

Murakkab: Resource-Efficient Agentic Workflow Orchestration in Cloud Platforms

Gohar Irfan Chaudhry, Esha Choukse, Haoran Qiu, Íñigo Goiri,
Rodrigo Fonseca, Adam Belay, Ricardo Bianchini



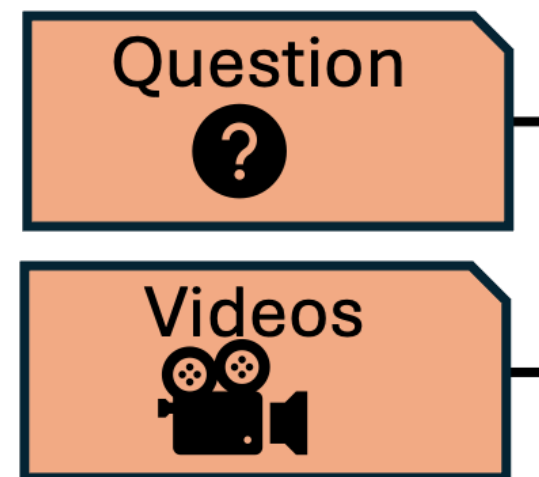
Agentic workflows are ubiquitous...



In this talk, we will focus on the
resource-efficiency of agentic workflows

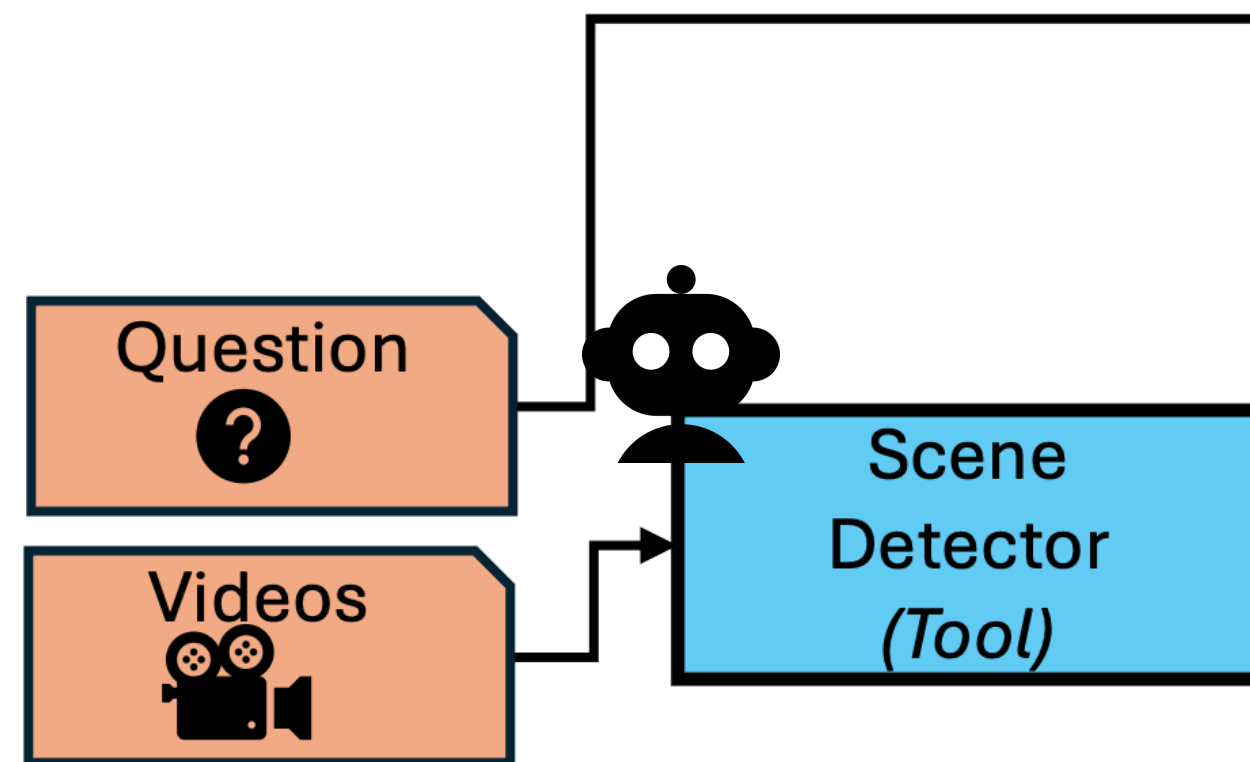
Example: Video Q/A workflow

Given a set of videos, answer the queries



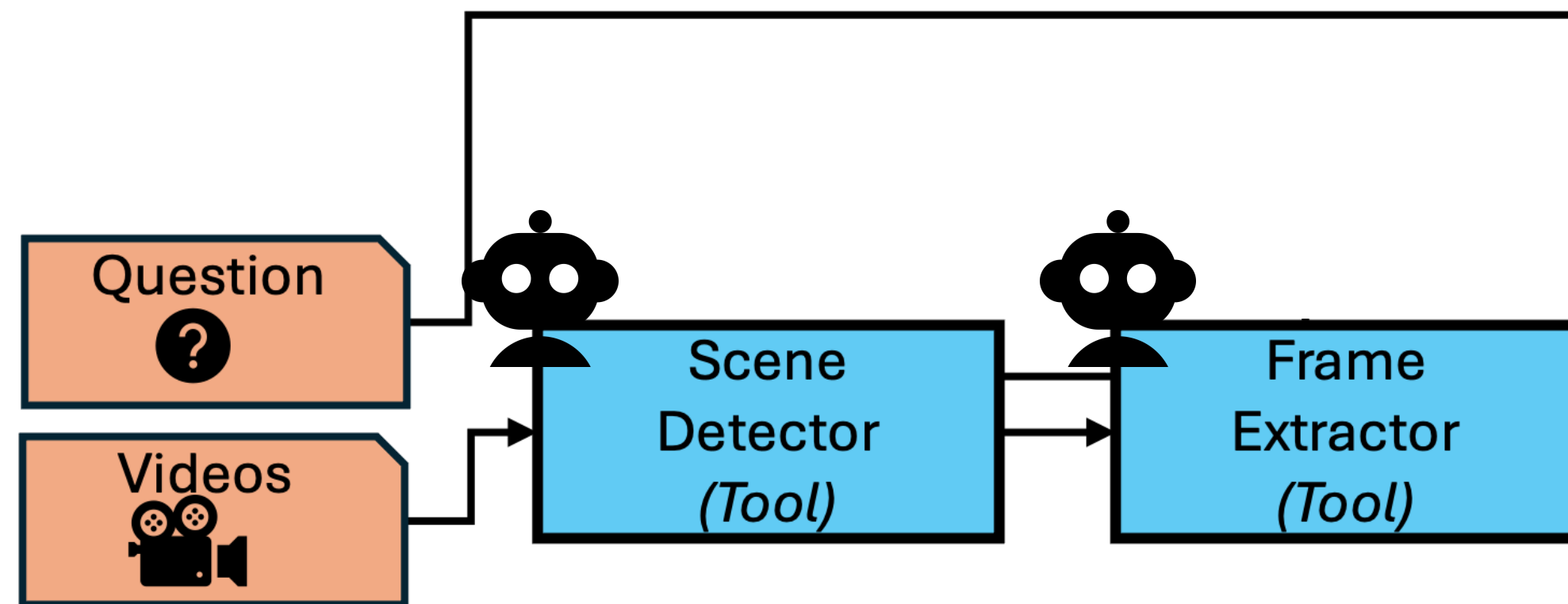
Example: Video Q/A workflow

Given a set of videos, answer the queries



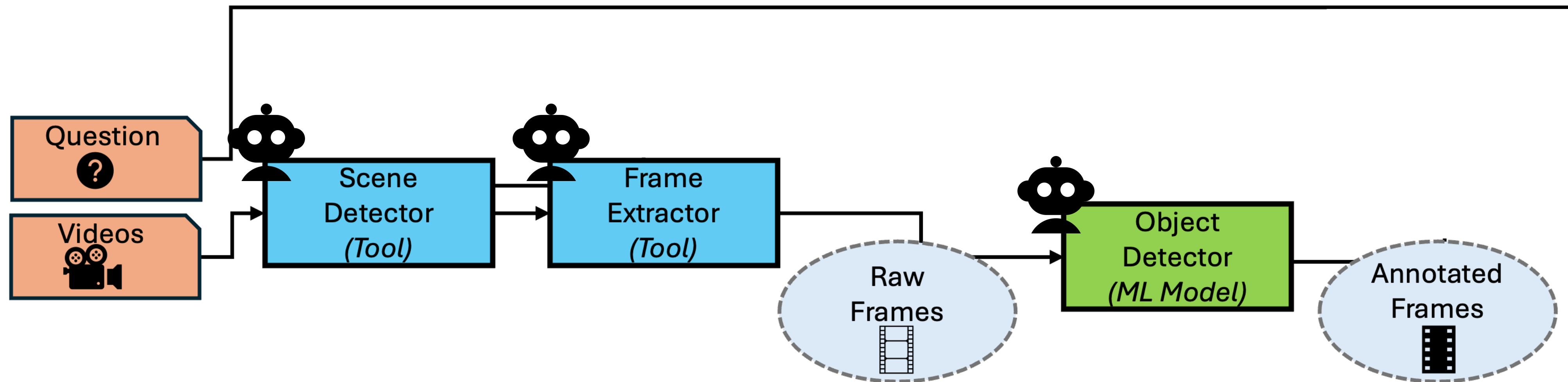
Example: Video Q/A workflow

Given a set of videos, answer the queries



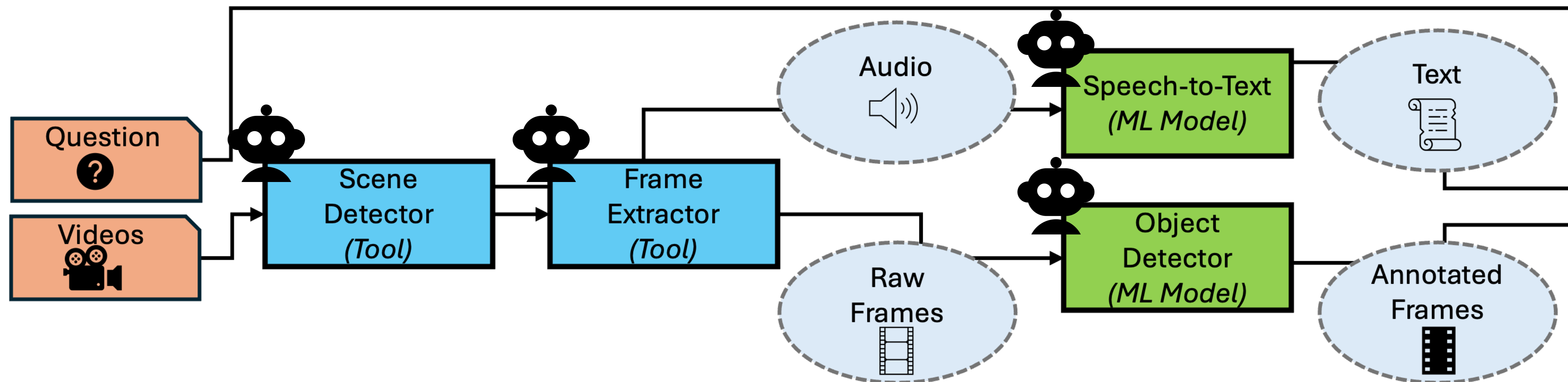
Example: Video Q/A workflow

Given a set of videos, answer the queries



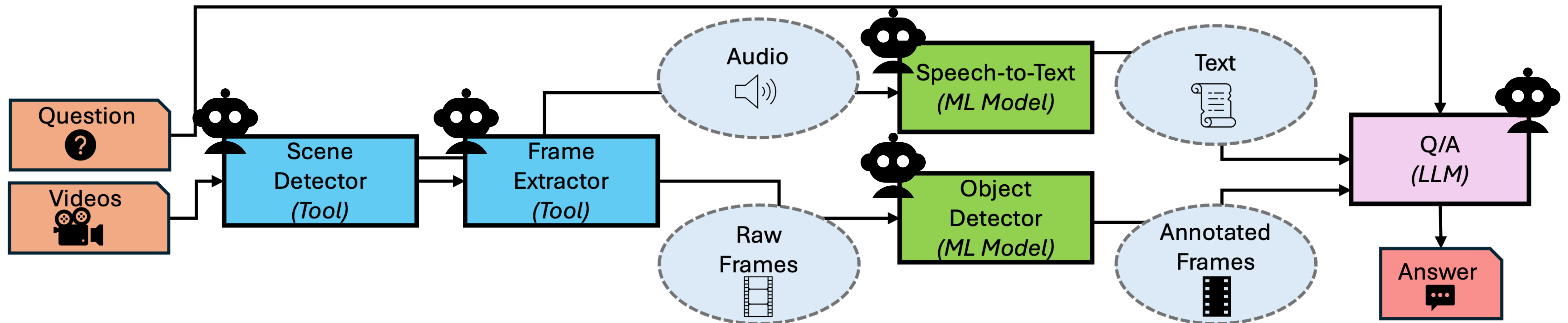
Example: Video Q/A workflow

Given a set of videos, answer the queries



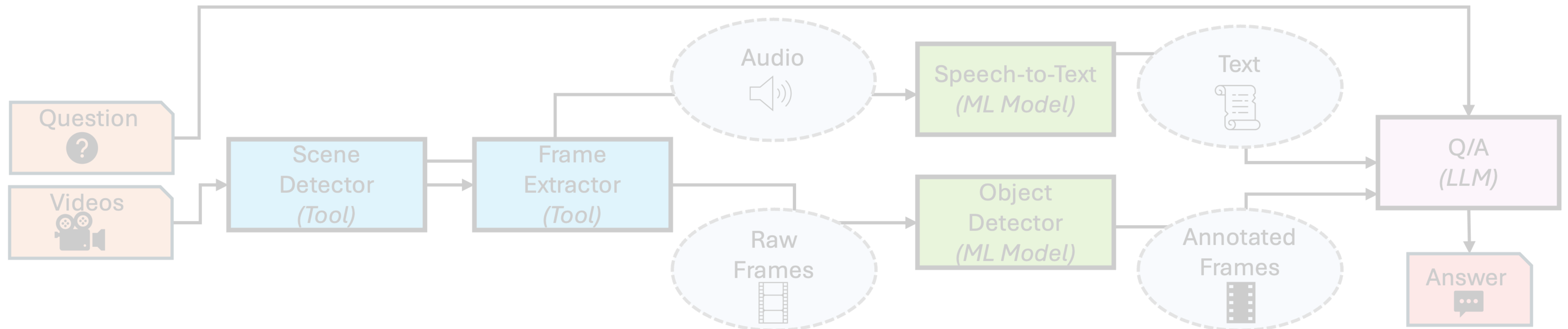
Example: Video Q/A workflow

Given a set of videos, answer the queries



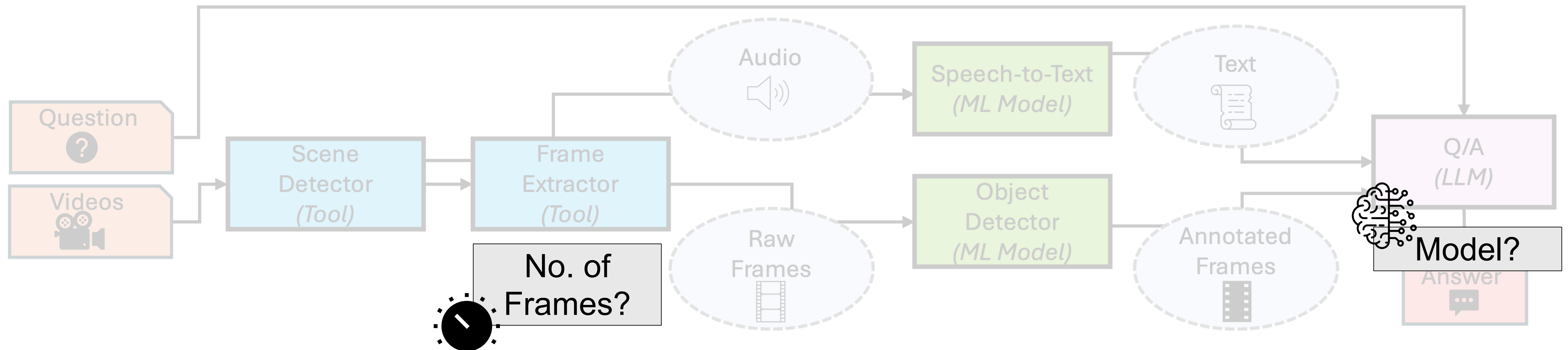
Example: Video Q/A workflow

Agent-, workflow-, hardware-level configurations



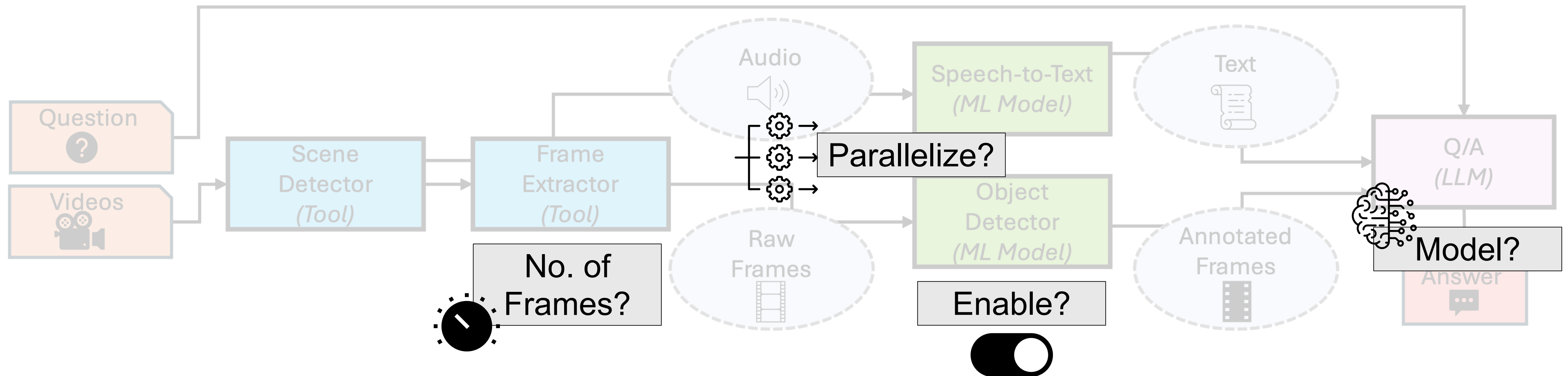
Example: Video Q/A workflow

Agent-, workflow-, and hardware-level configurations



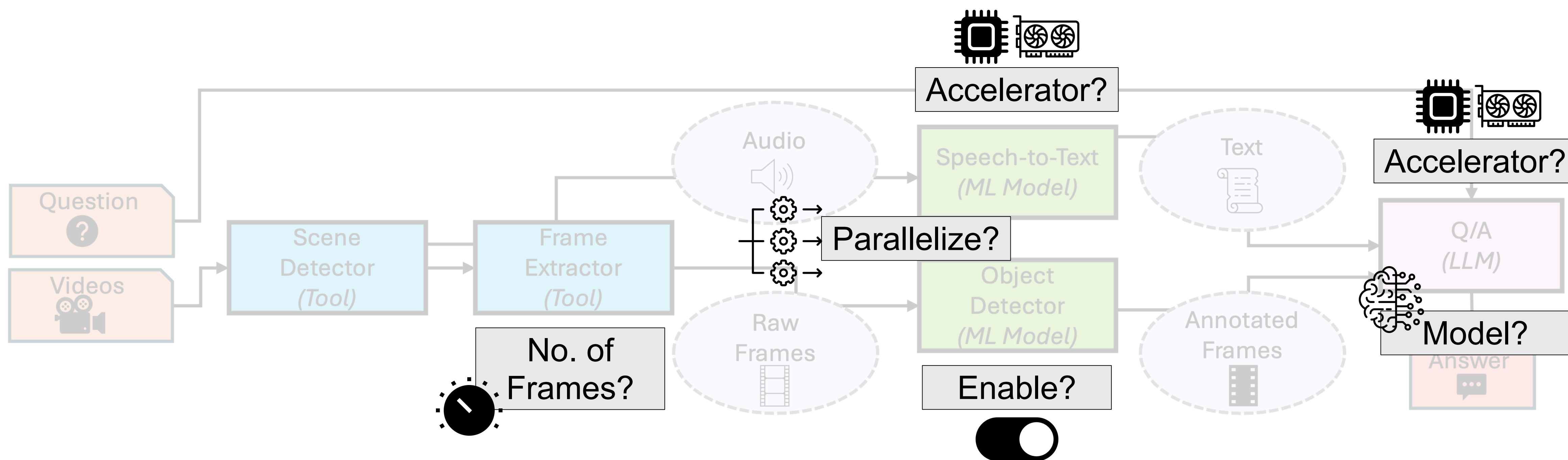
Agent-, workflow-, and hardware-level configurations

