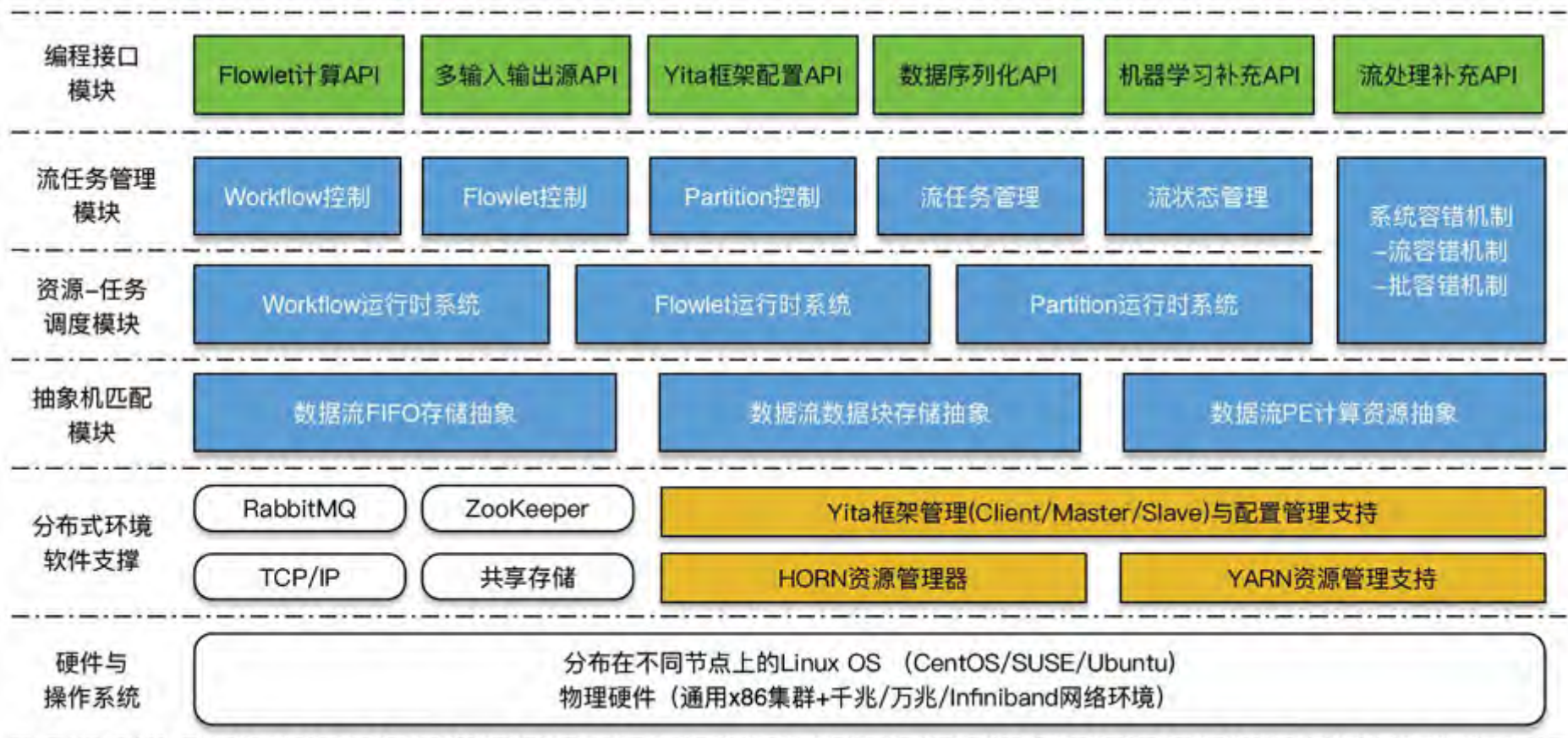


细粒度-流水线-数据流并行模式示意图

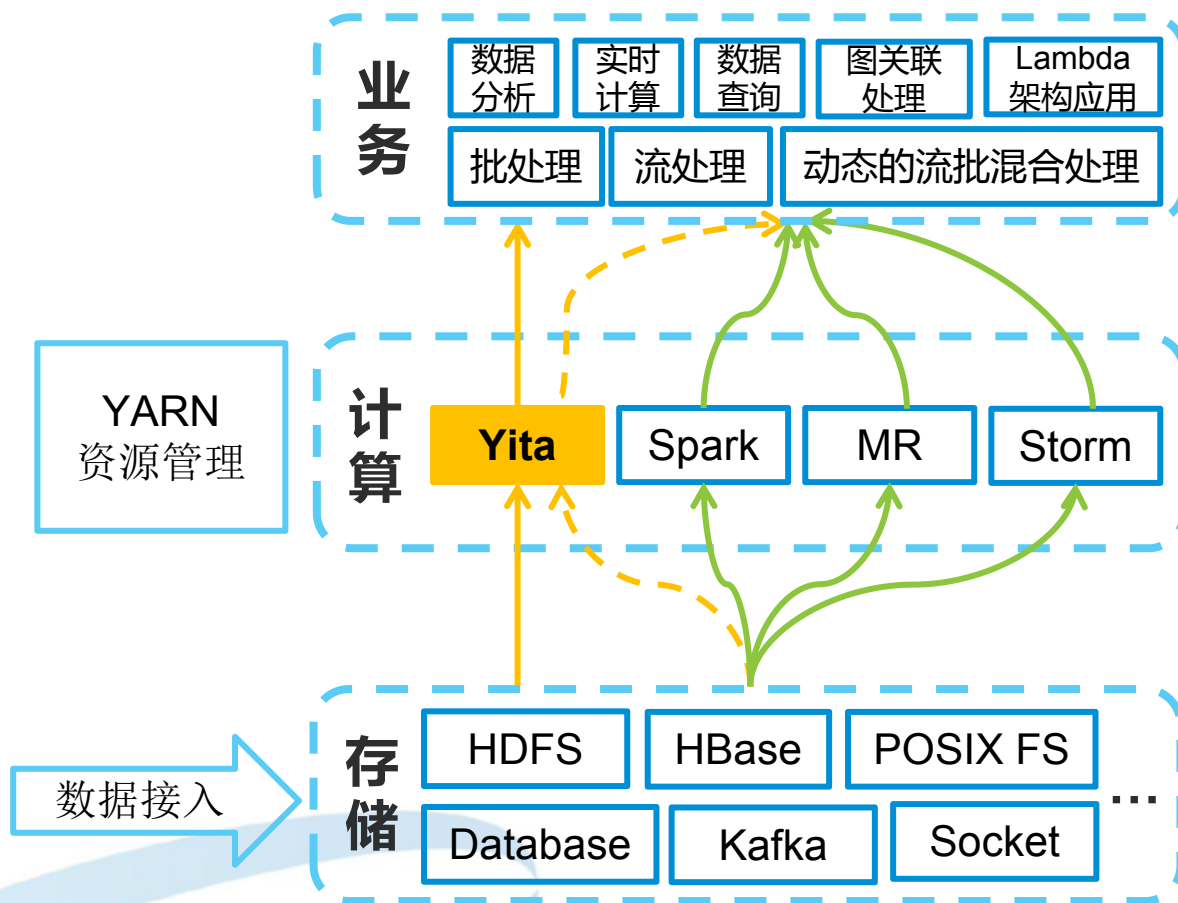
Yita系统运行流程



