



Challenges of training and deploying foundation models at scale

Maxime Hugues Ph.D.
Principal Applied Scientist GenAI

2025-06-16

Generative AI has potential business value



NEW EXPERIENCES

Create new innovative and engaging ways of interacting with your customers and employees



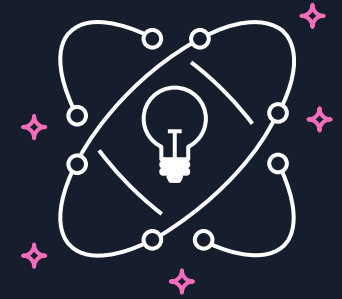
PRODUCTIVITY

Radically improve productivity across all lines of business



INSIGHTS

Extract insights and clear answers from all your corporate information, enabling faster and better decisions



CREATIVITY

Create new content and ideas, including conversations, stories, images, videos, and music

Generative AI Application Examples



Travel & Hospitality

Chatbot / assistants / conversational AI, question answering, search



Financial services

Risk management, fraud detection, summarize annual reports, review portfolio / risks



Healthcare

Protein folding, drug development, personalized medicine, improved medical imaging



Consumer goods

Optimize pricing and inventory, correctly flag product brand and category



Media and entertainment

Video game generation, upscaling content, face synthesis, film preservation and coloring