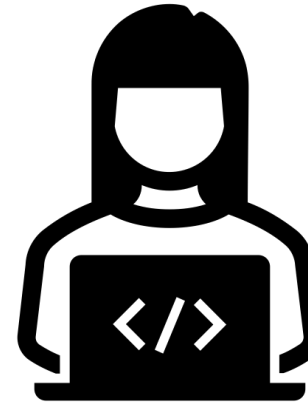
Agent-, workflow-, and hardware-level configurations



Example: Video Q/A workflow

Agent-, workflow-, and hardware-level configurations

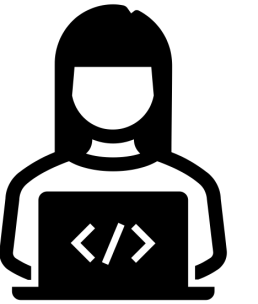




App. Developer

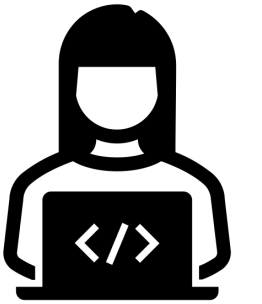
How is a workflow defined today?

Example: Video Q/A workflow definition



```
1  # == Workflow nodes (coupled app logic + configs) ==
2  scene_detection = Tool(
3      fn=SceneDetector(),      # Implemented earlier by the
4                               # developer.
5      key=AWS_KEY, resources={"CPUs": 32}
6  )
7  frame_extractor = Tool(
8      fn=FrameExtractor(),    # Implemented earlier by the
9                               # developer.
10     key=AWS_KEY, params={"num_frames": 15}, resources={"CPUs": 32}
11 )
12 speech_to_text = MLModel(
13     name="Whisper", key=OPENAI_API_KEY, resources={"PTUs": 50}
14 )
15 object_detection = MLModel(
16     name="CLIP", key=AZURE_KEY, resources={"CPUs": 128}
17 )
18 question_answer = LLM(
19     name="Llama-3.2", key=DATABRICKS_API_KEY,
20     params={"batch": 256}, resources={"GPUs": 8, "Type": "H100"},
21     system_prompt="You are an agent that can understand videos.",
22     user_prompt="Answer the given question about the video."
23 )
24 # == Query and data ==
25 ques = "What is the name of the person wearing the red dress?"
26 videos = ["road_trip.mp4"]
27 # == Workflow (ordering of nodes and data flow) ==
28 scenes, audio = scene_detection(videos)
29 frames = frame_extractor(scenes)
30 transcript = speech_to_text(audio)
31 obj_frames = object_detection(frames)
32 answer = question_answer(ques, transcript, obj_frames)
```

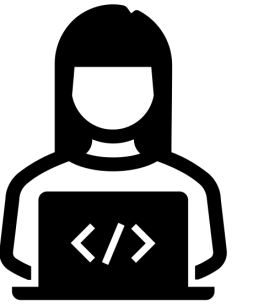
Example: Video Q/A workflow definition



Application Logic

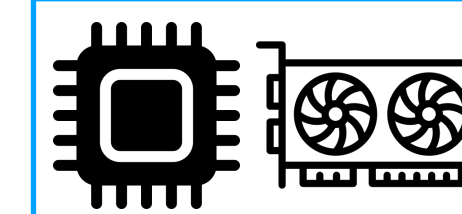
```
1  # == Workflow nodes (coupled app logic + configs) ==
2  scene_detection = Tool(
3      fn=SceneDetector(),      # Implemented earlier by the
4                               # developer.
5      key=AWS_KEY, resources={"CPUs": 32}
6  )
7  frame_extractor = Tool(
8      fn=FrameExtractor(),     # Implemented earlier by the
9                               # developer.
10     key=AWS_KEY, params={"num_frames": 15}, resources={"CPUs": 32}
11 )
12 speech_to_text = MLModel(
13     name="Whisper", key=OPENAI_API_KEY, resources={"PTUs": 50}
14 )
15 object_detection = MLModel(
16     name="CLIP", key=AZURE_KEY, resources={"CPUs": 128}
17 )
18 question_answer = LLM(
19     name="Llama-3.2", key=DATABRICKS_API_KEY,
20     params={"batch": 256}, resources={"GPUs": 8, "Type": "H100"},
21     system_prompt="You are an agent that can understand videos.",
22     user_prompt="Answer the given question about the video."
23 )
24 # == Query and data ==
25 ques    = "What is the name of the person wearing the red dress?"
26 videos  = ["road_trip.mp4"]
27 # == Workflow (ordering of nodes and data flow) ==
28 scenes, audio = scene_detection(videos)
29 frames        = frame_extractor(scenes)
30 transcript     = speech_to_text(audio)
31 obj_frames     = object_detection(frames)
32 answer        = question_answer(ques, transcript, obj_frames)
```

Example: Video Q/A workflow definition



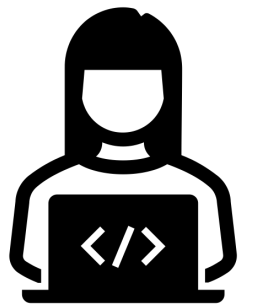
```
1  # == Workflow nodes (coupled app logic + configs) ==
2  scene_detection = Tool(
3      fn=SceneDetector(),      # Implemented earlier by the
4      key=AWS_KEY, resources={"CPUs": 32}
5  )
6  frame_extractor = Tool(
7      fn=FrameExtractor(),    # Implemented earlier by the
8      key=AWS_KEY, params={"num_frames": 15}, resources={"CPUs": 32}
9  )
10 speech_to_text = MLModel(
11     name="Whisper", key=OPENAI_API_KEY, resources={"PTUs": 50}
12 )
13 object_detection = MLModel(
14     name="CLIP", key=AZURE_KEY, resources={"CPUs": 128}
15 )
16 question_answer = LLM(
17     name="Llama-3.2", key=DATABRICKS_API_KEY,
18     params={"batch": 256}, resources={"GPUs": 8, "Type": "H100"},
19     system_prompt="You are an agent that can understand videos.",
20     user_prompt="Answer the given question about the video."
21 )
22 # == Query and data ==
23 ques    = "What is the name of the person wearing the red dress?"
24 videos  = ["road_trip.mp4"]
25 # == Workflow (ordering of nodes and data flow) ==
26 scenes, audio = scene_detection(videos)
27 frames        = frame_extractor(scenes)
28 transcript     = speech_to_text(audio)
29 obj_frames     = object_detection(frames)
30 answer        = question_answer(ques, transcript, obj_frames)
```

Application Logic



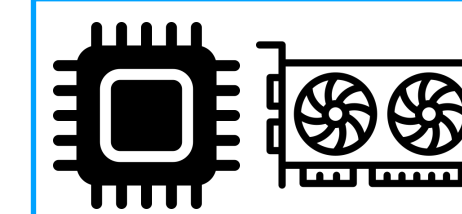
Hardware
Config.

Example: Video Q/A workflow definition

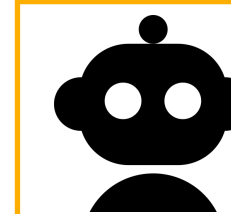


```
1 # == Workflow nodes (coupled app logic + configs) ==
2 scene_detection = Tool(
3     fn=SceneDetector(), # Implemented earlier by the
4     key=AWS_KEY, resources={"CPUs": 32}
5 )
6 frame_extractor = Tool(
7     fn=FrameExtractor(), # Implemented earlier by the
8     key=AWS_KEY, params={"num_frames": 15}, resources={"CPUs": 32}
9 )
10 speech_to_text = MLModel(
11     name="Whisper", key=OPENAI_API_KEY, resources={"PTUs": 50}
12 )
13 object_detection = MLModel(
14     name="CLIP", key=AZURE_KEY, resources={"CPUs": 128}
15 )
16 question_answer = LLM(
17     name="Llama-3.2", key=DATABRICKS_API_KEY,
18     params={"batch": 256}, resources={"GPUs": 8, "Type": "H100"},
19     system_prompt="You are an agent that can understand videos.",
20     user_prompt="Answer the given question about the video."
21 )
22 # == Query and data ==
23 ques = "What is the name of the person wearing the red dress?"
24 videos = ["road_trip.mp4"]
25 # == Workflow (ordering of nodes and data flow) ==
26 scenes, audio = scene_detection(videos)
27 frames = frame_extractor(scenes)
28 transcript = speech_to_text(audio)
29 obj_frames = object_detection(frames)
30 answer = question_answer(ques, transcript, obj_frames)
```

Application Logic

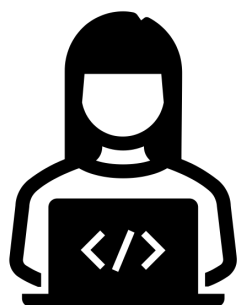


Hardware
Config.



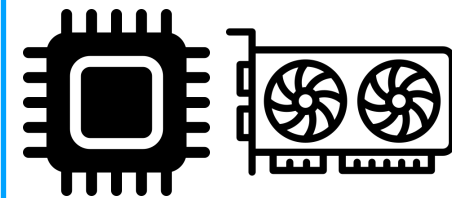
Agent
Config.

Example: Video Q/A workflow definition

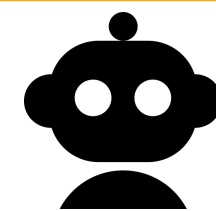


```
1 # == Workflow nodes (coupled app logic + configs) ==
2 scene_detection = Tool(
3     fn=SceneDetector(), # Implemented earlier by the
4     developer.
5     key=AWS_KEY, resources={"CPUs": 32}
6 )
7 frame_extractor = Tool(
8     fn=FrameExtractor(), # Implemented earlier by the
9     developer.
10    key=AWS_KEY, params={"num_frames": 15}, resources={"CPUs": 32}
11 )
12 speech_to_text = MLModel(
13     name="Whisper", key=OPENAI_API_KEY, resources={"PTUs": 50}
14 )
15 object_detection = MLModel(
16     name="CLIP", key=AZURE_KEY, resources={"CPUs": 128}
17 )
18 question_answer = LLM(
19     name="Llama-3.2", key=DATABRICKS_API_KEY,
20     params={"batch": 256}, resources={"GPUs": 8, "Type": "H100"},
21     system_prompt="You are an agent that can understand videos.",
22     user_prompt="Answer the given question about the video."
23 )
24 # == Query and data ==
25 ques = "What is the name of the person wearing the red dress?"
26 videos = ["road_trip.mp4"]
27 # == Workflow (ordering of nodes and data flow) ==
28 scenes, audio = scene_detection(videos)
29 frames = frame_extractor(scenes)
30 transcript = speech_to_text(audio)
31 obj_frames = object_detection(frames)
32 answer = question_answer(ques, transcript, obj_frames)
```

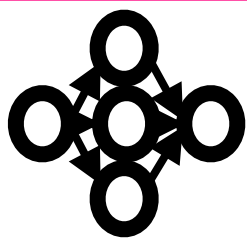
Application Logic



Hardware
Config.

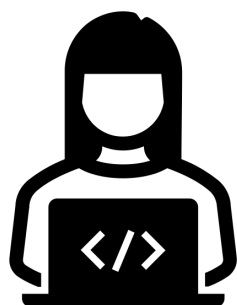


Agent
Config.



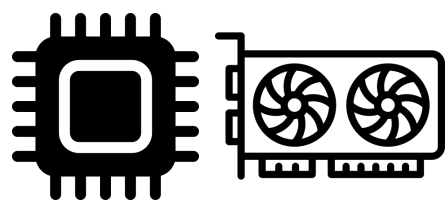
Workflow
Config.

Example: Video Q/A workflow definition

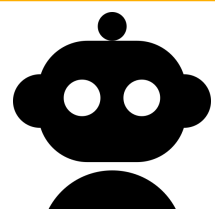


```
1 # == Workflow nodes (coupled app logic + configs) ==
2 scene_detection = Tool(
3     fn=SceneDetector(), # Implemented earlier by the
4     developer.
5     key=AWS_KEY, resources={"CPUs": 32}
6 )
7 frame_extractor = Tool(
8     fn=FrameExtractor(), # Implemented earlier by the
9     developer.
10    key=AWS_KEY, params={"num_frames": 15}, resources={"CPUs": 32}
11 )
12 speech_to_text = MLModel(
13     name="Whisper", key=OPENAI_API_KEY, resources={"PTUs": 50}
14 )
15 object_detection = MLModel(
16     name="CLIP", key=AZURE_KEY, resources={"CPUs": 128}
17 )
18 question_answer = LLM(
19     name="Llama-3.2", key=DATABRICKS_API_KEY,
20     params={"batch": 256}, resources={"GPUs": 8, "Type": "H100"},
21     system_prompt="You are an agent that can understand videos.",
22     user_prompt="Answer the given question about the video."
23 )
24 # == Query and data ==
25 ques = "What is the name of the person wearing the red dress?"
26 videos = ["road_trip.mp4"]
27 # == Workflow (ordering of nodes and data flow) ==
28 scenes, audio = scene_detection(videos)
29 frames = frame_extractor(scenes)
30 transcript = speech_to_text(audio)
31 obj_frames = object_detection(frames)
32 answer = question_answer(ques, transcript, obj_frames)
```

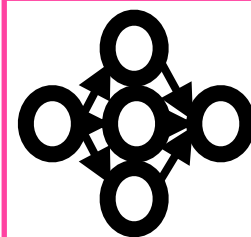
Application Logic



Hardware
Config.



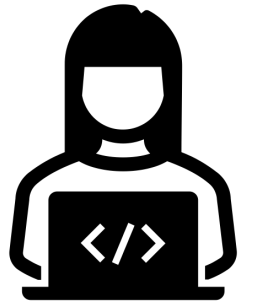
Agent
Config.



Workflow
Config.

Execution Details

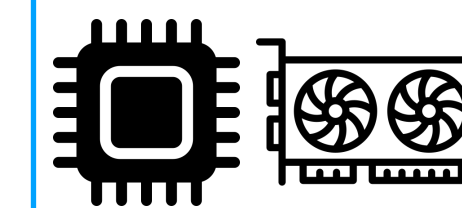
Example: Video Q/A workflow definition



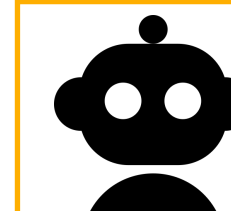
Tightly coupled application logic and execution details...

```
1 # == Workflow nodes (coupled app logic + configs) ==
2 scene_detection = Tool(
3     fn=SceneDetector(), # Implemented earlier by the
4     key=AWS_KEY, resources={"CPUs": 32}
5 )
6 frame_extractor = Tool(
7     fn=FrameExtractor(), # Implemented earlier by the
8     key=AWS_KEY, params={"num_frames": 15}, resources={"CPUs": 32}
9 )
10 speech_to_text = MLModel(
11     name="Whisper", key=OPENAI_API_KEY, resources={"PTUs": 50}
12 )
13 object_detection = MLModel(
14     name="CLIP", key=AZURE_KEY, resources={"CPUs": 128}
15 )
16 question_answer = LLM(
17     name="Llama-3.2", key=DATABRICKS_API_KEY,
18     params={"batch": 256}, resources={"GPUs": 8, "Type": "H100"},
19     system_prompt="You are an agent that can understand videos.",
20     user_prompt="Answer the given question about the video."
21 )
22 # == Query and data ==
23 ques = "What is the name of the person wearing the red dress?"
24 videos = ["road_trip.mp4"]
25 # == Workflow (ordering of nodes and data flow) ==
26 scenes, audio = scene_detection(videos)
27 frames = frame_extractor(scenes)
28 transcript = speech_to_text(audio)
29 obj_frames = object_detection(frames)
30 answer = question_answer(ques, transcript, obj_frames)
```

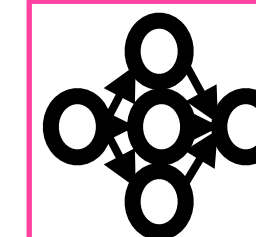
Application Logic



Hardware Config.



Agent Config.