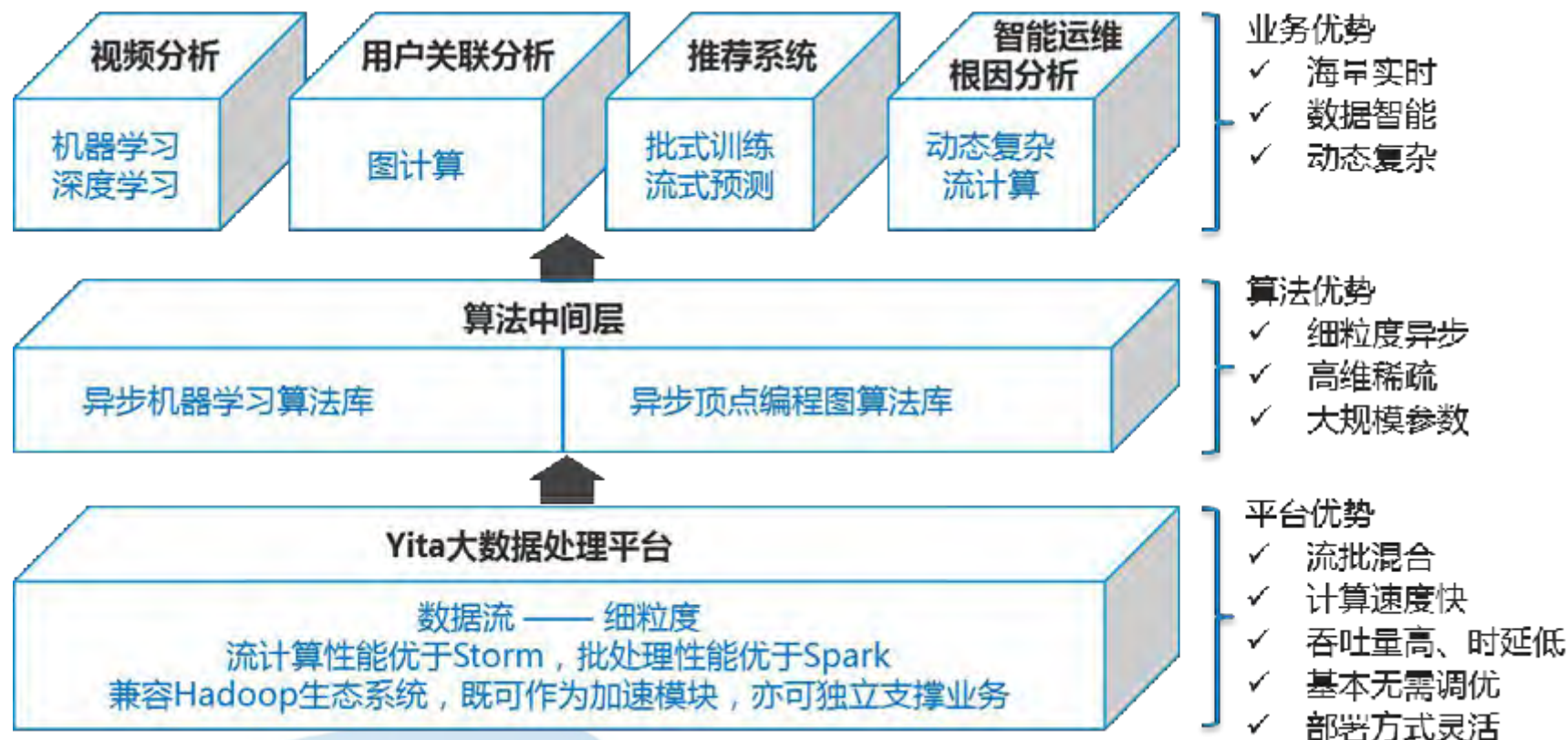
Yita系统组件架构



Yita与大数据生态圈



Yita的业务场景





省内存
细粒度



耗内存
粗粒度

集群硬

控制流 (以