Energy and utilities

Design renewable energy sources optimized for geo, predictive maintenance



Automotive

Autonomous vehicles, design parts for fuel efficiency



Technology hardware

Chip design, robotics

Different GenAI player

Model Builder



Startup, Research
Institute, Enterprise

Model Tuner



Integrated Software
Vendors, Startups, GSI

Model Consumer



Everyone

Most have inexistent or limited access to accelerators

Persona training, deploying and using model

Machine Learning Engineer



Engineer



Scientist



Data scientist



Analyst



Accelerated computing for ML

LARGEST SELECTION OF ML INSTANCES IN THE CLOUD

MACHINE LEARNING

TRAINING

P4de

NVIDIA A100
EFA v1

P5

NVIDIA H100
EFA v2

P5e

NVIDIA H200
EFA v2

P5en

NVIDIA H200
EFA v3

P6-B200

NVIDIA B200
EFA v4

TRN1

AWS TRAINIUM
EFA v2

TRN2

AWS TRAINIUM2
EFA v3

INFERENCE

G5

NVIDIA A10G

G6

NVIDIA L4

G6e

NVIDIA L40S

INF1

AWS
INFERENCE

INF2

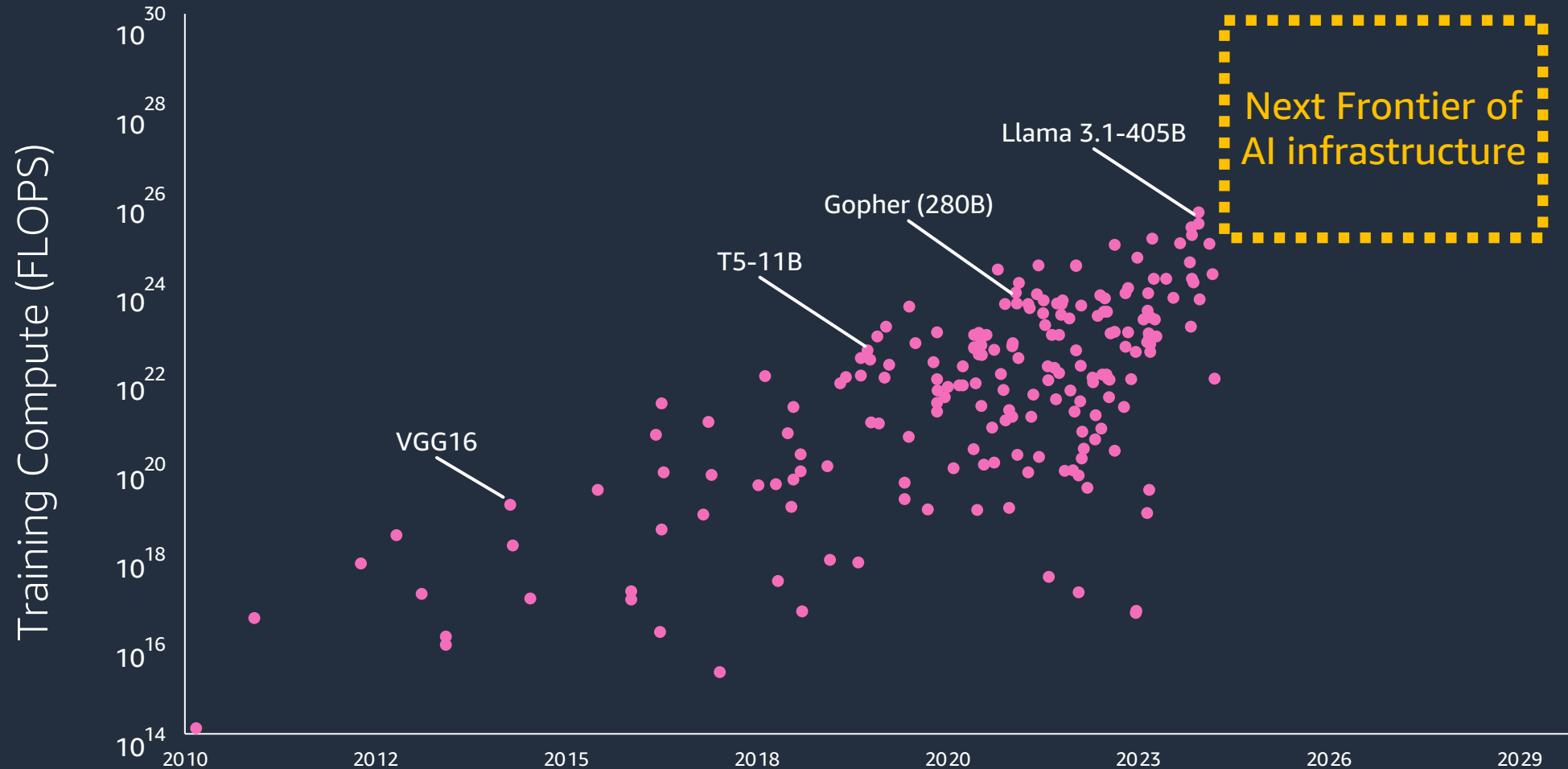
AWS
INFERENCE2



Model Training Challenges



Scaling compute requires next-generation AI infrastructure



Training at scale

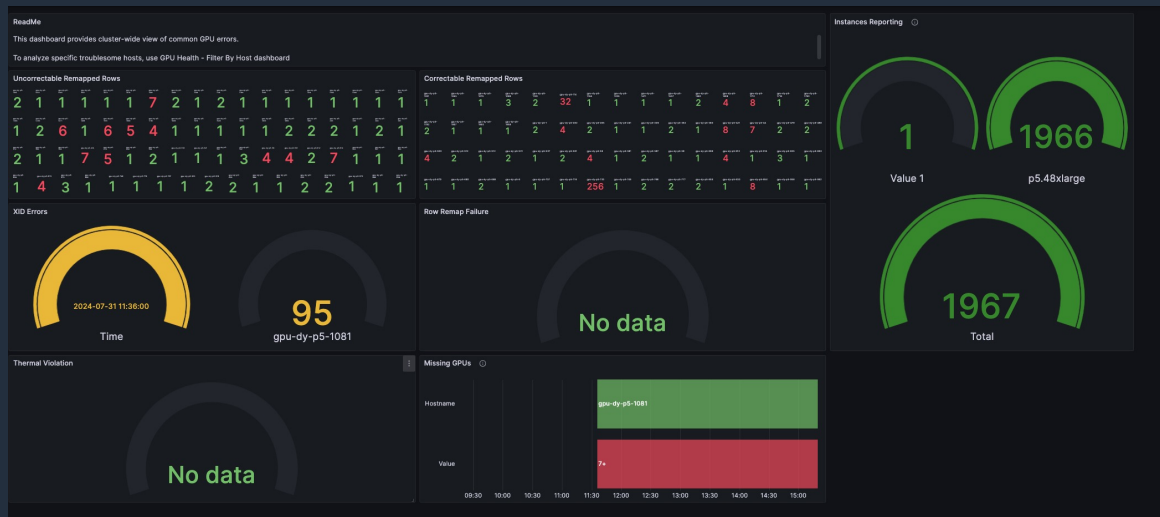
Llama 3 405B was trained during 54 days on 16,000 H100 GPUs on a multilingual corpus of 15T tokens corpus