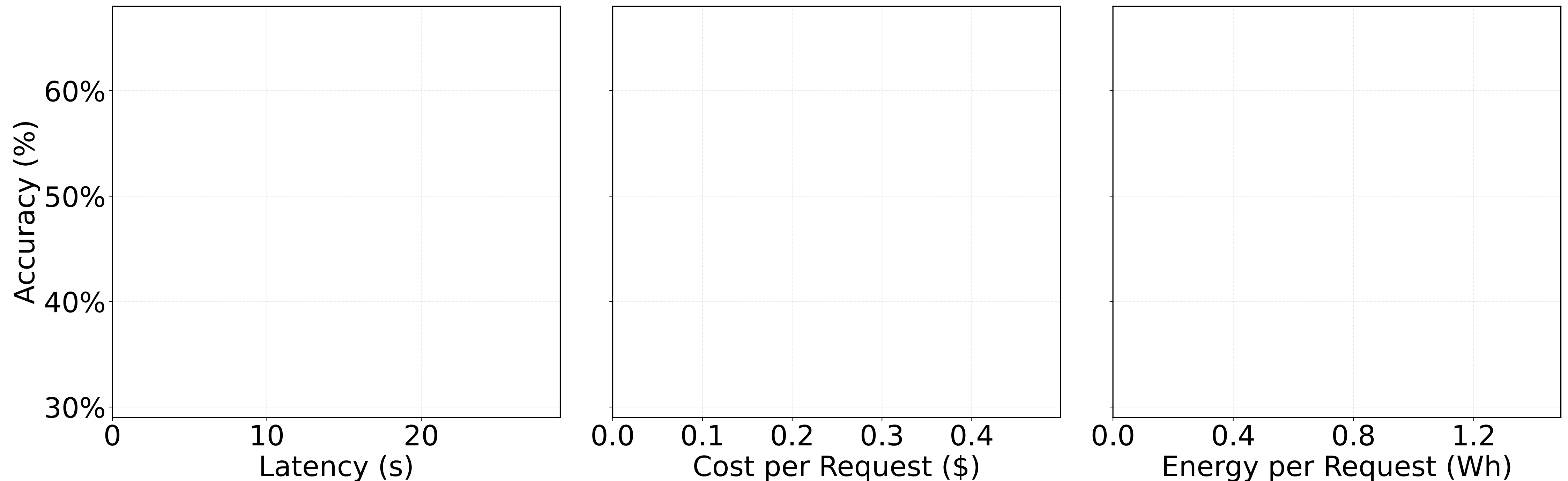


Workflow Config.

Execution Details

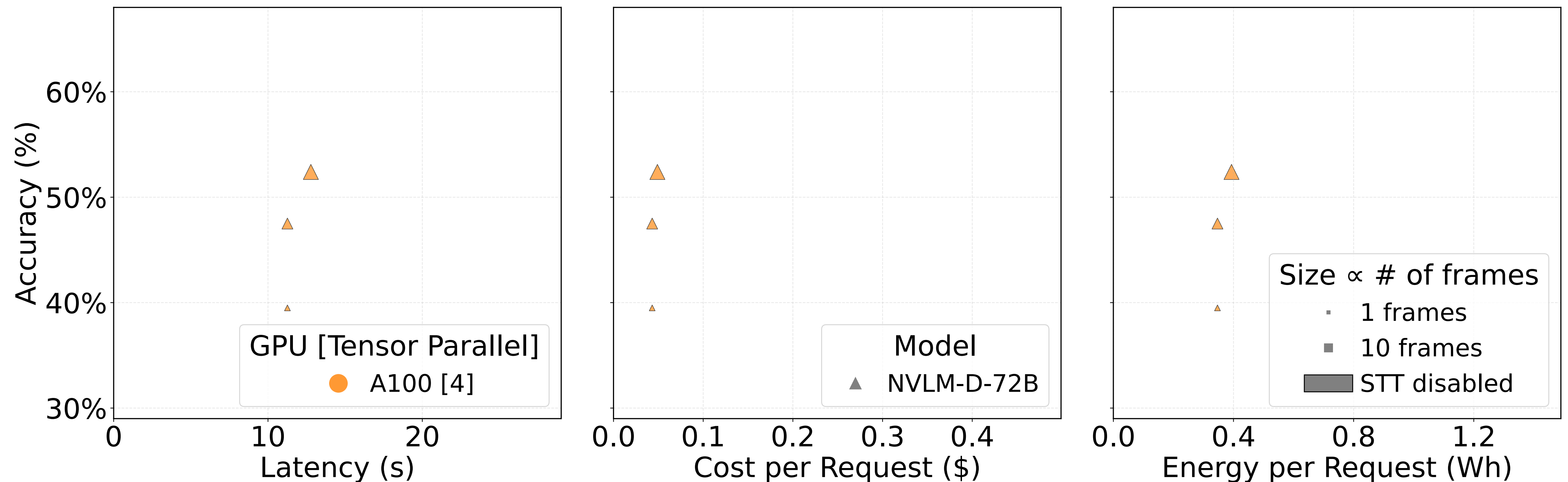
High-dimensional configuration space

Multiple objectives/SLOs



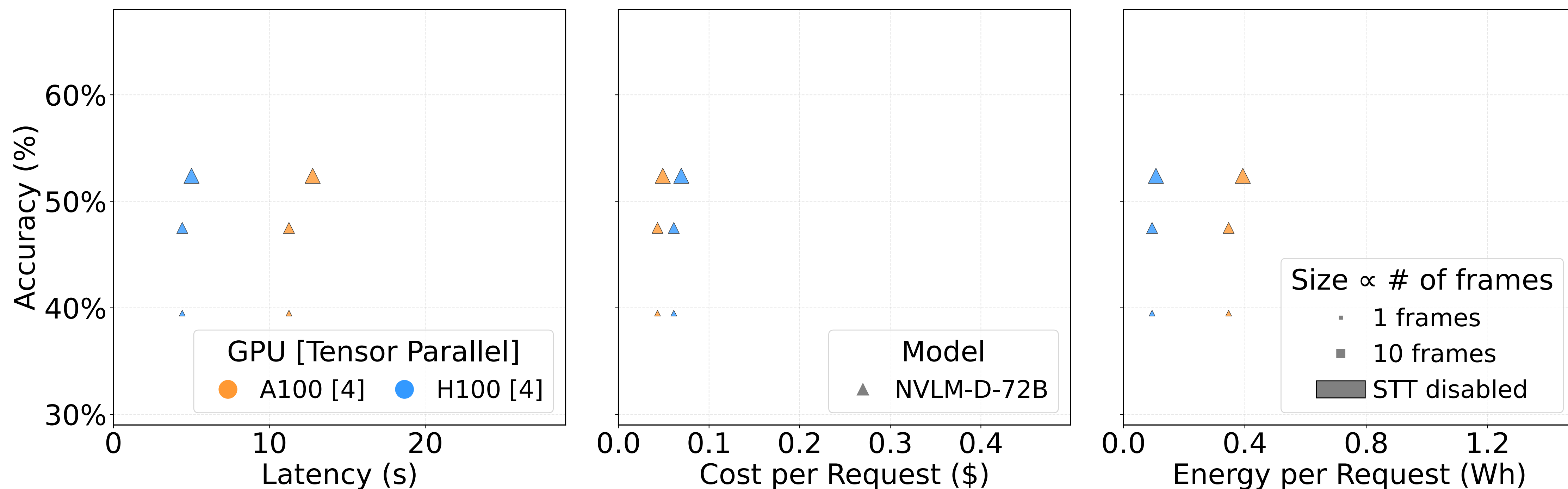
High-dimensional configuration space

Multiple objectives/SLOs



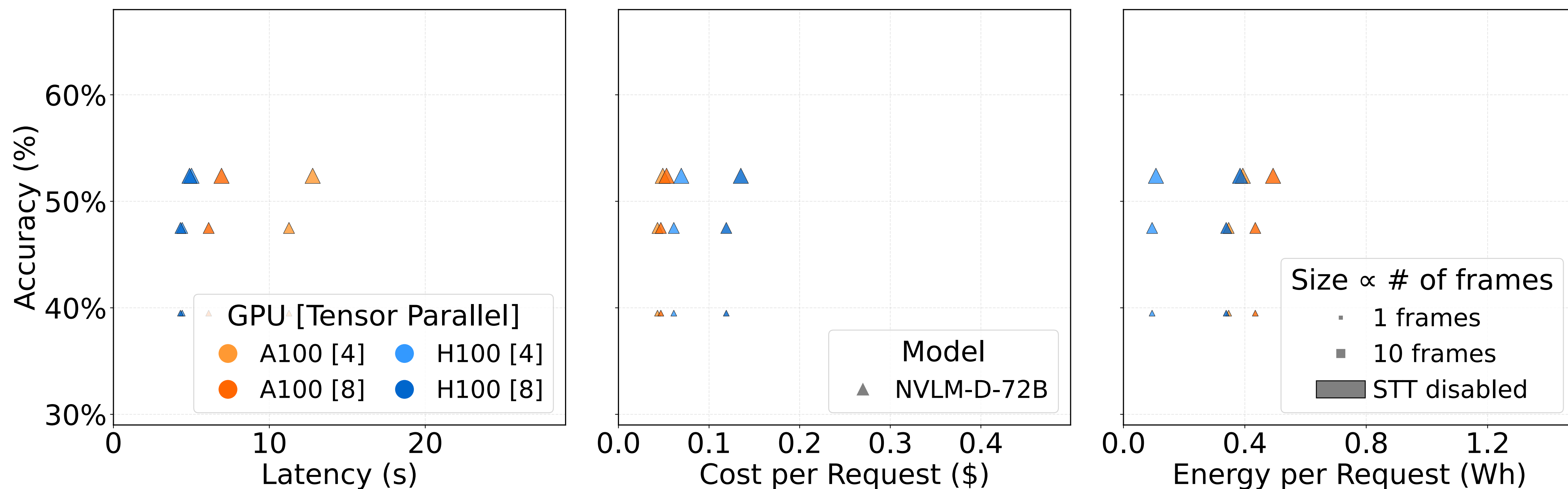
High-dimensional configuration space

Multiple objectives/SLOs



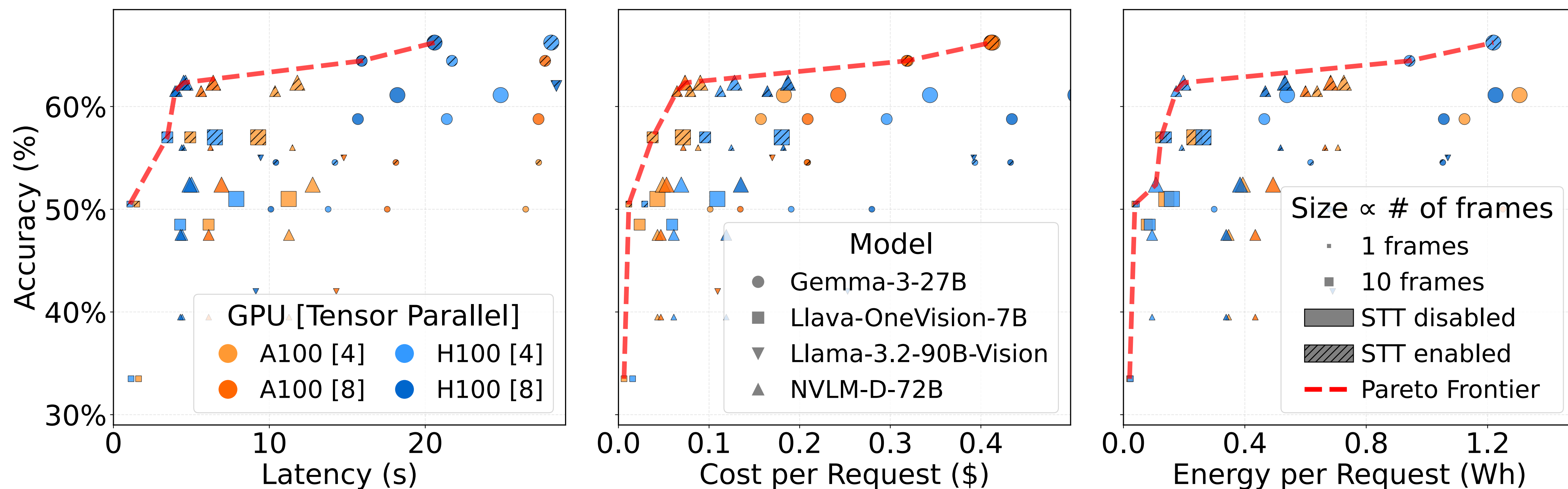
High-dimensional configuration space

Multiple objectives/SLOs



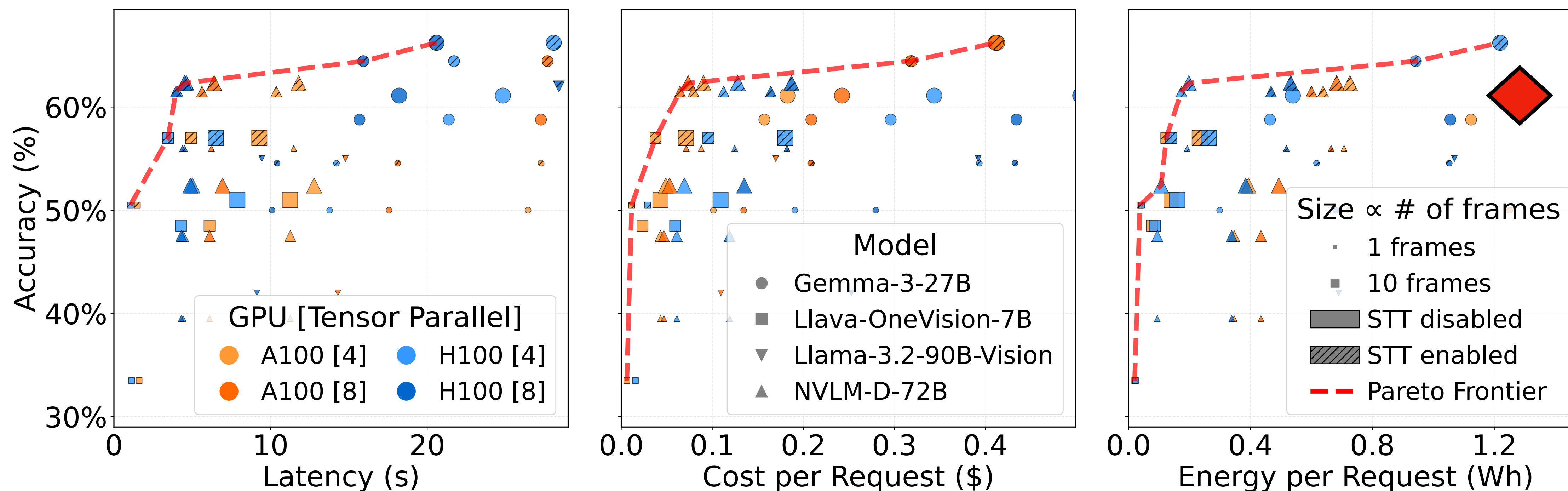
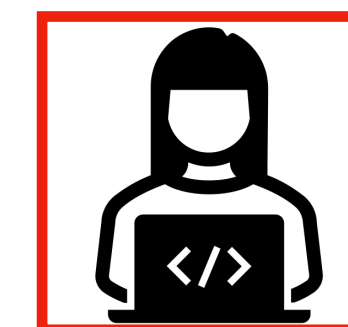
High-dimensional configuration space

Multiple objectives/SLOs



High-dimensional configuration space

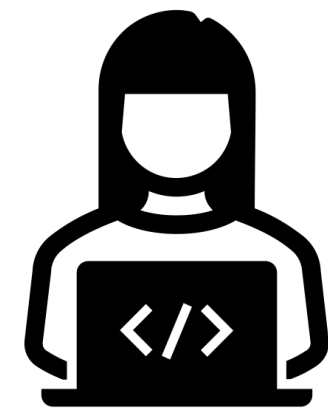
Manual configuration is intractable and often suboptimal...



Sources of resource-inefficiency today:
(1) Rigid workflows → difficult to adapt

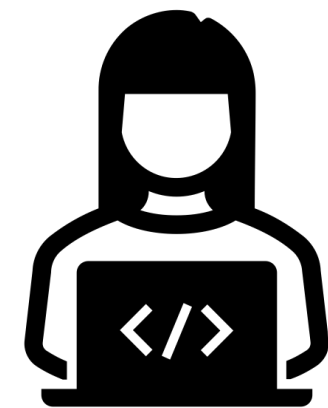
Entities involved in an agentic workflow's lifecycle

Entities involved in an agentic workflow's lifecycle

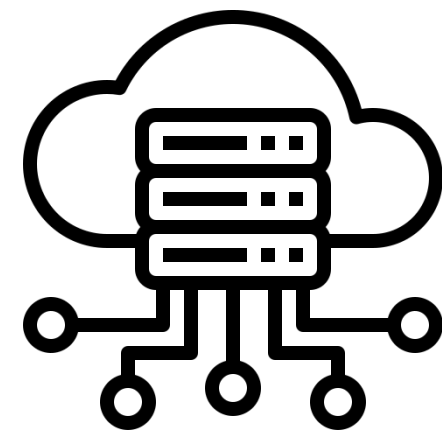


App. Developer

Entities involved in an agentic workflow's lifecycle

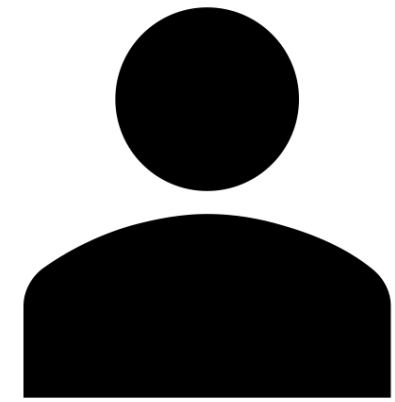


App. Developer

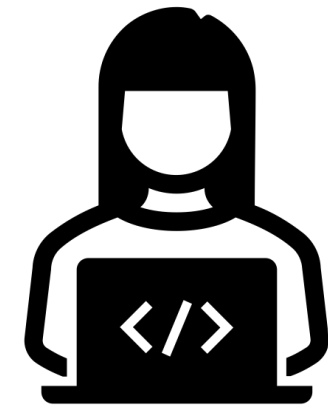


Cloud Platform

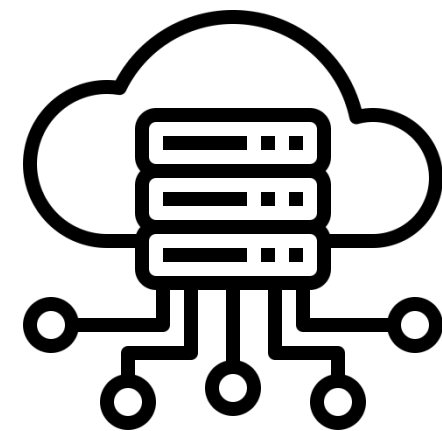
Entities involved in an agentic workflow's lifecycle