L=

Yita : 细粒

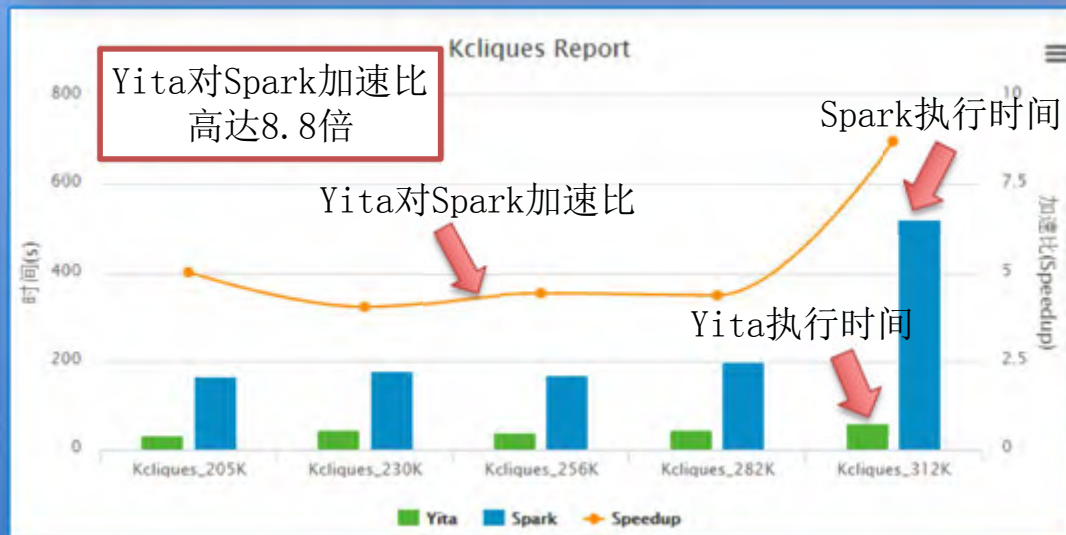
变

L=

Yita : 细粒

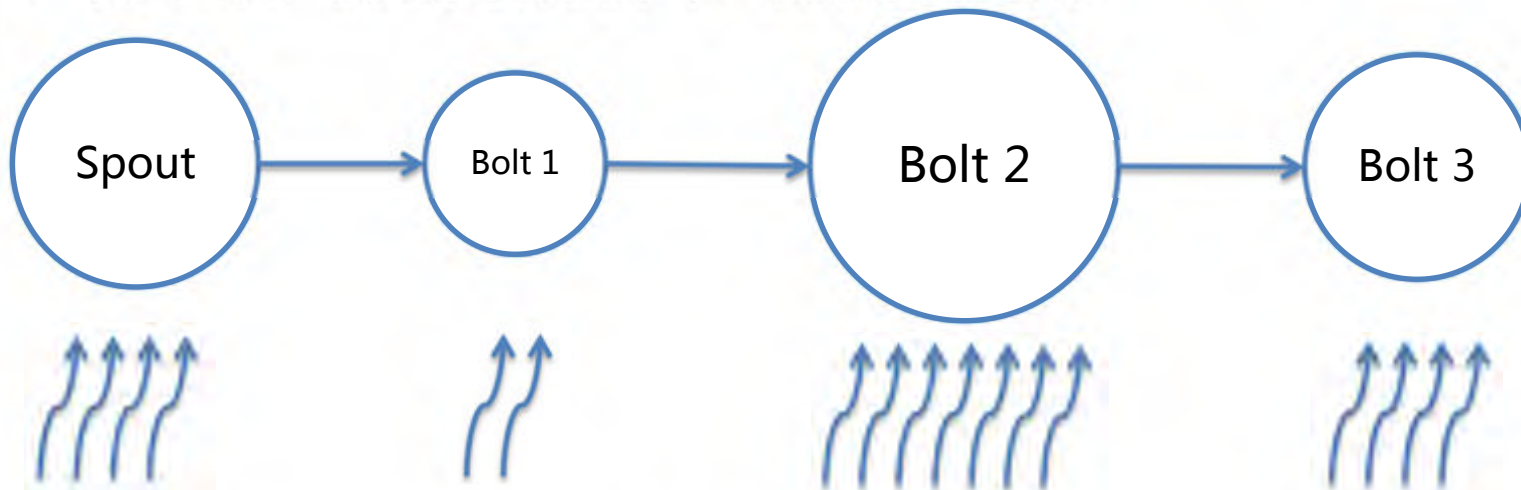