Interrupted every ~2.8h
GPUs accounted for 58.7% of failure

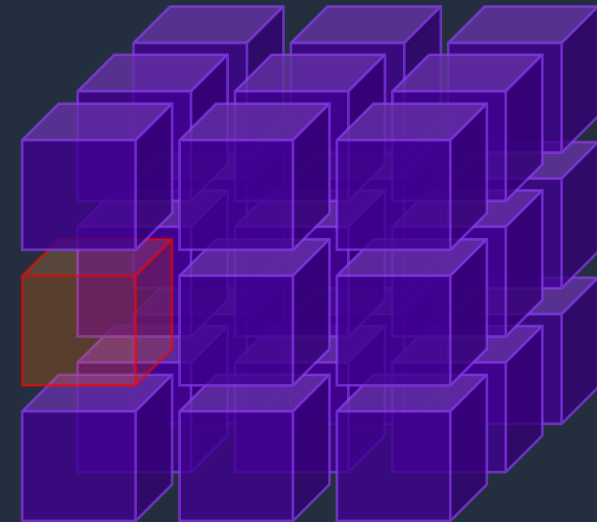


Infrastructure resilience

Observability



Fault detection, hardware replacement



SageMaker Hyperpod
EKS NodeRepair

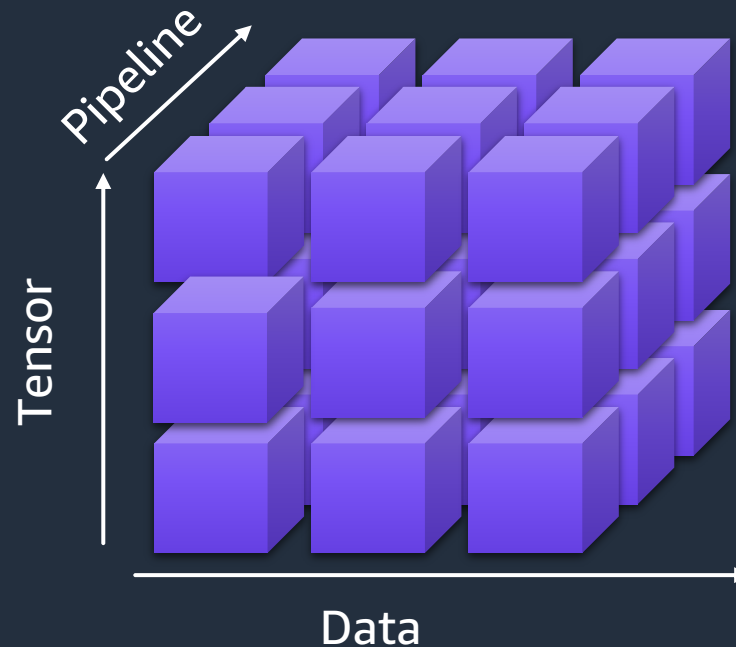
Customer focus on training and business value