Users



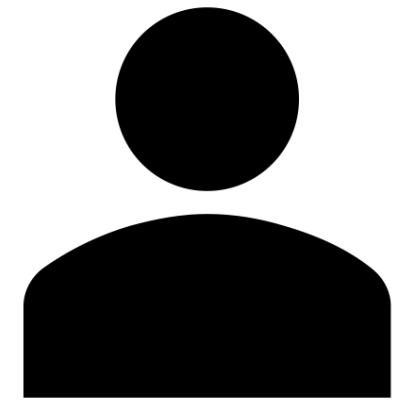
App. Developer



Cloud Platform

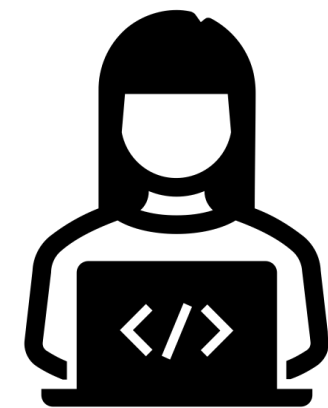
Entities involved in an agentic workflow's lifecycle

Objectives for each entity in the stack



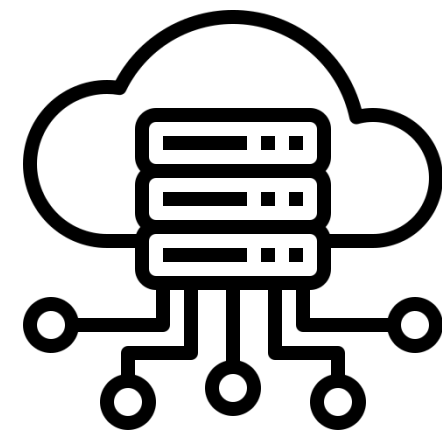
Users

- Cost
- Latency
- Accuracy



App. Developer

- Profitability
- User-experience

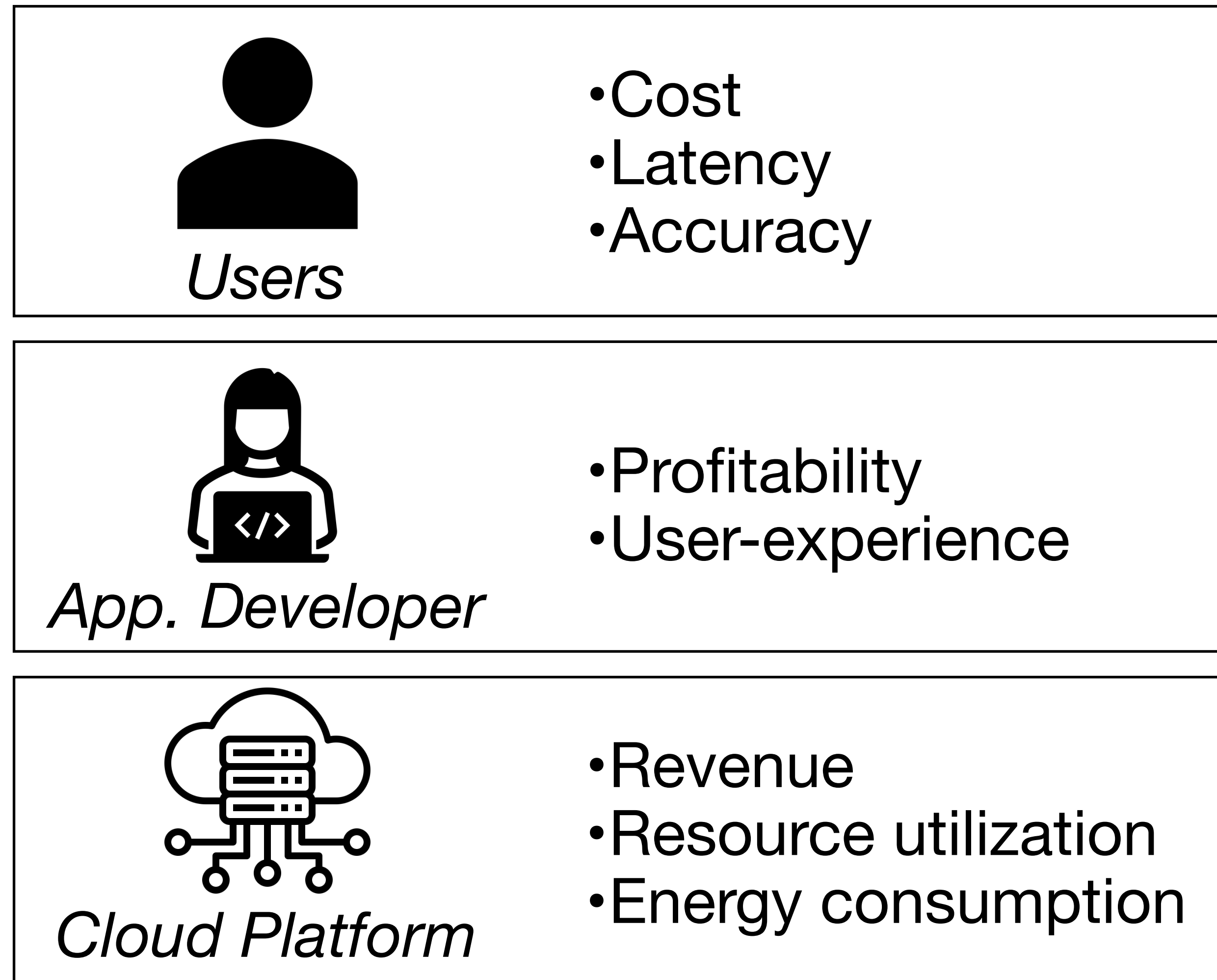


Cloud Platform

- Revenue
- Resource utilization
- Energy consumption

Entities involved in an agentic workflow's lifecycle

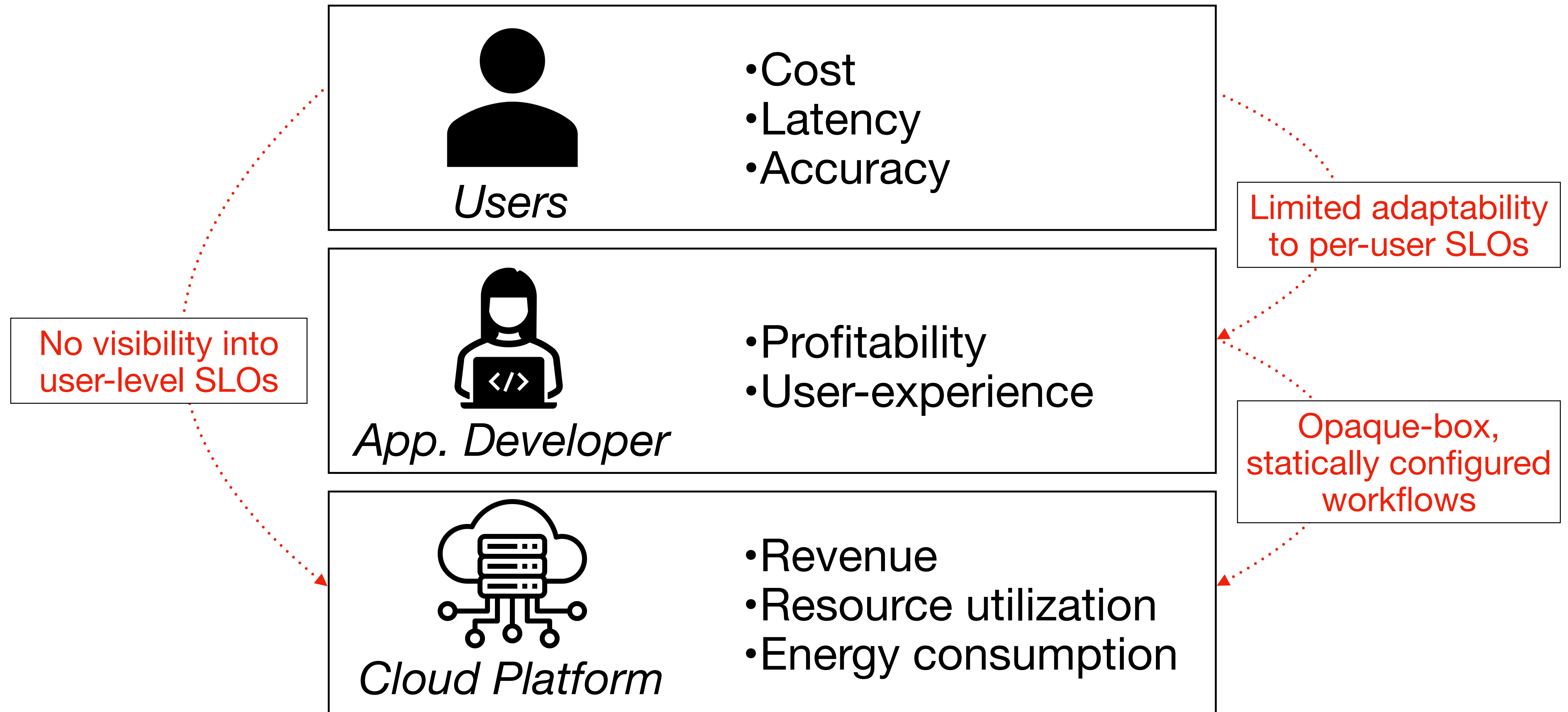
Exist in silos, limited cross-stack visibility...



Opaque-box,
statically configured
workflows

Entities involved in an agentic workflow's lifecycle

Exist in silos, limited cross-stack visibility...



Sources of resource-inefficiency today:

- (1) Rigid workflows → difficult to adapt
- (2) Fragmented stack → limits cross-stack optimization

We need:
Unified, end-to-end orchestration to
deliver high resource-efficiency

Murakkab

Resource-Efficient Agentic Workflow Orchestration in Cloud Platforms

Murakkab

Resource-Efficient Agentic Workflow Orchestration in Cloud Platforms

1. Declarative Workflow Specification

```
1  # == Sub-tasks in the workflow ==
2  scene_detect = "Given a list of videos, identify scenes in each."
3  frame_extract = "Given a list of scenes, extract frames."
4  stt           = "Given a list of scenes, convert audio to text."
5  q_a          = "Answer the query given some context."
6  # == Workflow description (sub-tasks and data flow) ==
7  def workflow(query, videos):
8      scenes = scene_detect(videos)
9      frames = frame_extract(scenes)
10     transcript = stt(scenes)
11     answer = q_a(query, [frames, transcript])
12     return answer
13 # == Execution with example request ==
14 query = "What is the name of the person wearing the red dress?"
15 videos = ["road_trip.mp4"]
16 result = run(workflow(query, videos), slo=LOW_LATENCY)
```

Separate the “what” from the “how”.

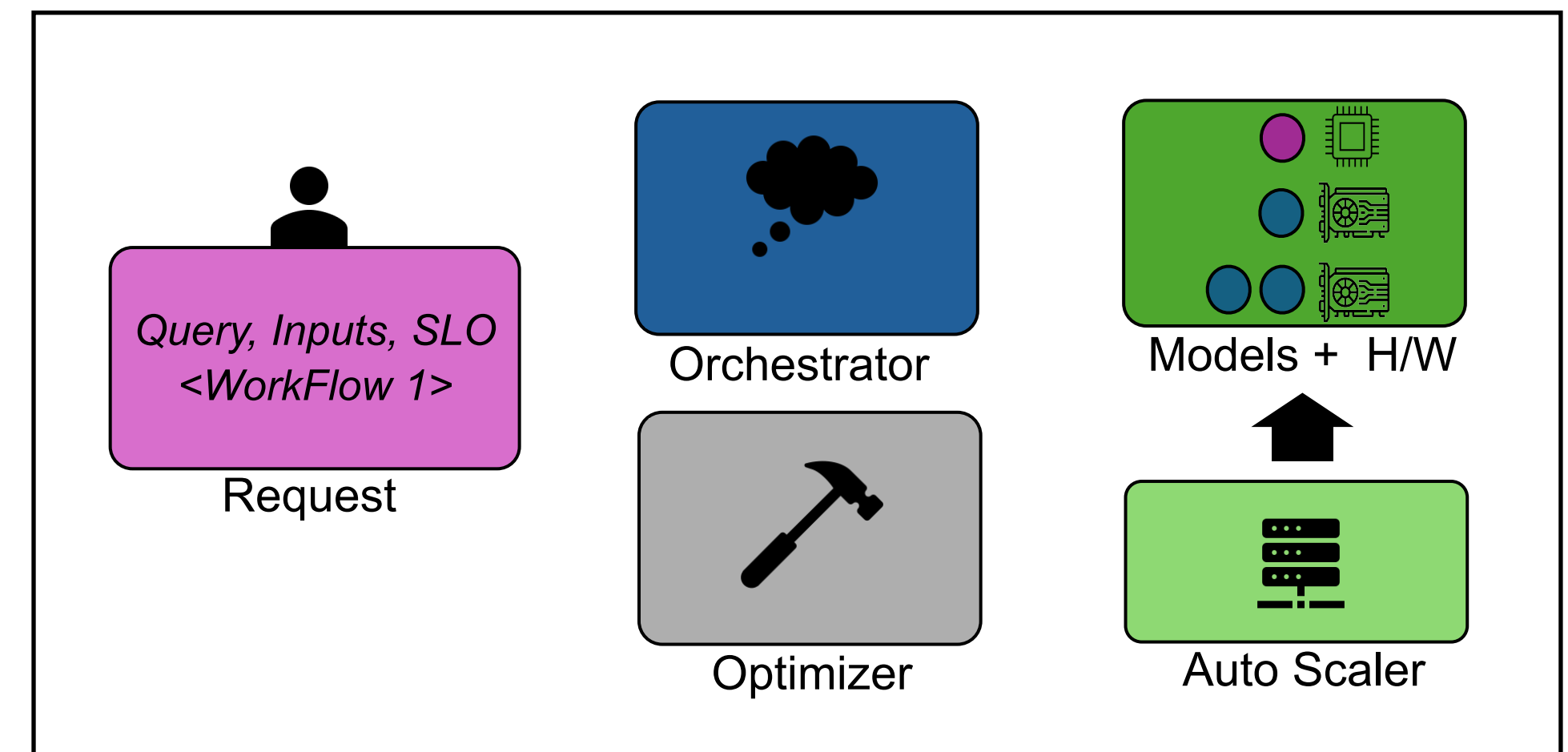
Murakkab

Resource-Efficient Agentic Workflow Orchestration in Cloud Platforms

1. Declarative Workflow Specification

```
1 # == Sub-tasks in the workflow ==
2 scene_detect = "Given a list of videos, identify scenes in each."
3 frame_extract = "Given a list of scenes, extract frames."
4 stt = "Given a list of scenes, convert audio to text."
5 q_a = "Answer the query given some context."
6 # == Workflow description (sub-tasks and data flow) ==
7 def workflow(query, videos):
8     scenes = scene_detect(videos)
9     frames = frame_extract(scenes)
10    transcript = stt(scenes)
11    answer = q_a(query, [frames, transcript])
12    return answer
13 # == Execution with example request ==
14 query = "What is the name of the person wearing the red dress?"
15 videos = ["road_trip.mp4"]
16 result = run(workflow(query, videos), slo=LOW_LATENCY)
```

2. Adaptive, SLO-aware Orchestration



Separate the “what” from the “how”.

**Configuration is a runtime decision,
not a development decision.**

Phase 1: Development

Declarative Workflow Specification

```
1  # == Sub-tasks in the workflow ==
2  scene_detect    = "Given a list of videos, identify scenes in each."
3  frame_extract   = "Given a list of scenes, extract frames."
4  stt             = "Given a list of scenes, convert audio to text."
5  q_a            = "Answer the query given some context."
6  # == Workflow description (sub-tasks and data flow) ==
7  def workflow(query, videos):
8      scenes      = scene_detect(videos)
9      frames      = frame_extract(scenes)
10     transcript   = stt(scenes)
11     answer       = q_a(query, [frames, transcript])
12     return answer
13 # == Execution with example request ==
14 query    = "What is the name of the person wearing the red dress?"
15 videos   = ["road_trip.mp4"]
16 result   = run(workflow(query, videos), slo=LOW_LATENCY)
```

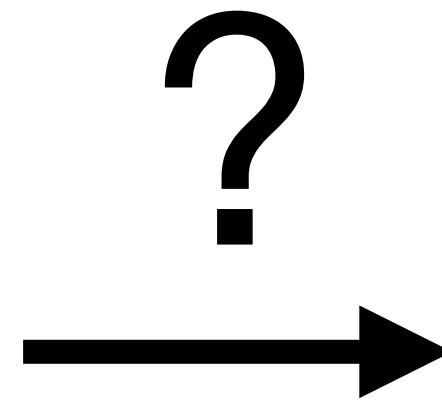
Application Logic

SLOs

Challenge

How do we capture the missing details?

```
1 # == Workflow nodes (coupled app logic + configs) ==
2 scene_detection = Tool(
3     fn=SceneDetector(), # Implemented earlier by the
4     # developer.
5     key=AWS_KEY, resources={"CPUs": 32}
6 )
7 frame_extractor = Tool(
8     fn=FrameExtractor(), # Implemented earlier by the
9     # developer.
10    key=AWS_KEY, params={"num_frames": 15}, resources={"CPUs": 32}
11 )
12 speech_to_text = MLModel(
13     name="Whisper", key=OPENAI_API_KEY, resources={"PTUs": 50}
14 )
15 object_detection = MLModel(
16     name="CLIP", key=AZURE_KEY, resources={"CPUs": 128}
17 )
18 question_answer = LLM(
19     name="Llama-3.2", key=DATABRICKS_API_KEY,
20     params={"batch": 256}, resources={"GPUs": 8, "Type": "H100"},
21     system_prompt="You are an agent that can understand videos.",
22     user_prompt="Answer the given question about the video."
23 )
24 # == Query and data ==
25 ques = "What is the name of the person wearing the red dress?"
26 videos = ["road_trip.mp4"]
27 # == Workflow (ordering of nodes and data flow) ==
28 scenes, audio = scene_detection(videos)
29 frames = frame_extractor(scenes)
30 transcript = speech_to_text(audio)
31 obj_frames = object_detection(frames)
32 answer = question_answer(ques, transcript, obj_frames)
```