L

L=

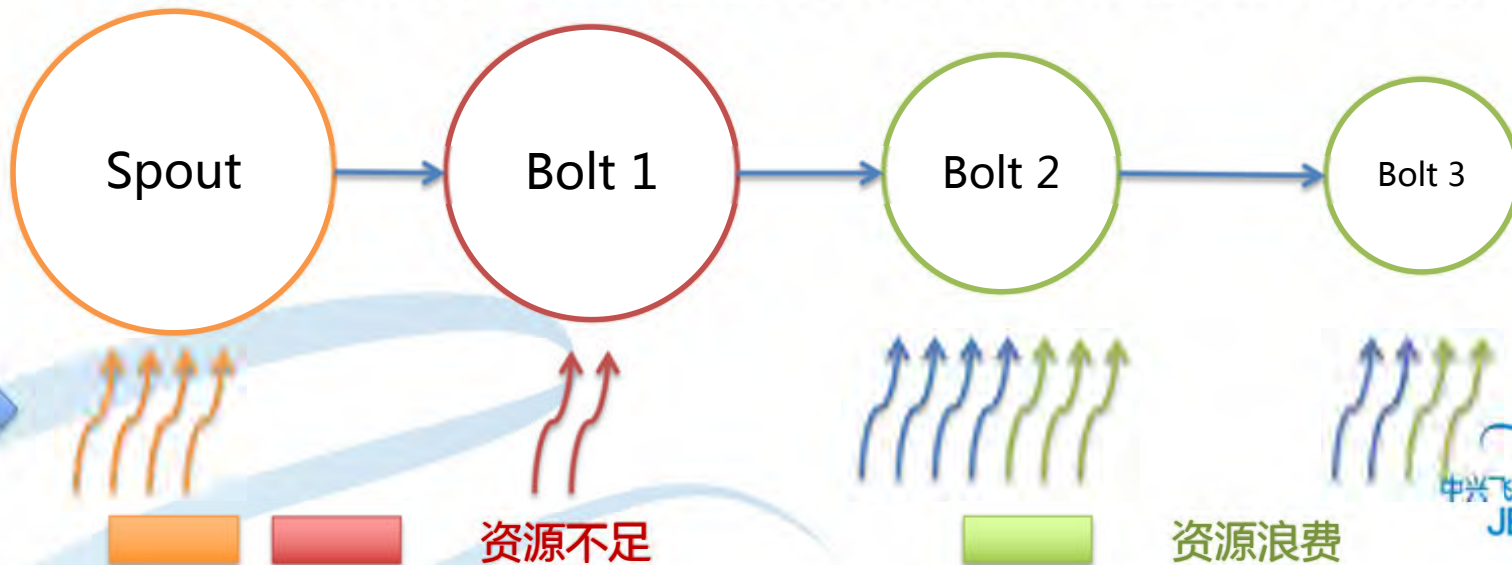


Yita优势：动态实时 (1/2)

- 按照经验，预先静态定义好每一级所需计算资源



- 实际情况中，每一级负载往往处于动态变化中，影响吞吐和延迟