Multi-knobs performance

Parallelism

NCCL Parameters

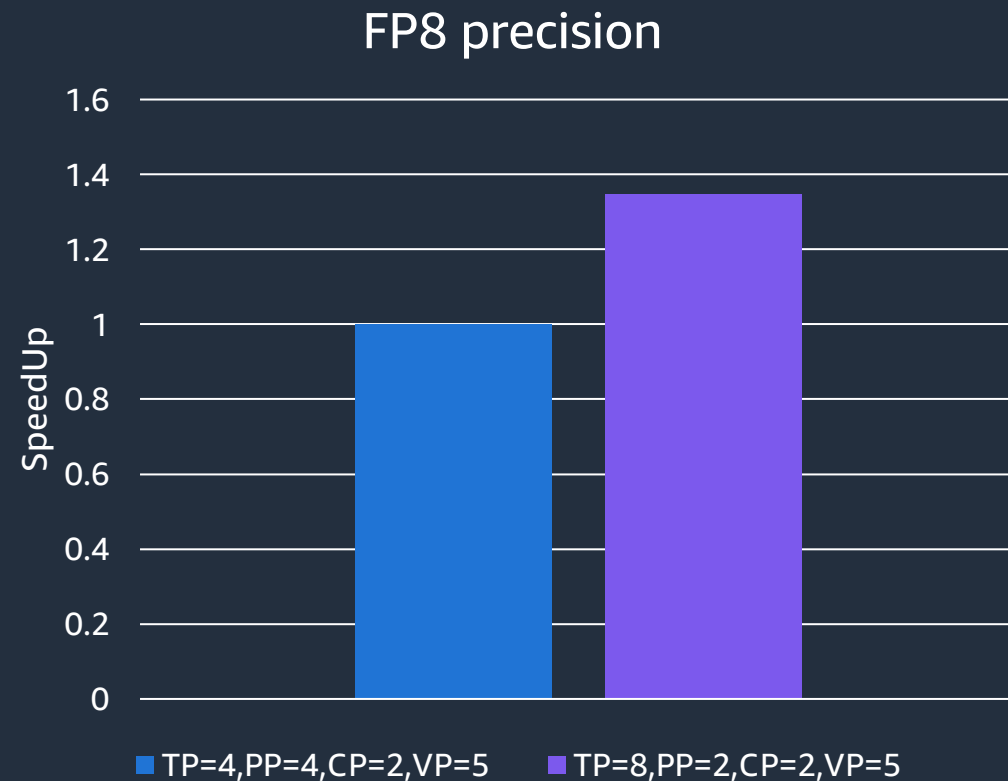
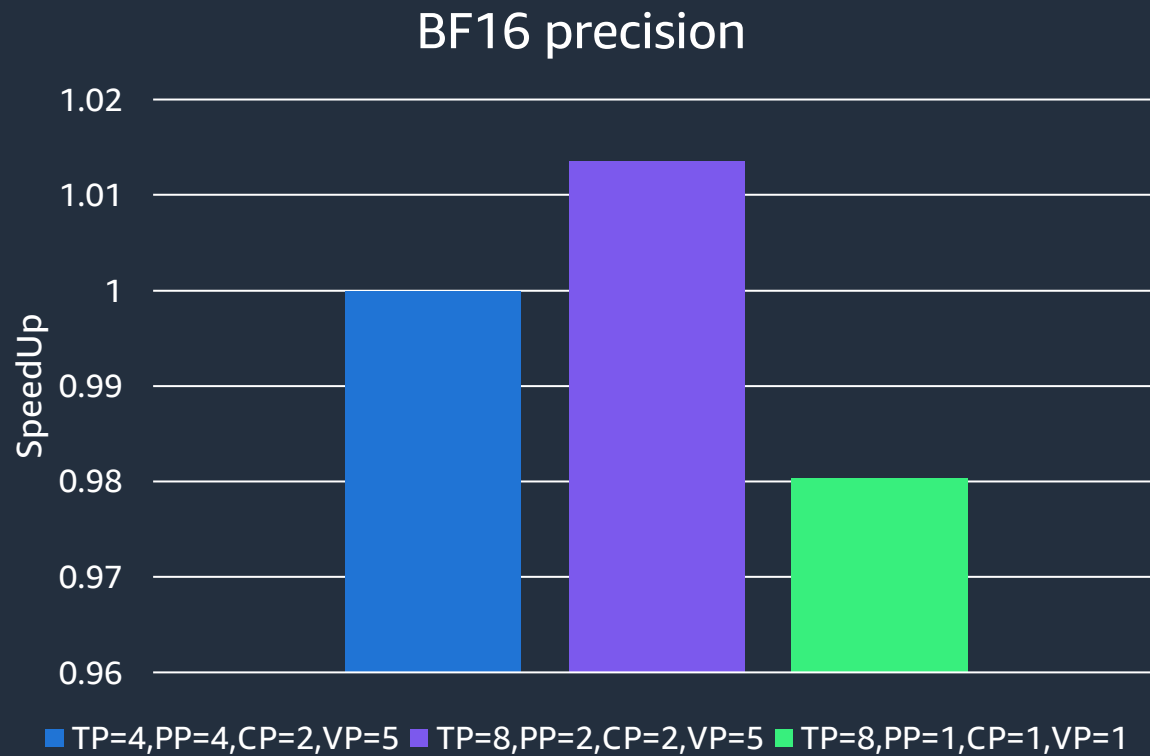
NCCL_COLLNET
NCCL_BUFFSIZE
NCCL_P2P
MCCL_MNVLS



Precision

BF16
FP8+BF16
MoE FP4 + BF16

Parallelism impact on training performance



Llama 3.1 70B on 128x H100

LLama3.1 70B Profile

CUDA HW (0000:53:00.0 - NVIDIA H100 80GB HBM3)

[All Streams]

▼ 96.4% Kernels

▶ 38.8% ncclDevKernel_SendRecv

▶ 15.4% sm90_xmma_gemm_bf16f32_bf16f32_f32_nt_n_tilsize128x128x64_warpgroupsize1x1x1_execu

▶ 8.3% sm90_xmma_gemm_bf16bf16_bf16f32_f32_tn_n_tilsize128x128x64_warpgroupsize1x1x1_execu

▶ 6.5% sm90_xmma_gemm_bf16bf16_bf16f32_f32_nn_n_tilsize256x128x64_warpgroupsize2x1x1_execu

▶ 4.6% sm90_xmma_gemm_bf16bf16_bf16f32_f32_nn_n_tilsize128x128x64_warpgroupsize1x1x1_execu

▶ 4.0% sm90_xmma_gemm_bf16bf16_bf16f32_f32_tn_n_tilsize128x128x64_warpgroupsize1x1x1_execu