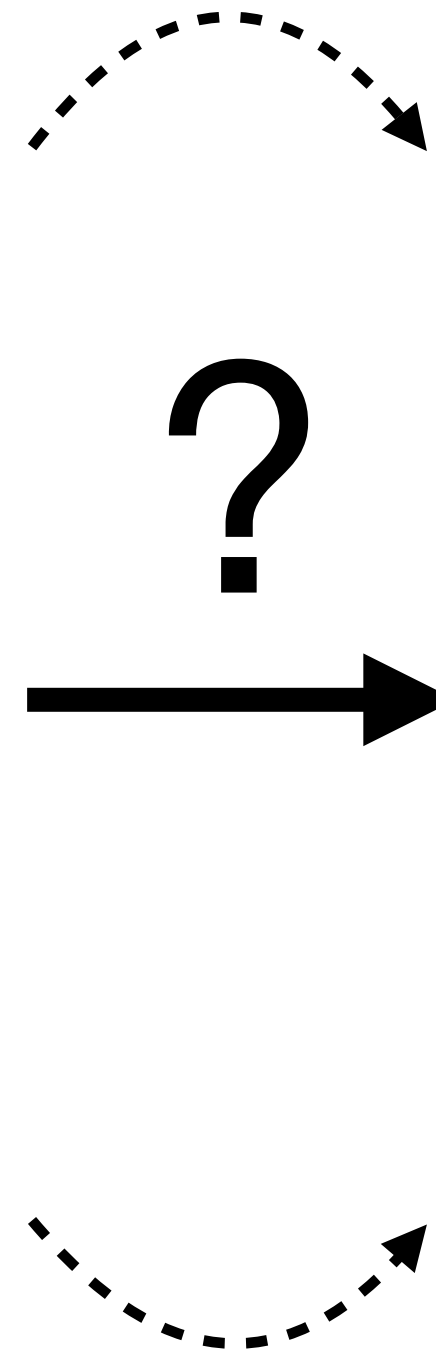
```
1 # == Sub-tasks in the workflow ==
2 scene_detect = "Given a list of videos, identify scenes in each."
3 frame_extract = "Given a list of scenes, extract frames."
4 stt = "Given a list of scenes, convert audio to text."
5 q_a = "Answer the query given some context."
6 # == Workflow description (sub-tasks and data flow) ==
7 def workflow(query, videos):
8     scenes = scene_detect(videos)
9     frames = frame_extract(scenes)
10    transcript = stt(scenes)
11    answer = q_a(query, [frames, transcript])
12    return answer
13 # == Execution with example request ==
14 query = "What is the name of the person wearing the red dress?"
15 videos = ["road_trip.mp4"]
16 result = run(workflow(query, videos), slo=LOW_LATENCY)
```

Challenge

How do we capture the missing details?

```
1 # == Workflow nodes (coupled app logic + configs) ==
2 scene_detection = Tool(
3     fn=SceneDetector(),      # Implemented earlier by the
4                               developer.
5     key=AWS_KEY, resources={"CPUs": 32}
6 )
7 frame_extractor = Tool(
8     fn=FrameExtractor(),     # Implemented earlier by the
9                               developer.
10    key=AWS_KEY, params={"num_frames": 15}, resources={"CPUs": 32}
11 )
12 speech_to_text = MLModel(
13     name="Whisper", key=OPENAI_API_KEY, resources={"PTUs": 50}
14 )
15 object_detection = MLModel(
16     name="CLIP", key=AZURE_KEY, resources={"CPUs": 128}
17 )
18 question_answer = LLM(
19     name="Llama-3.2", key=DATABRICKS_API_KEY,
20     params={"batch": 256}, resources={"GPUs": 8, "Type": "H100"},
21     system_prompt="You are an agent that can understand videos.",
22     user_prompt="Answer the given question about the video."
23 )
24 # == Query and data ==
25 ques      = "What is the name of the person wearing the red dress?"
26 videos    = ["road_trip.mp4"]
27 # == Workflow (ordering of nodes and data flow) ==
28 scenes, audio = scene_detection(videos)
29 frames        = frame_extractor(scenes)
30 transcript     = speech_to_text(audio)
31 obj_frames    = object_detection(frames)
32 answer        = question_answer(ques, transcript, obj_frames)
```

Which agents to use?



```
1 # == Sub-tasks in the workflow ==
2 scene_detect      = "Given a list of videos, identify scenes in each."
3 frame_extract     = "Given a list of scenes, extract frames."
4 stt               = "Given a list of scenes, convert audio to text."
5 q_a              = "Answer the query given some context."
6 # == Workflow description (sub-tasks and data flow) ==
7 def workflow(query, videos):
8     scenes        = scene_detect(videos)
9     frames        = frame_extract(scenes)
10    transcript     = stt(scenes)
11    answer         = q_a(query, [frames, transcript])
12    return answer
13 # == Execution with example request ==
14 query           = "What is the name of the person wearing the red dress?"
15 videos          = ["road_trip.mp4"]
16 result          = run(workflow(query, videos), slo=LOW_LATENCY)
```

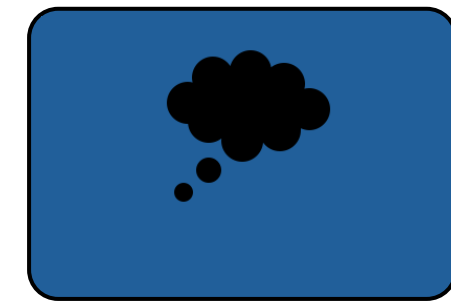
How to configure the workflow?

Challenge

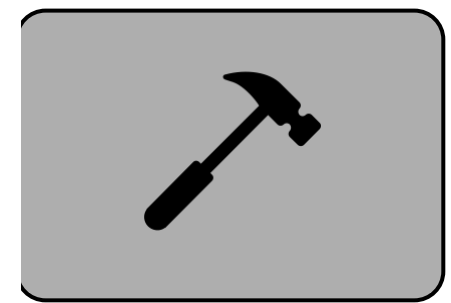
Murakkab: Orchestrator + Optimizer

```
1 # == Workflow nodes (coupled app logic + configs) ==
2 scene_detection = Tool(
3     fn=SceneDetector(),      # Implemented earlier by the
4                               developer.
5     key=AWS_KEY, resources={"CPUs": 32}
6 )
7 frame_extractor = Tool(
8     fn=FrameExtractor(),     # Implemented earlier by the
9                               developer.
10    key=AWS_KEY, params={"num_frames": 15}, resources={"CPUs": 32}
11 )
12 speech_to_text = MLModel(
13     name="Whisper", key=OPENAI_API_KEY, resources={"PTUs": 50}
14 )
15 object_detection = MLModel(
16     name="CLIP", key=AZURE_KEY, resources={"CPUs": 128}
17 )
18 question_answer = LLM(
19     name="Llama-3.2", key=DATABRICKS_API_KEY,
20     params={"batch": 256}, resources={"GPUs": 8, "Type": "H100"},
21     system_prompt="You are an agent that can understand videos.",
22     user_prompt="Answer the given question about the video."
23 )
24 # == Query and data ==
25 ques      = "What is the name of the person wearing the red dress?"
26 videos    = ["road_trip.mp4"]
27 # == Workflow (ordering of nodes and data flow) ==
28 scenes, audio = scene_detection(videos)
29 frames        = frame_extractor(scenes)
30 transcript     = speech_to_text(audio)
31 obj_frames    = object_detection(frames)
32 answer        = question_answer(ques, transcript, obj_frames)
```

Which agents to use?



Orchestrator



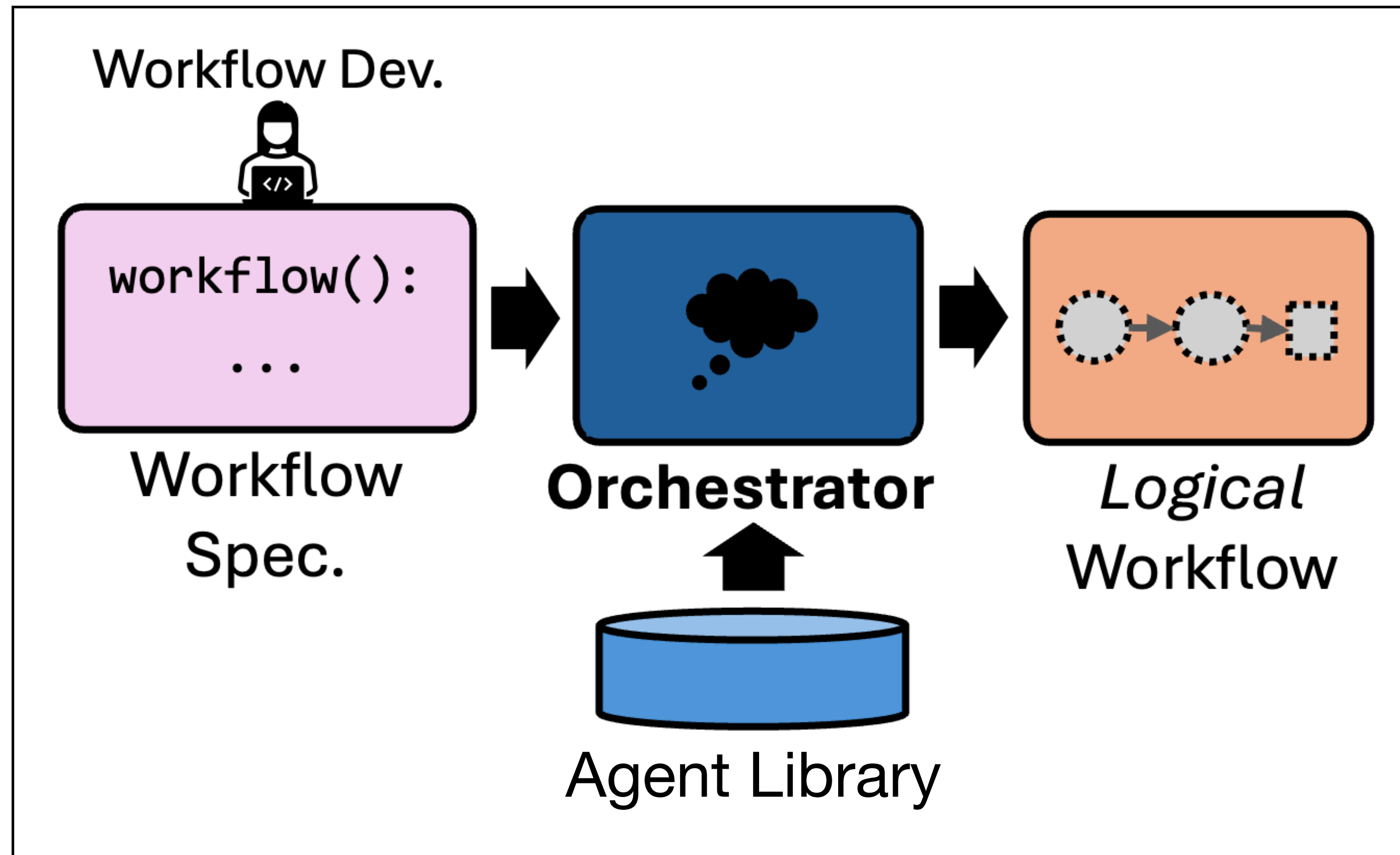
Optimizer

```
1 # == Sub-tasks in the workflow ==
2 scene_detect = "Given a list of videos, identify scenes in each."
3 frame_extract = "Given a list of scenes, extract frames."
4 stt           = "Given a list of scenes, convert audio to text."
5 q_a          = "Answer the query given some context."
6 # == Workflow description (sub-tasks and data flow) ==
7 def workflow(query, videos):
8     scenes      = scene_detect(videos)
9     frames      = frame_extract(scenes)
10    transcript   = stt(scenes)
11    answer       = q_a(query, [frames, transcript])
12    return answer
13 # == Execution with example request ==
14 query         = "What is the name of the person wearing the red dress?"
15 videos        = ["road_trip.mp4"]
16 result        = run(workflow(query, videos), slo=LOW_LATENCY)
```

How to configure the workflow?

Phase 1: Development

Orchestrator: Workflow Specification to Logical Workflow



Phase 1: Development

Separate the “what” from the “how”

```
1  # == Sub-tasks in the workflow ==
2  scene_detect    = "Given a list of videos, identify scenes in each."
3  frame_extract   = "Given a list of scenes, extract frames."
4  stt             = "Given a list of scenes, convert audio to text."
5  q_a            = "Answer the query given some context."
6  # == Workflow description (sub-tasks and data flow) ==
7  def workflow(query, videos):
8      scenes      = scene_detect(videos)
9      frames      = frame_extract(scenes)
10     transcript   = stt(scenes)
11     answer       = q_a(query, [frames, transcript])
12     return answer
13 # == Execution with example request ==
14 query    = "What is the name of the person wearing the red dress?"
15 videos   = ["road_trip.mp4"]
16 result   = run(workflow(query, videos), slo=LOW_LATENCY)
```

FrameExtractor
(videos, <num_frames>)

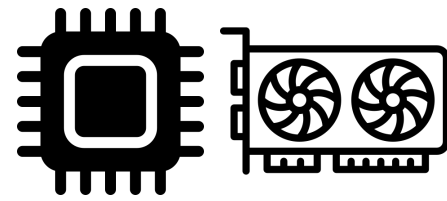
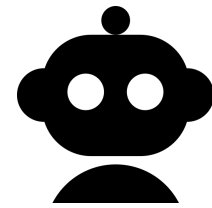
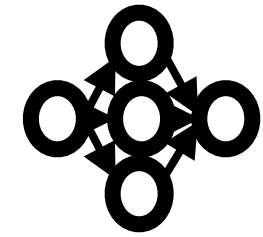
QuestionAnswer
(query, ctx, <model>)

Model
(<gpu_type>, <TP>)

Phase 2: Optimization + Execution

Profile-Guided Optimization

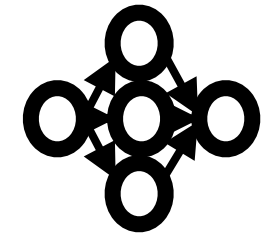
- **Workflow-level** profiles
- **Agent-level** profiles
- **Hardware-level** profiles



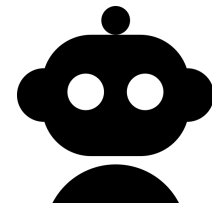
Phase 2: Optimization + Execution

Profile-Guided Optimization

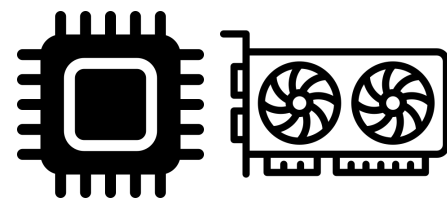
- **Workflow-level** profiles



- **Agent-level** profiles



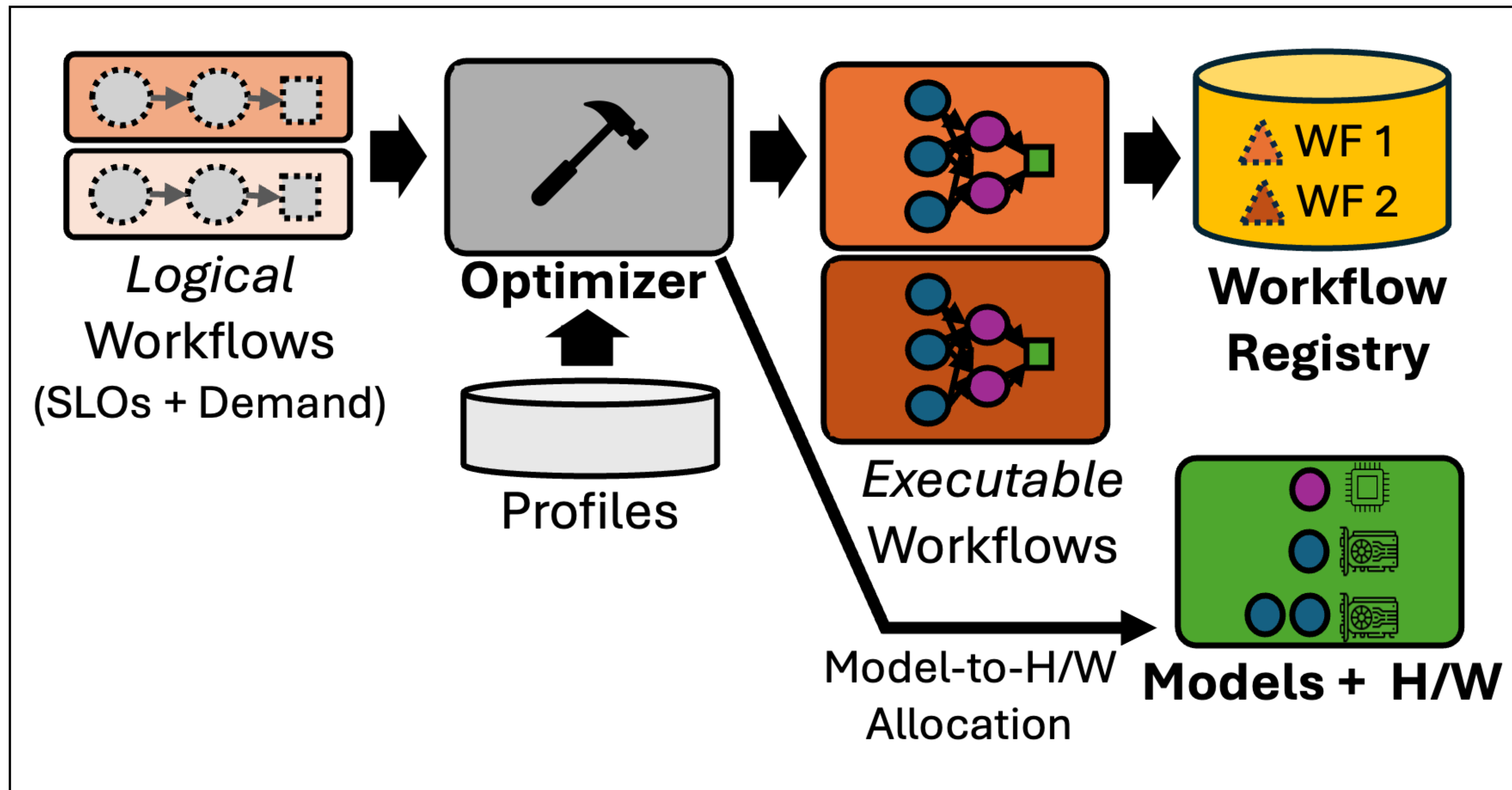
- **Hardware-level** profiles



- Composable profiles across levels:
 - Make runtime optimization decisions (e.g., quality, latency, cost, energy)
 - Maximizes **reuse** across workflows
 - Minimizes re-profiling