高吞吐
低延迟



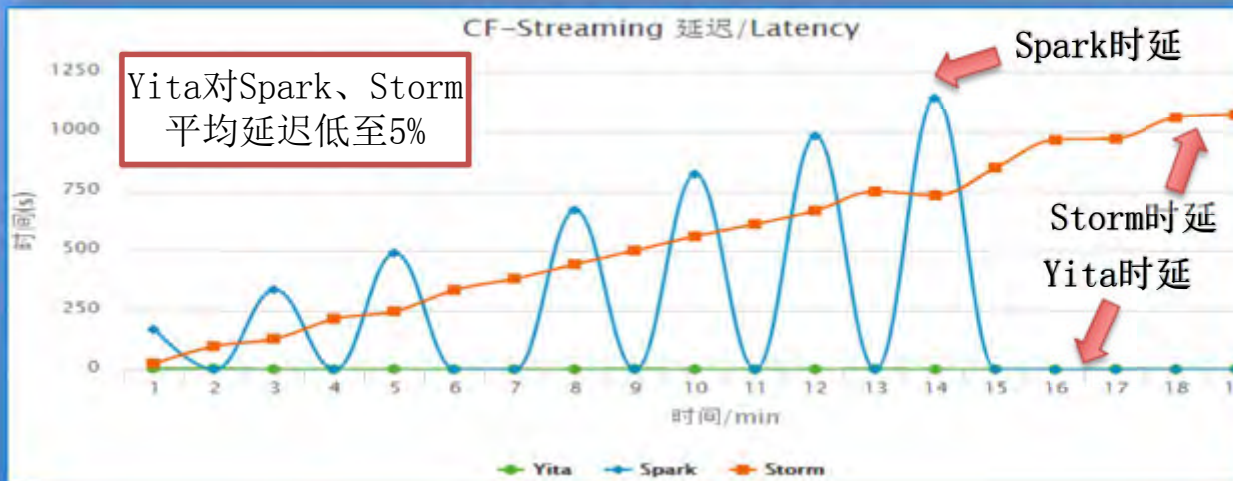
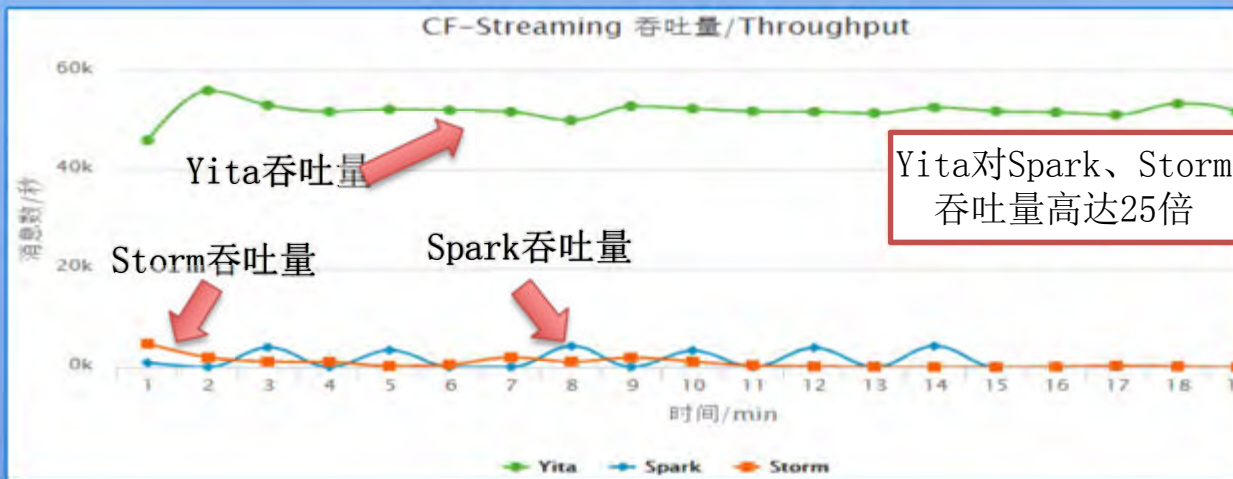
高吞吐
高延迟