▶ 3.9% flash_bwd_dq_dk_dv_loop_seqk_parallel_kernel

▶ 3.1% userbuffers_fp16_sum_inplace_gpu_mc_rs_oop

▶ 2.4% userbuffers_fp16_sum_inplace_gpu_mc_rs

▶ 2.3% userbuffers_fp16_sum_inplace_gpu_mc_ag

▶ 2.1% sm90_xmma_gemm_bf16bf16_bf16f32_f32_tn_n_tilsize128x256x64_warpgroupsize2x1x1_execu

▶ 1.6% flash_fwd_kernel

▶ 1.3% sm90_xmma_gemm_bf16bf16_bf16f32_f32_nn_n_tilsize128x256x64_warpgroupsize2x1x1_execu

Send/Recv dominant

Model Serving



Model serving type

Offline



Batch processing
Image
Text
Database

Online



ChatBot
Assistant
Medical Imaging
Image creation

Serving model closer to end-users



600+

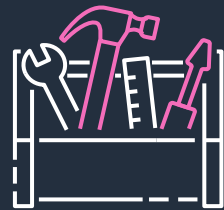
Amazon
CloudFront
Points of
Presence



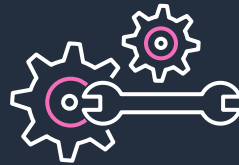
Build and Scale AI with Amazon Bedrock



Accelerate development of generative AI applications using FMs through an API, without managing infrastructure



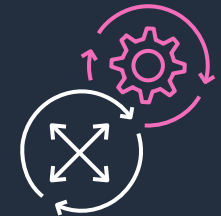
Choose FMs from AI21 Labs, Anthropic, Stability AI, and Amazon to find the right FM for your use case



Privately customize FMs using your organization's data

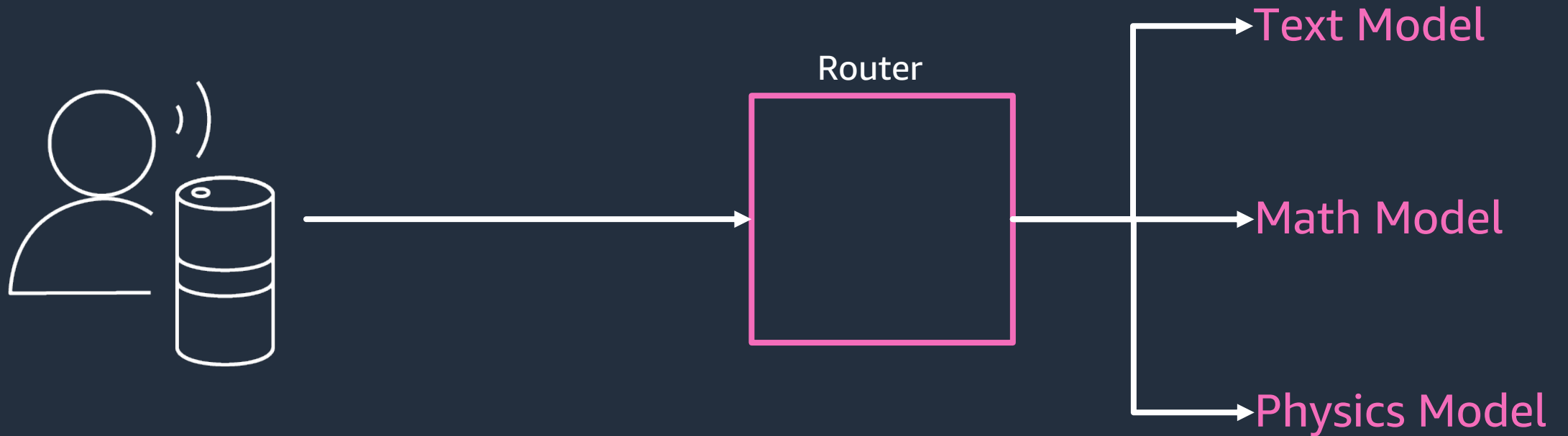


Enhance your data protection using comprehensive AWS security capabilities



Use AWS tools and capabilities that you are familiar with to deploy scalable, reliable, and secure generative AI applications

In the wait of AGI

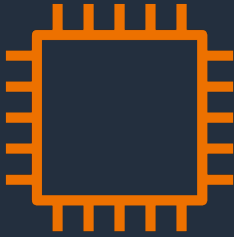


What's next



Ease of use for performance and results

Abstraction from infrastructure
is reality



AI optimizing AI ?



Tuning Infrastructure

Tuning parallelism

HPC Agents



Thank you!