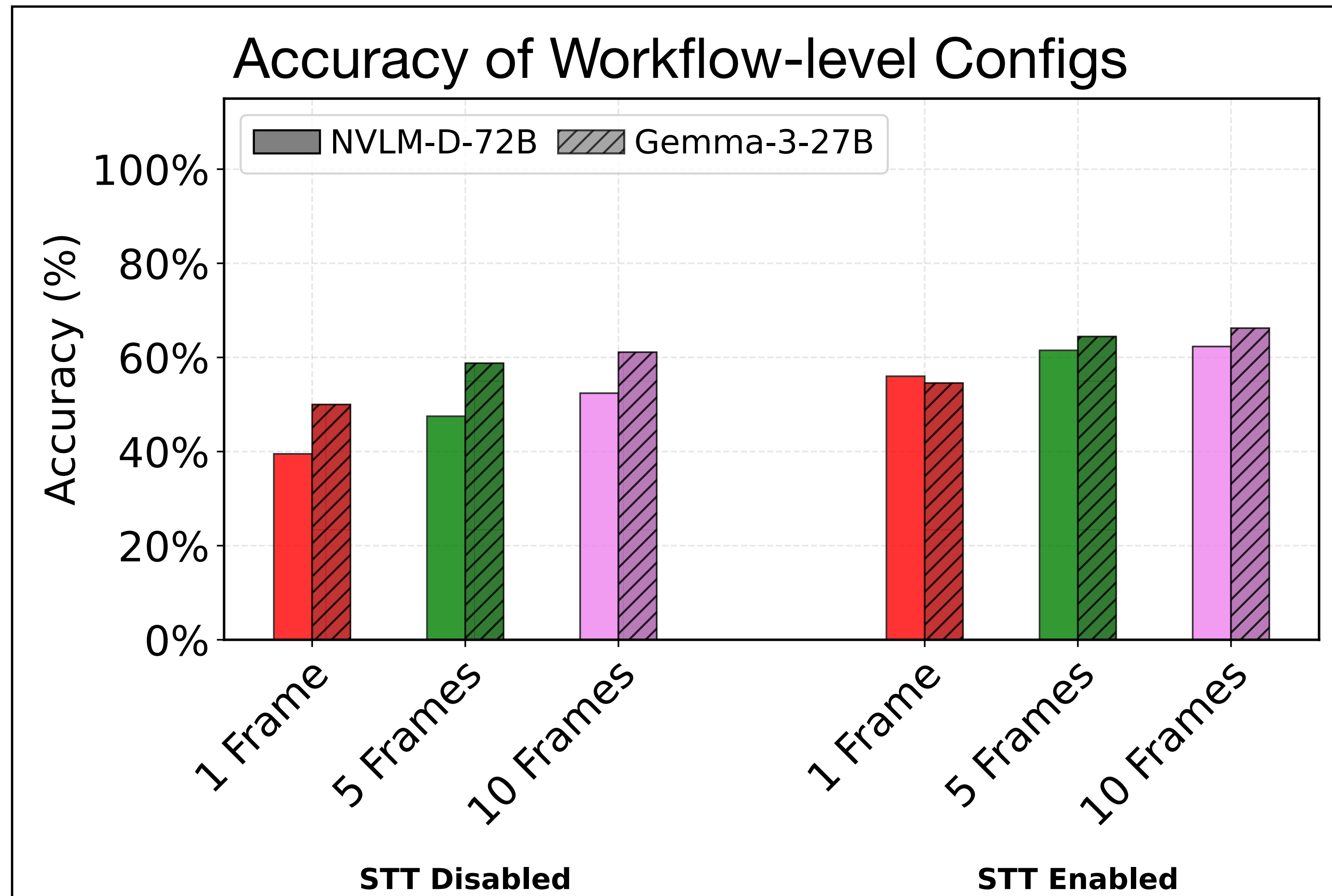
Phase 2: Optimization + Execution

Optimizer: Logical Workflow to Executable Workflow



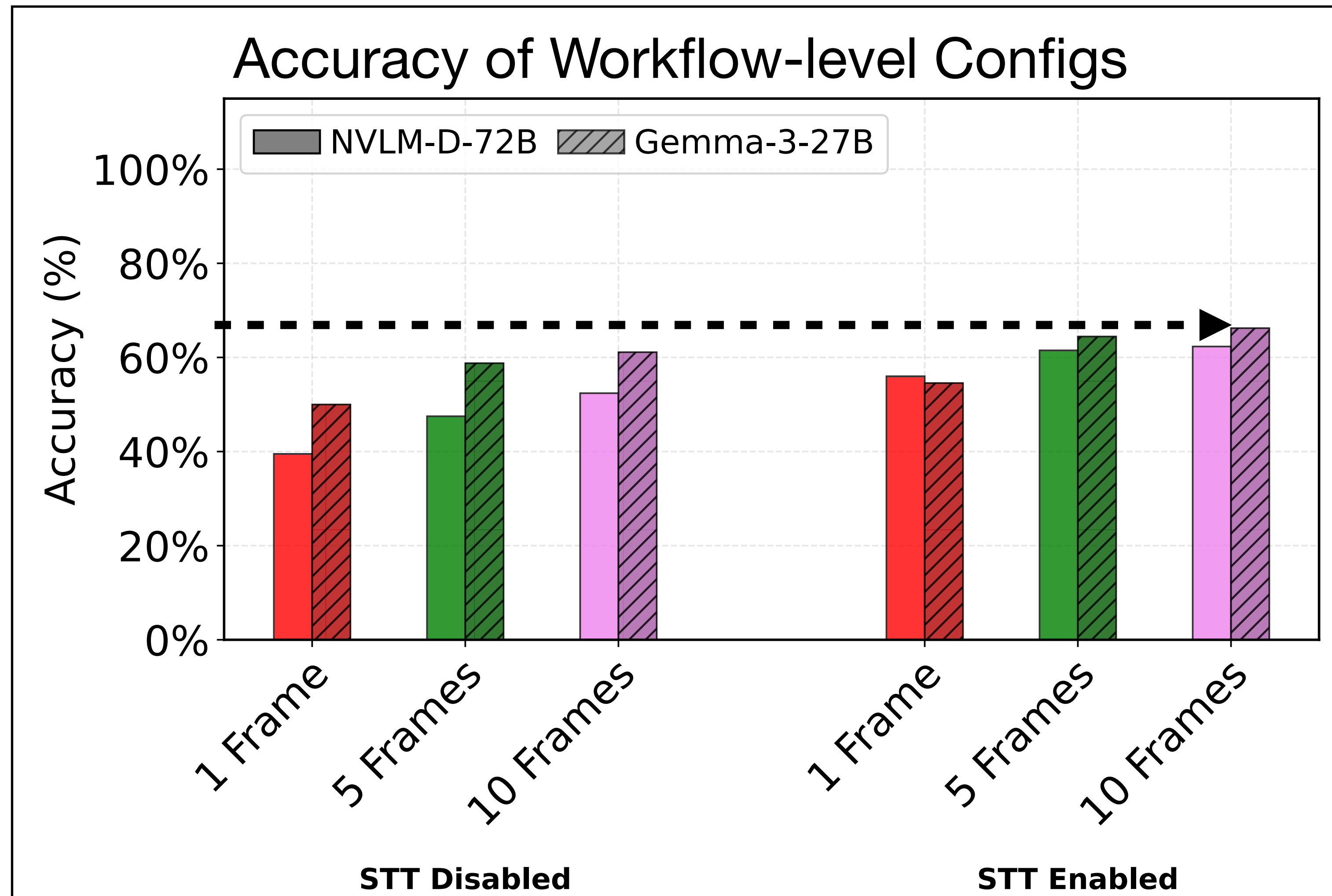
Phase 2: Optimization + Execution

Quality is determined using **workflow-level profiles**



Phase 2: Optimization + Execution

Quality is determined using **workflow-level profiles**

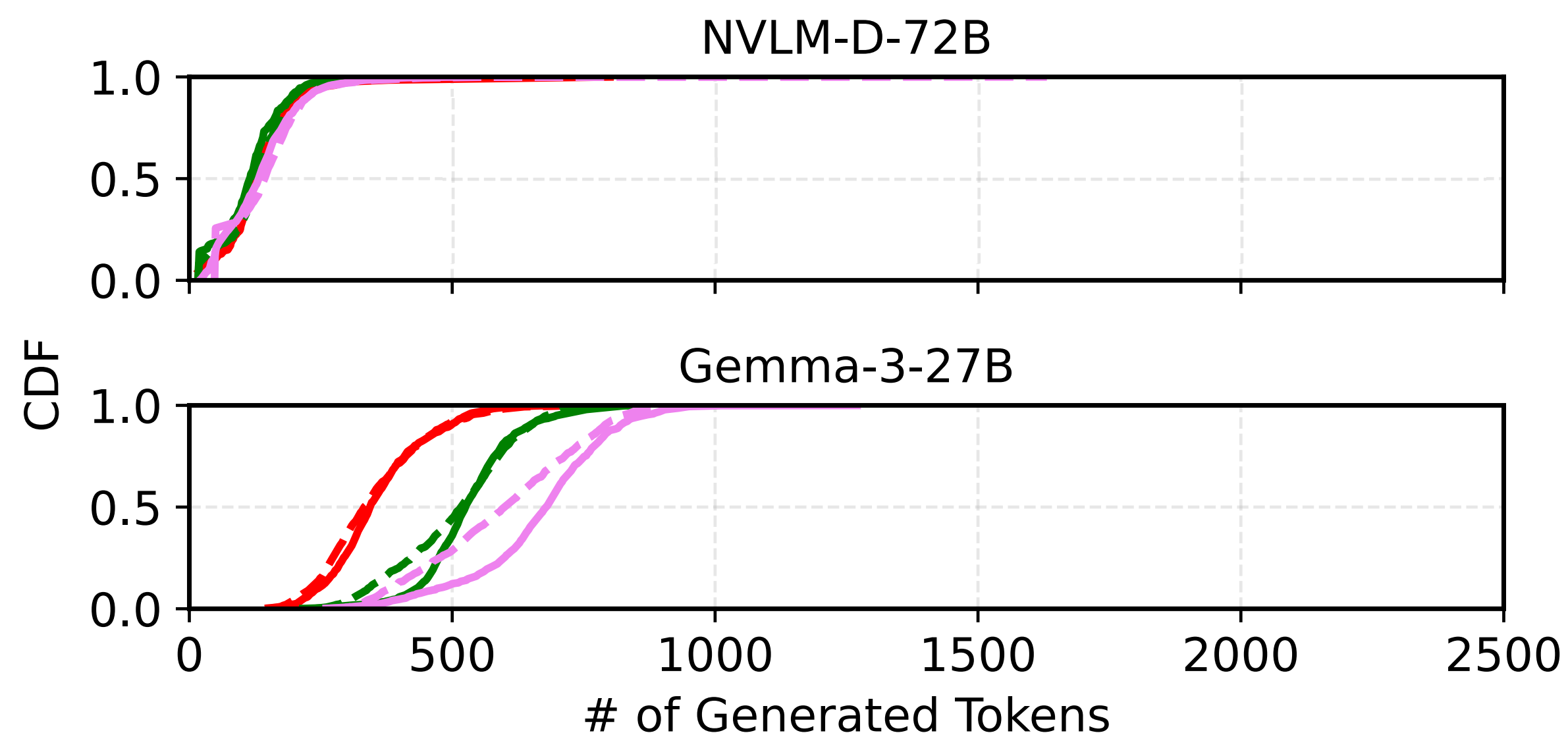


Phase 2: Optimization + Execution

Latency/cost is determined by agent- and hardware-level profiles

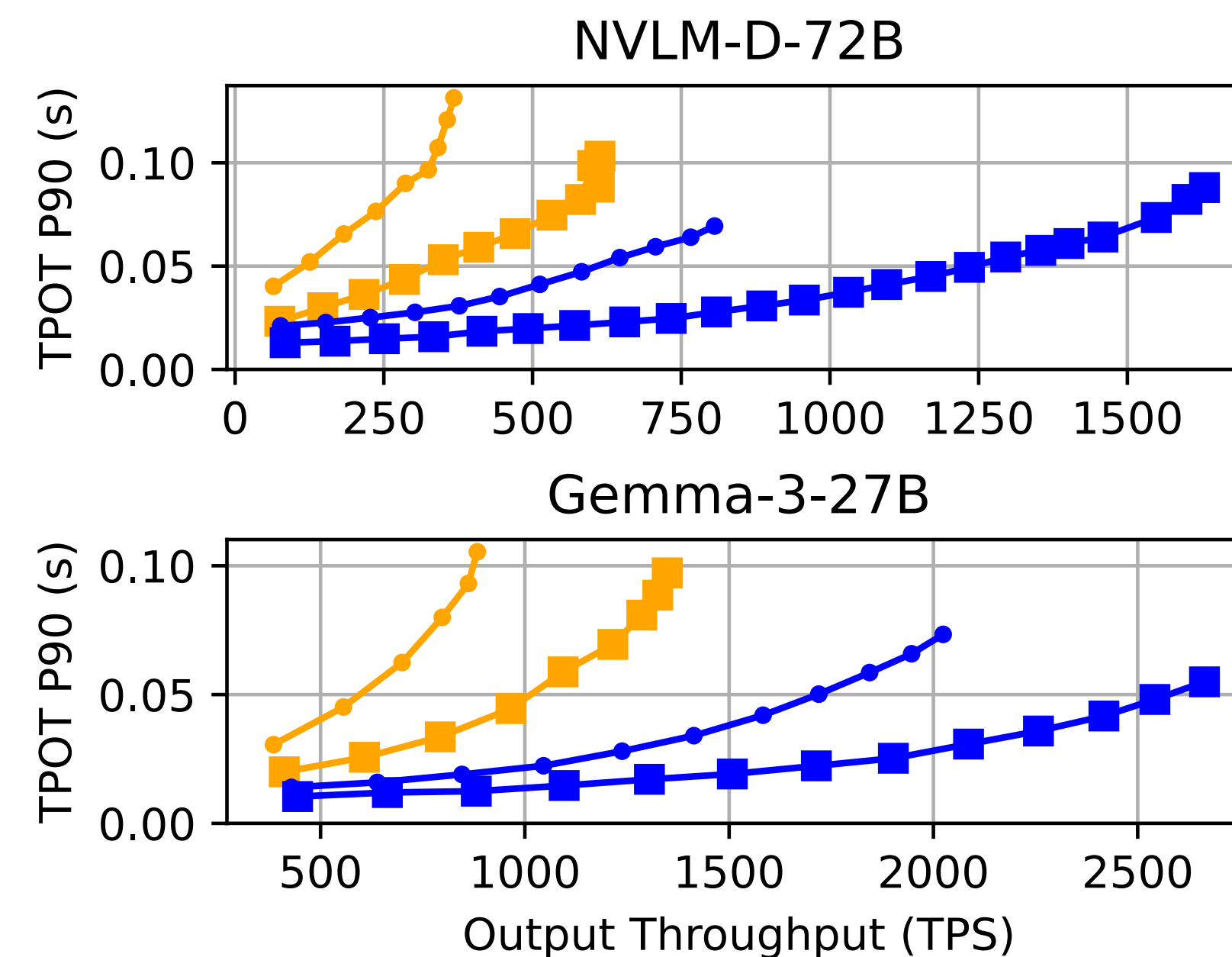
Token-Generation Load

Frames: 1 # Frames: 10 — STT: Y
Frames: 5 — STT: N



Token-Generation Latency

A100, TP=1 A100, TP=2 A100, TP=4 A100, TP=8
H100, TP=1 H100, TP=2 H100, TP=4 H100, TP=8

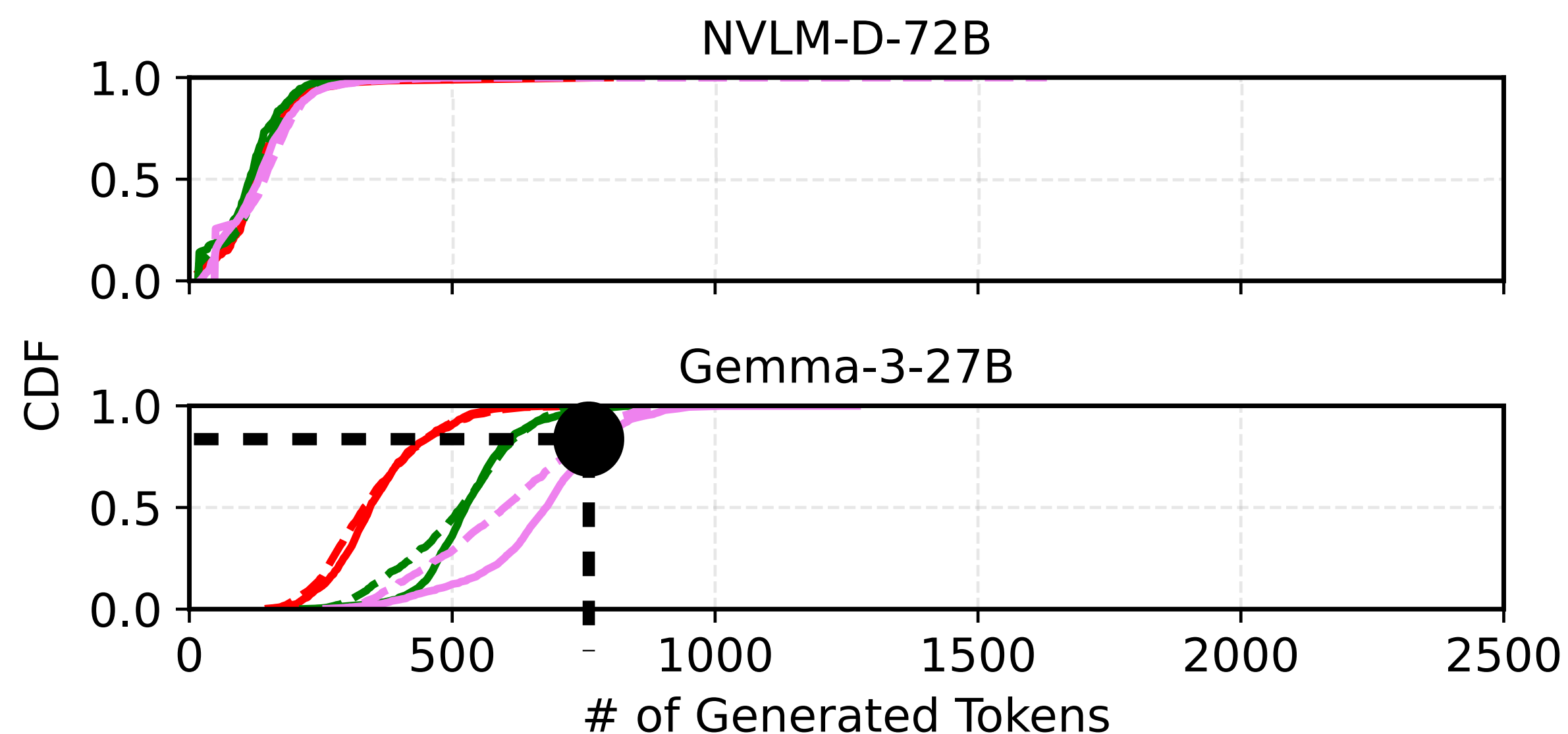


Phase 2: Optimization + Execution

Latency/cost is determined by agent- and hardware-level profiles

Token-Generation Load

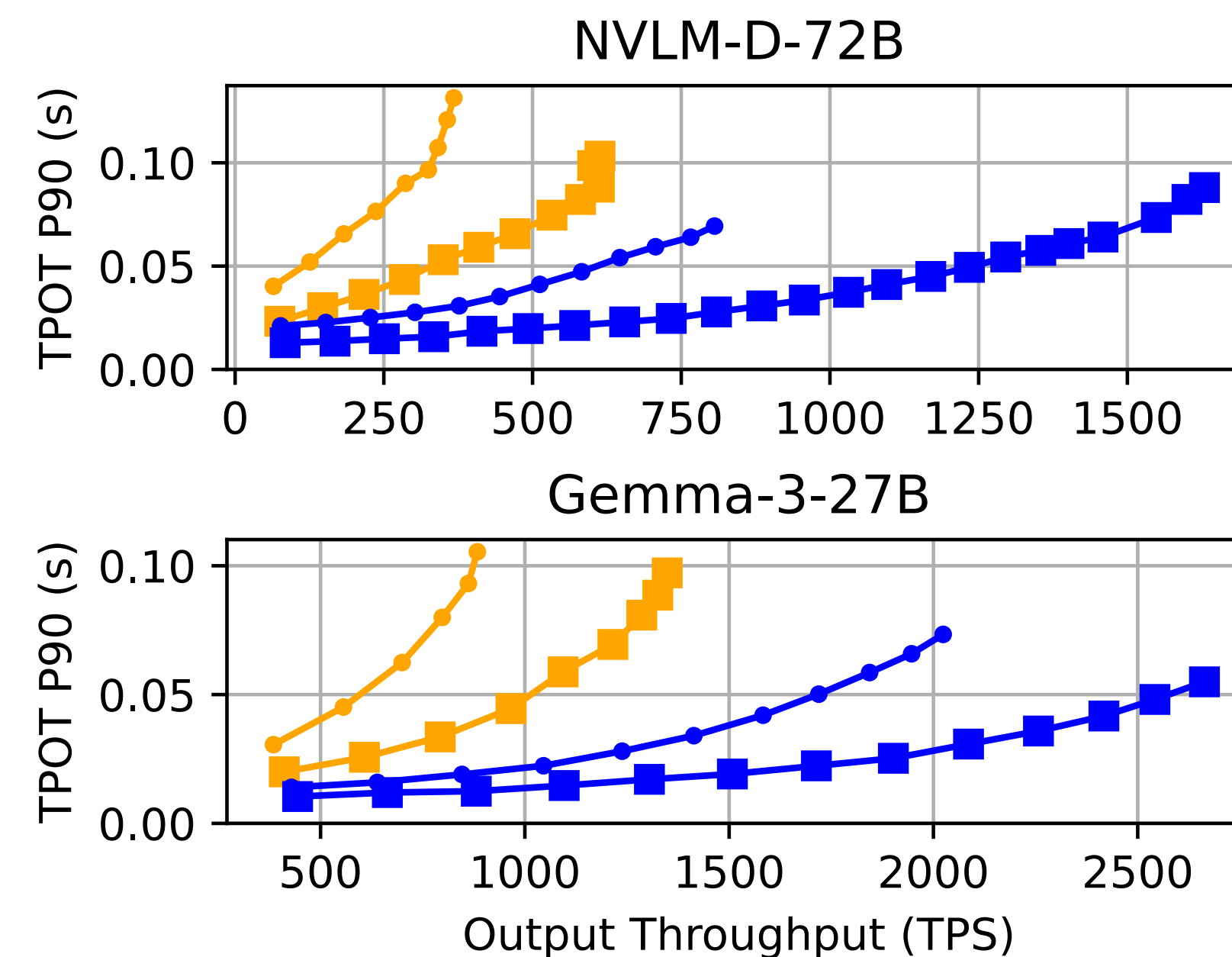
Frames: 1 # Frames: 10 — STT: Y
Frames: 5 — STT: N



~700 tokens @ P90

Token-Generation Latency