低吞吐
低延迟



2016-12-10



任务;
要求。





异步
细粒度



同步
粗粒度

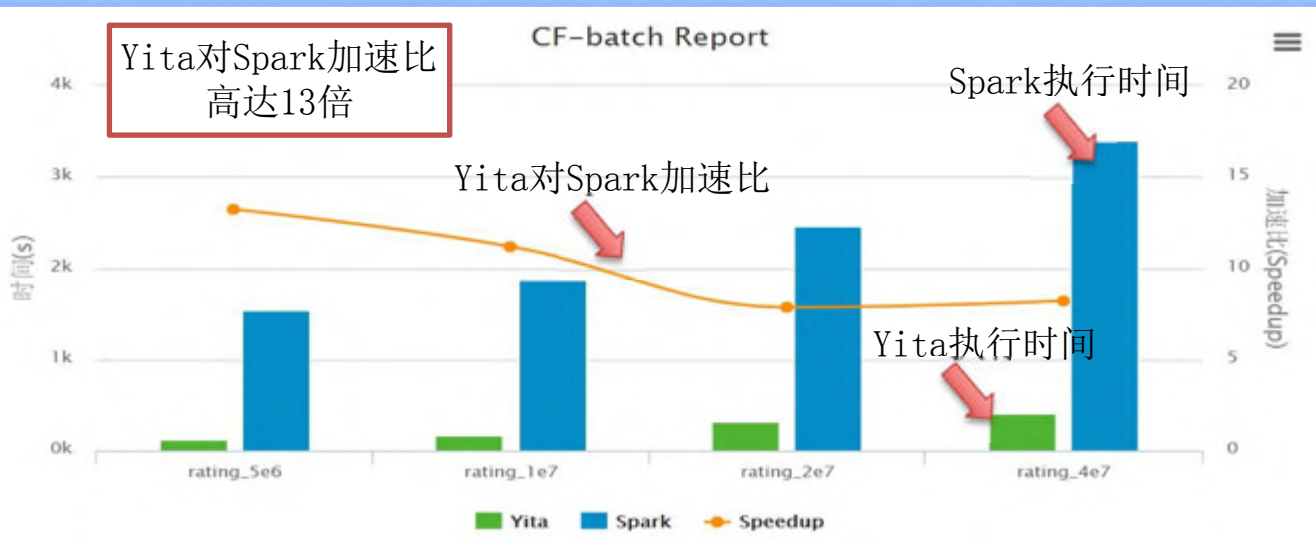
通用系统



专用系统



✓ 解决大场景



稀疏的大数

算法支持
动态调度

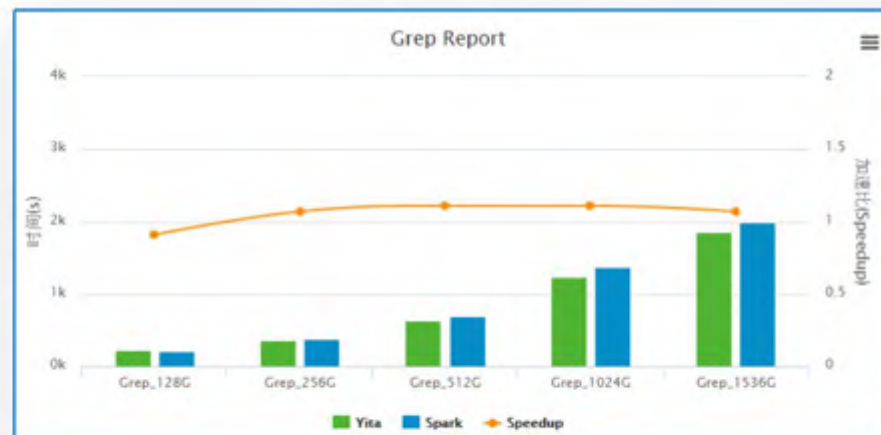
学习



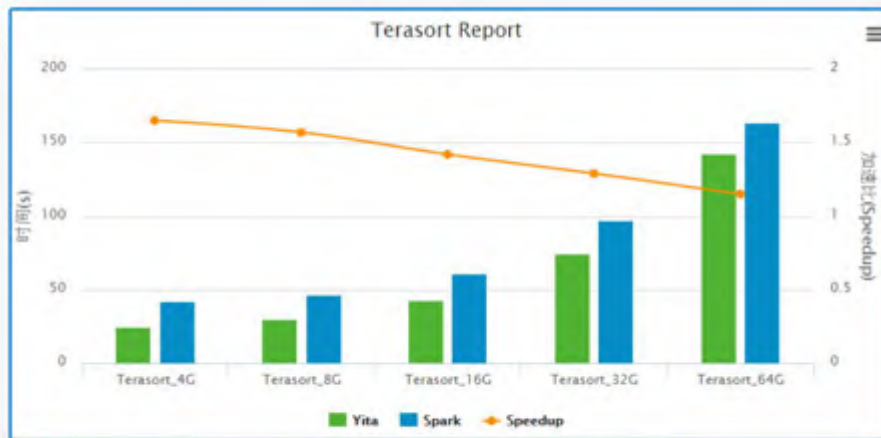
Yita在IO密集型上的性能的表现



WordCount 数据量增加



Grep 数据量增加



TeraSort 数据量增加

- 随着数据量的增大，IO操作比重逐渐增大，Yita针对计算效果的优化，反映在整体计算性能加速上表现甚微。



清华大学



北京大学



中科院计算所



华中科技大学



上海交通大学



中国地质大学



美国特拉华大学



IEEE STC
Parallel Execution
Model and System:
Dataflow and
Beyond

共同从事和推进数据
流-大数据技术研究
与产业化工作

商业案例

- 某省级电信运营商视频监控分析业务
- 某省级电信运营商海量流式数据实时分析
- 某省级电信运营商快速用户画像业务
- 某金融信息公司智能投顾业务
- 某省高速公路实时车辆视频分析系统
- 某市公安实时智能视频巡逻系统

- 某大学校级大数据处理平台

电信

金融

政府

教育

Yita版本状态

- Yita V1.0 : 8月15日发布

- Yita的内核稳定；Yita基本容错支持；
- 对YARN的全面支持，对其他生态系统的兼容和支持；
- 形成Yita图形化管理组件，YiBench及其图形化套件；
- 形成对比Spark、Storm在批处理、流处理高达10余倍的性能优势。

- Yita V1.1 : 即将发布

- 发布Yita机器学习算法库（第一阶段）；
- 逻辑回归、随机森林、协同过滤、Lasso、话题模型；
- 支持基于参数服务器抽象的异步细粒度机器学习算法实现；
- 对高维稀疏、大规模模型参数等场景下的机器学习性能表现优于Spark系统2~10倍。

中兴飞流展位 三层B08



欢迎各位莅临
中兴飞流展台

谢谢！ Q & A

郑龙

zheng.long@zte.com.cn

136-5186-6290