A100, TP=1 A100, TP=2 A100, TP=4 A100, TP=8
H100, TP=1 H100, TP=2 H100, TP=4 H100, TP=8

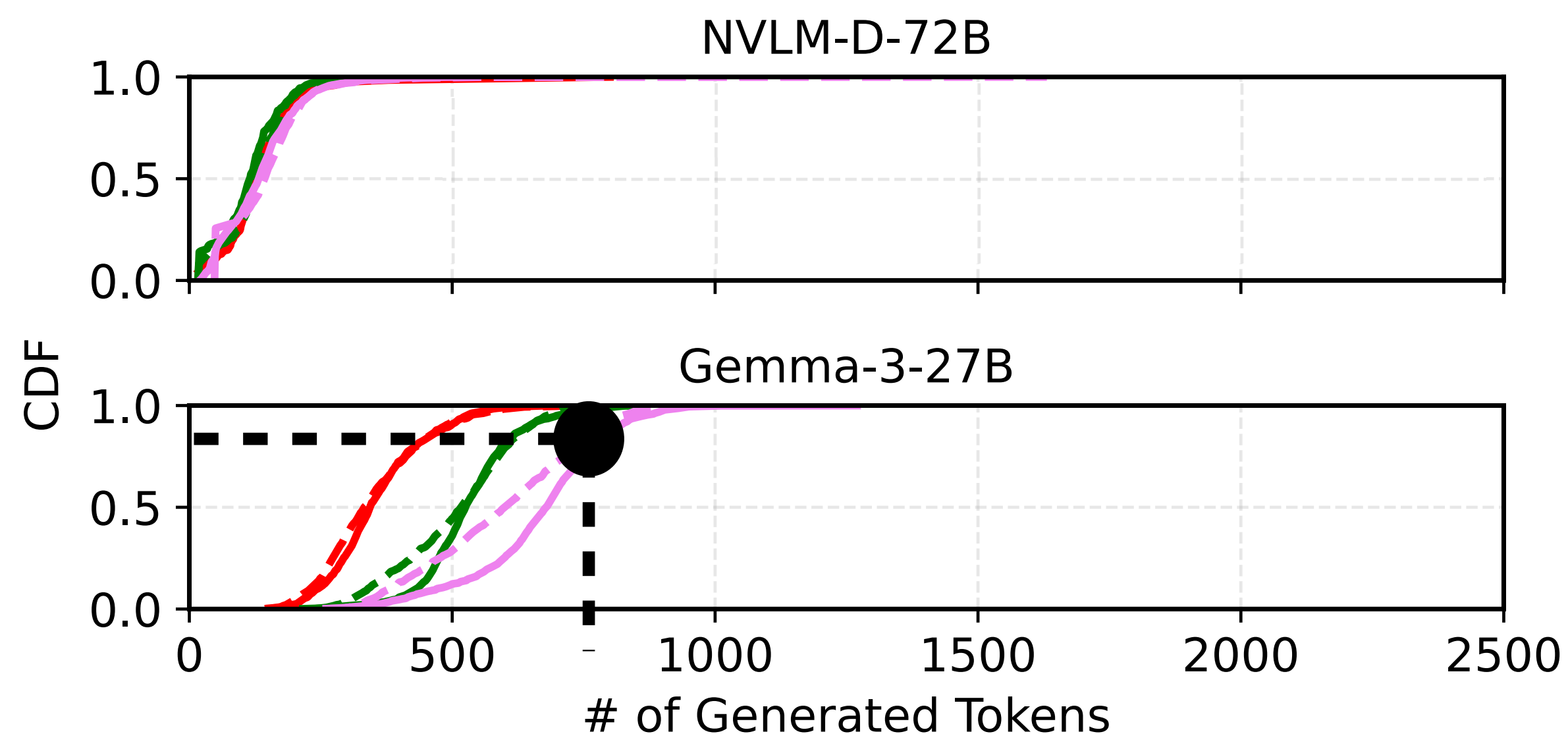


Phase 2: Optimization + Execution

Latency/cost is determined by agent- and hardware-level profiles

Token-Generation Load

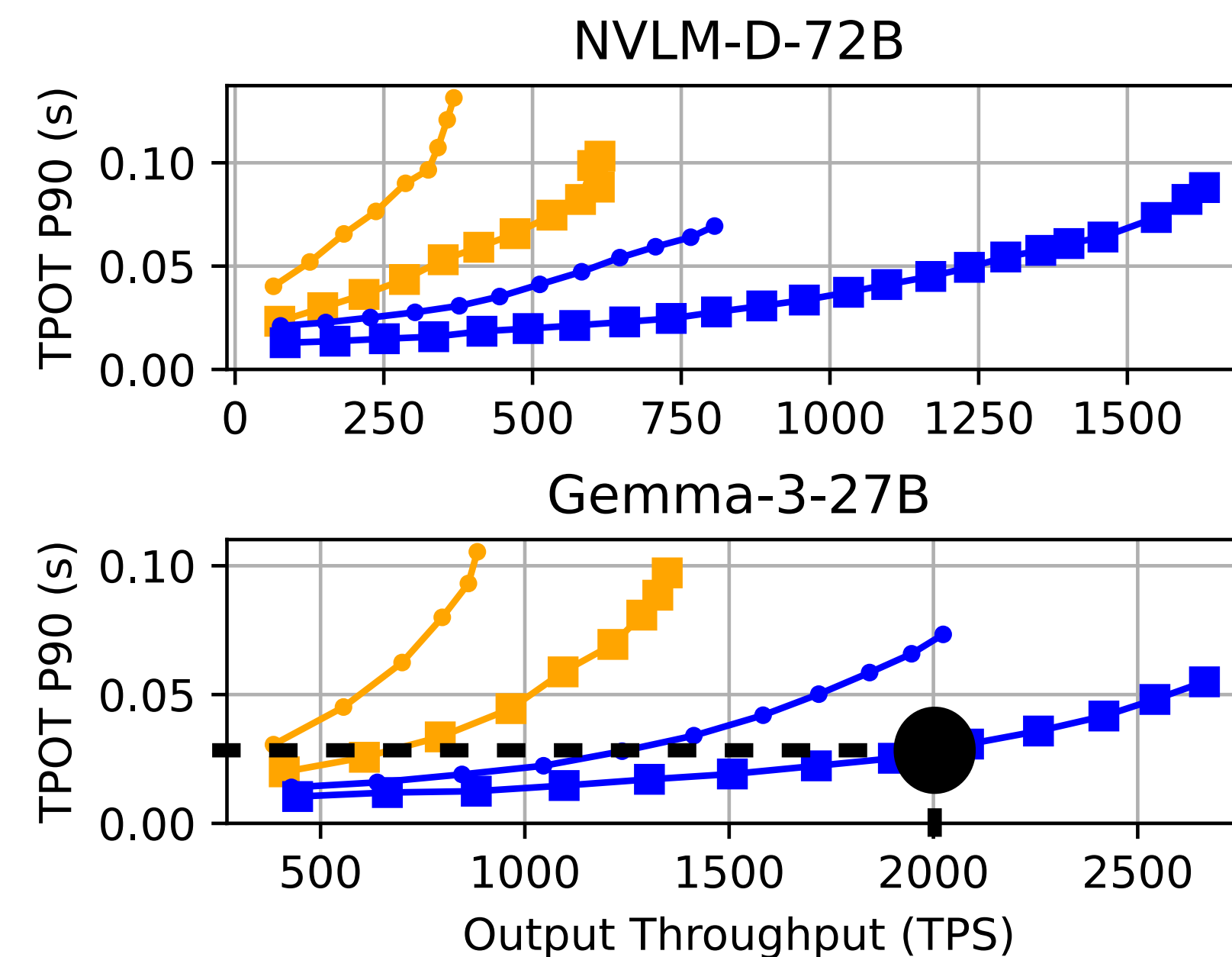
Frames: 1 # Frames: 10 — STT: Y
Frames: 5 — STT: N



~700 tokens @ P90

Token-Generation Latency

A100, TP=1 A100, TP=2 A100, TP=4 A100, TP=8
H100, TP=1 H100, TP=2 H100, TP=4 H100, TP=8



~0.03s @ 2000 TPS — 700 x 0.03s = 21s

Phase 2: Optimization + Execution

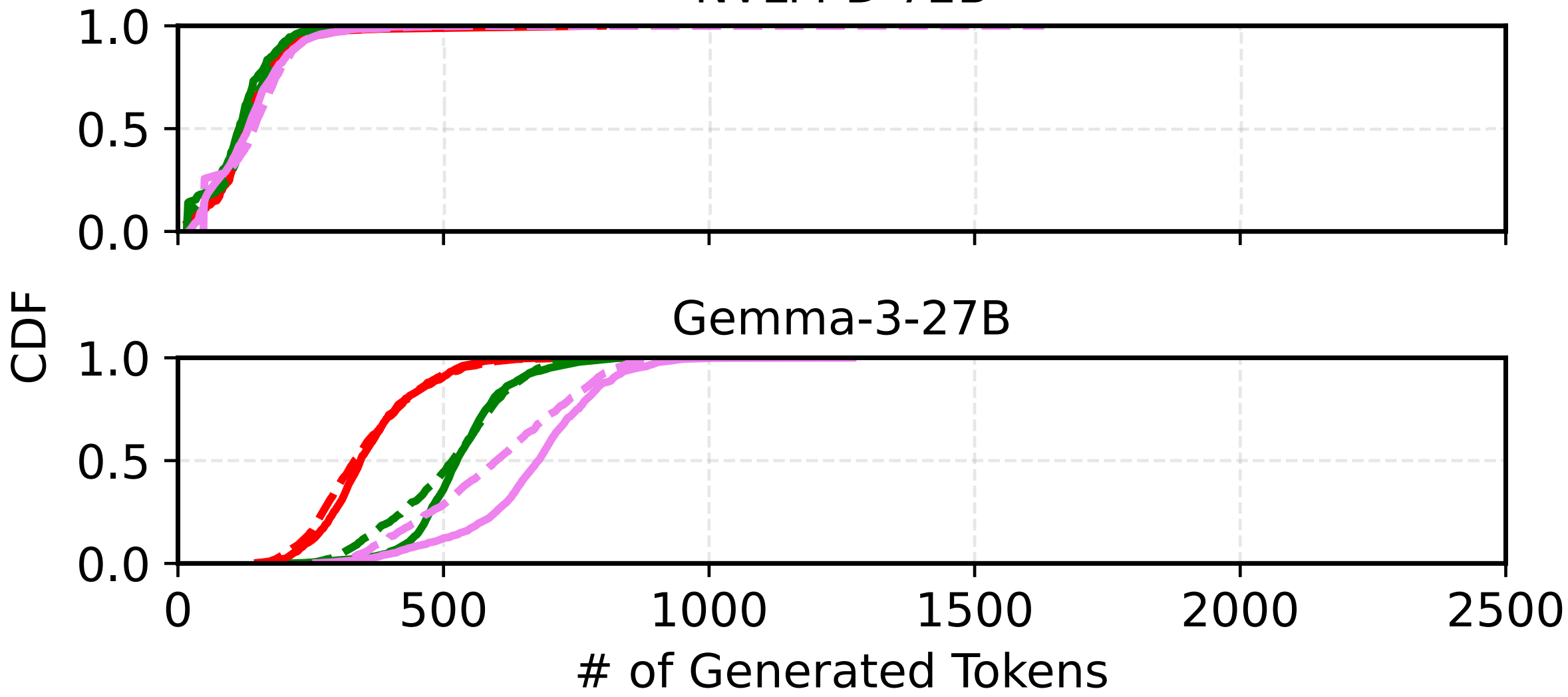
Energy efficiency is determined by agent- and hardware-level profiles

Token-Generation Load

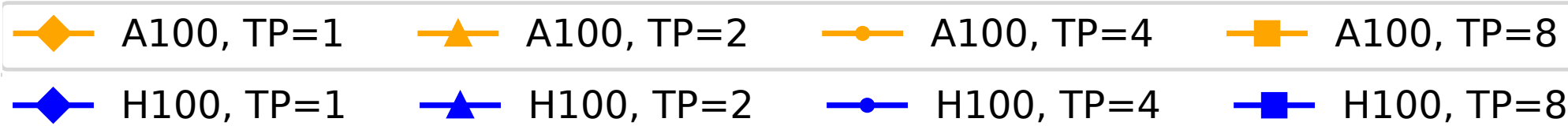


NVLM-D-72B

Gemma-3-27B

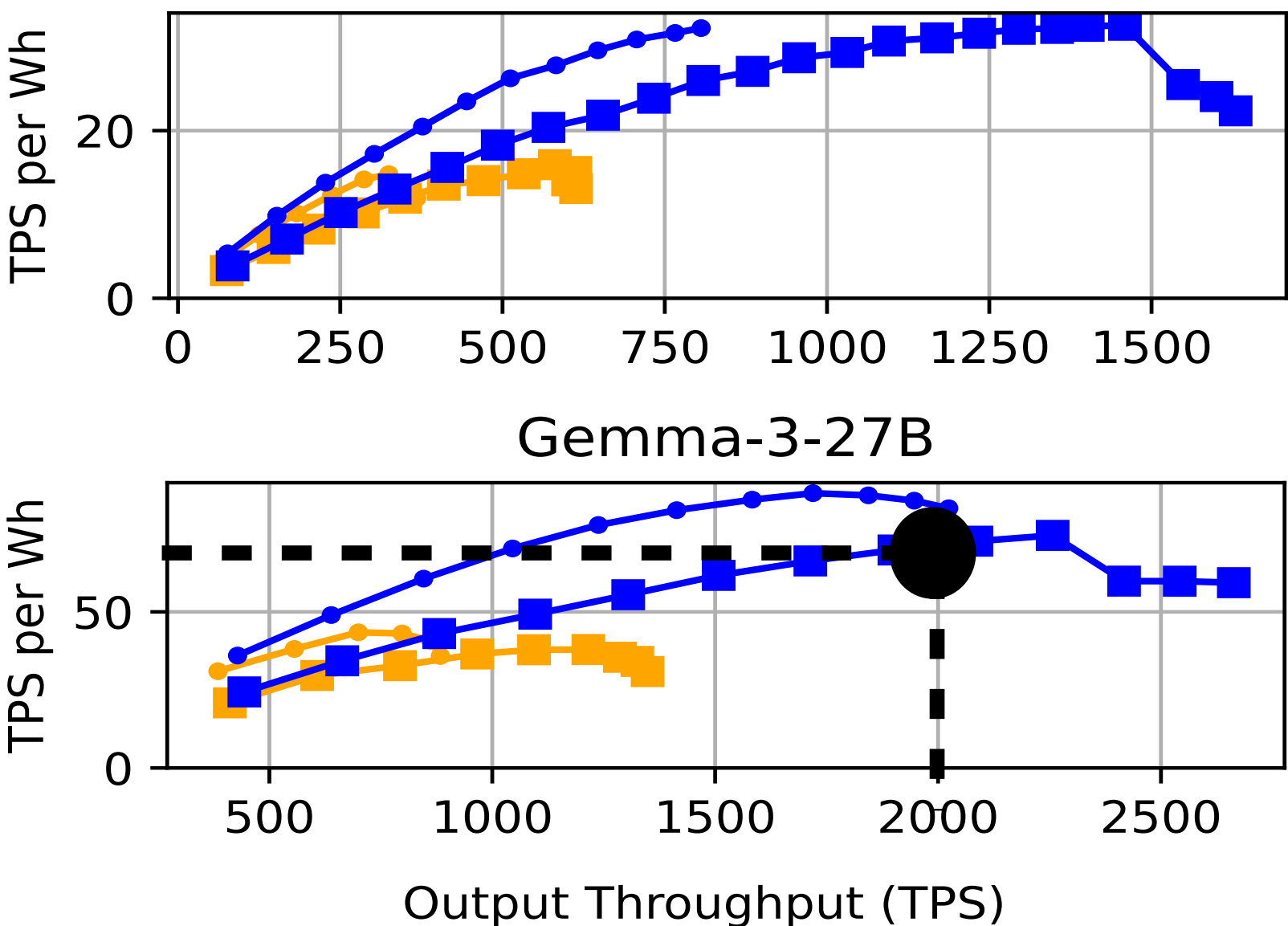


Token-Generation Energy Efficiency



NVLM-D-72B

Gemma-3-27B



Phase 2: Optimization + Execution

Configuration is a runtime decision, not a development decision

Optimizer

Mixed-Integer Linear Program

Objectives:

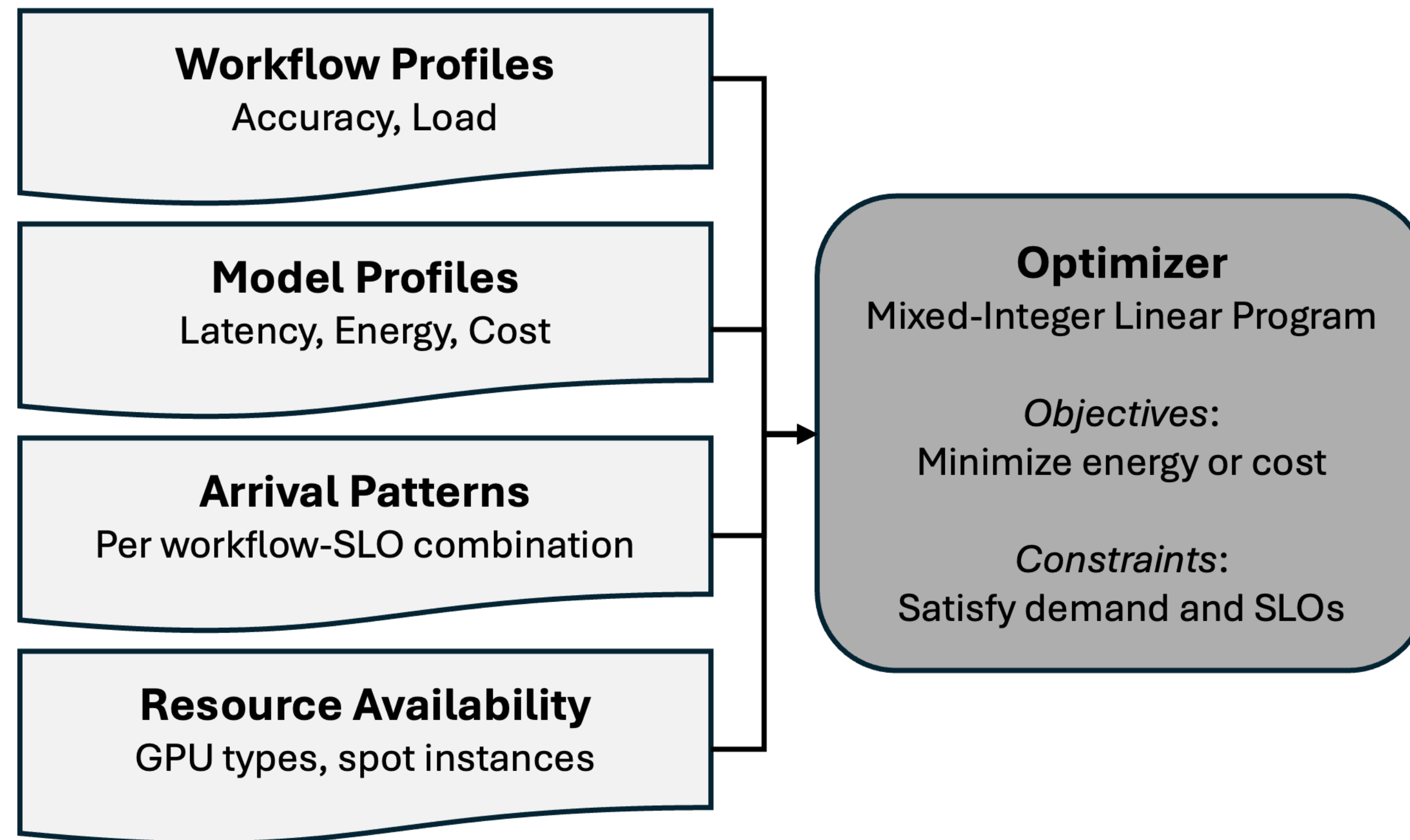
Minimize energy or cost

Constraints:

Satisfy demand and SLOs

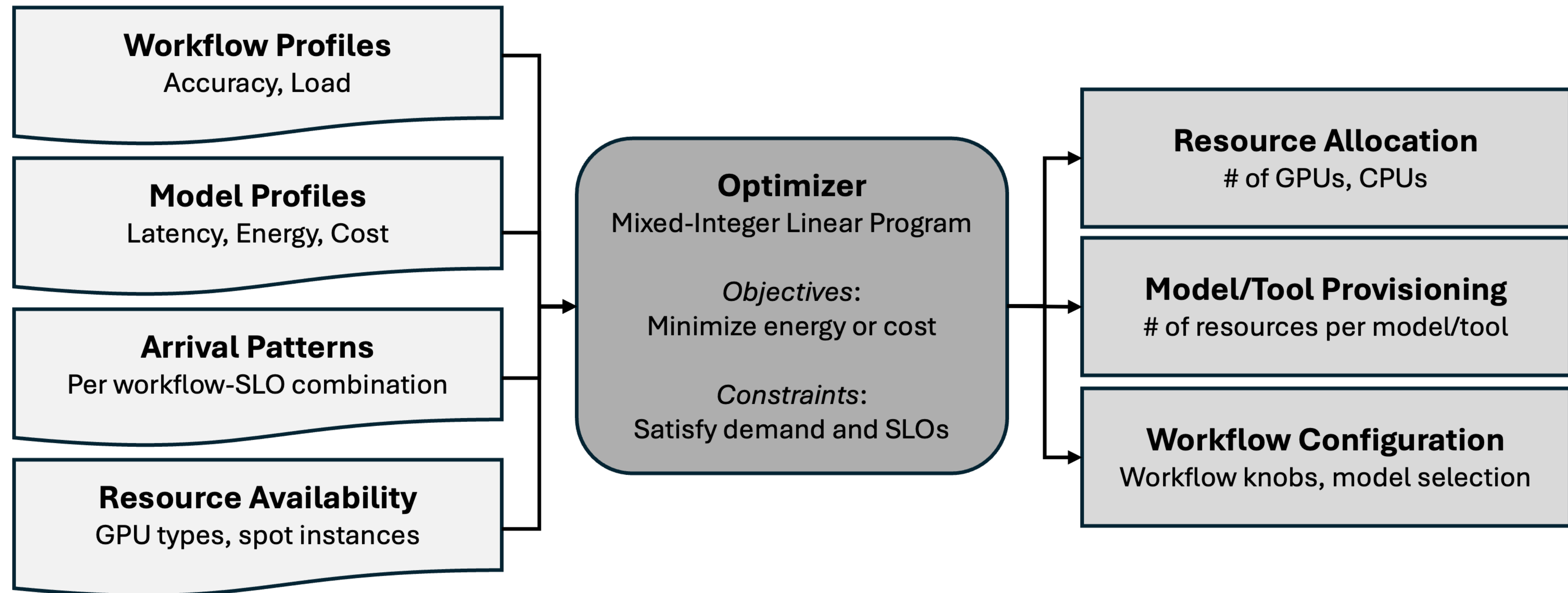
Phase 2: Optimization + Execution

Configuration is a runtime decision, not a development decision



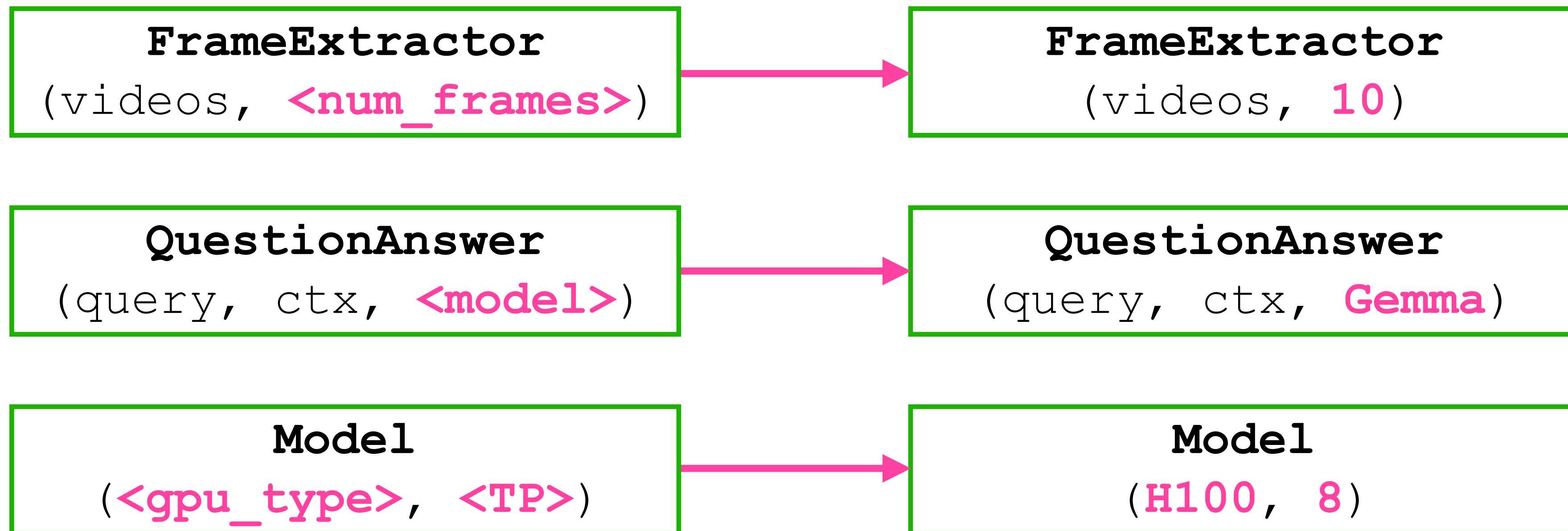
Phase 2: Optimization + Execution

Configuration is a runtime decision, not a development decision



Phase 2: Optimization + Execution

Plugging in the execution details



Phase 2: Optimization + Execution

Multi-workflow joint optimization

Phase 2: Optimization + Execution

Multi-workflow joint optimization

- Each **(workflow, SLO)-combination** has many feasible **configurations**

Phase 2: Optimization + Execution

Multi-workflow joint optimization

- Each **(workflow, SLO)-combination** has many feasible **configurations**
- Optimizer chooses **jointly**: steers different workflows towards **overlapping configurations**:
 - Same models, tools
 - Better, efficient hardware; amortizes per-user cost

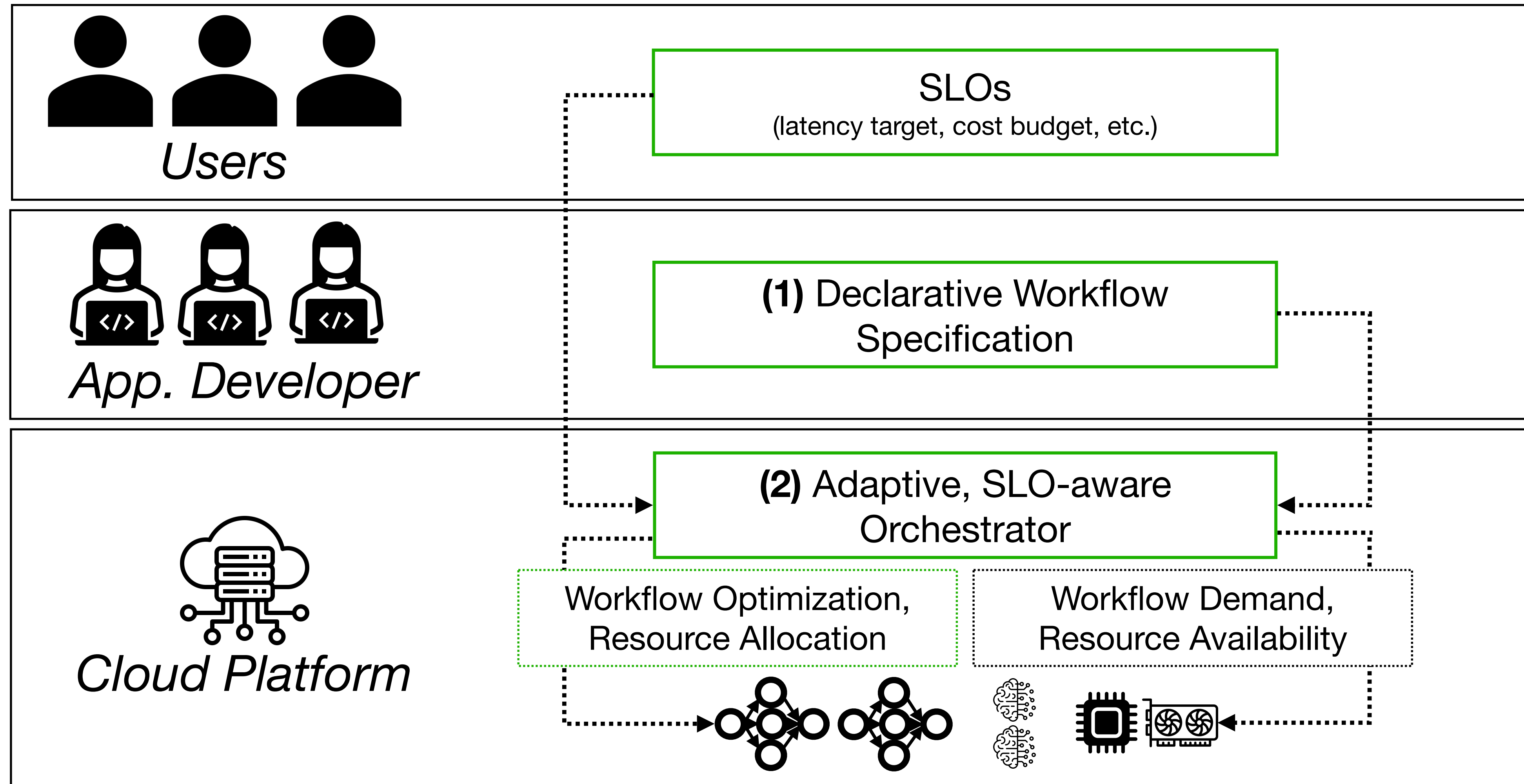
Phase 2: Optimization + Execution

Multi-workflow joint optimization

- Each **(workflow, SLO)-combination** has many feasible **configurations**
- Optimizer chooses **jointly**: steers different workflows towards **overlapping configurations**:
 - Same models, tools
 - Better, efficient hardware; amortizes per-user cost
- **Adapts** to dynamism in:
 - **Load** per (workflow, SLO)-combination
 - **Resource availability** (spot instances, energy sources etc.)

Murakkab

Cross-stack visibility + runtime orchestration across *different* workflows

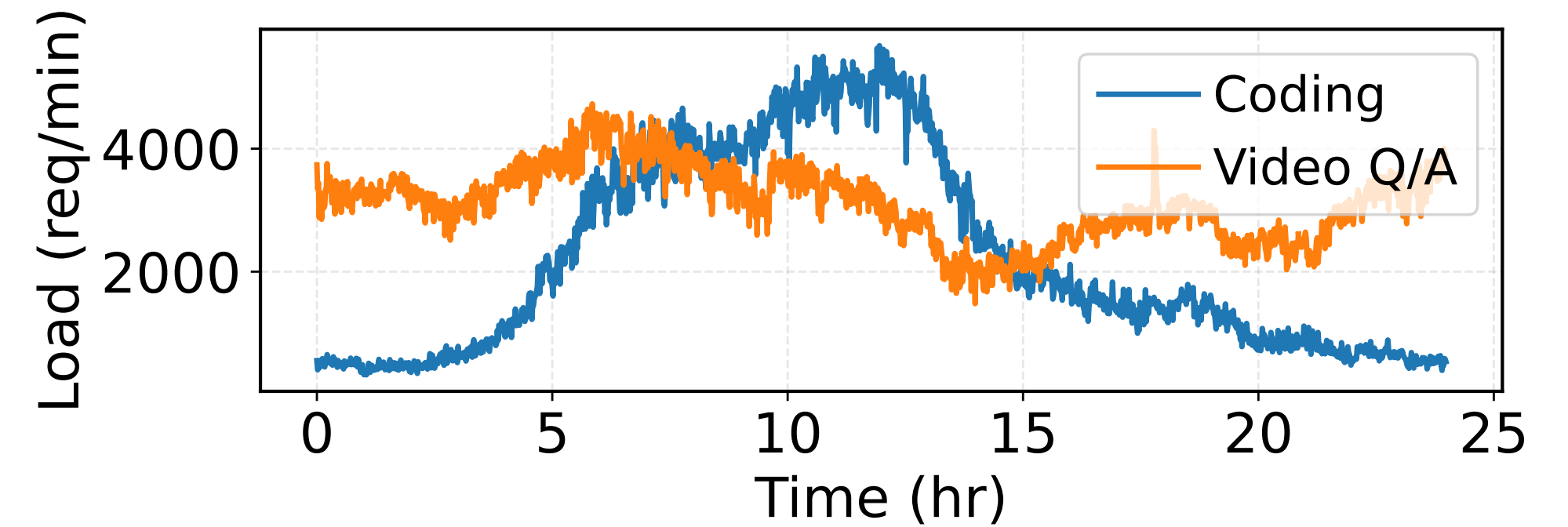


How much **resource-efficiency**
does Murakkab unlock?

Experiment Setup

Workload:

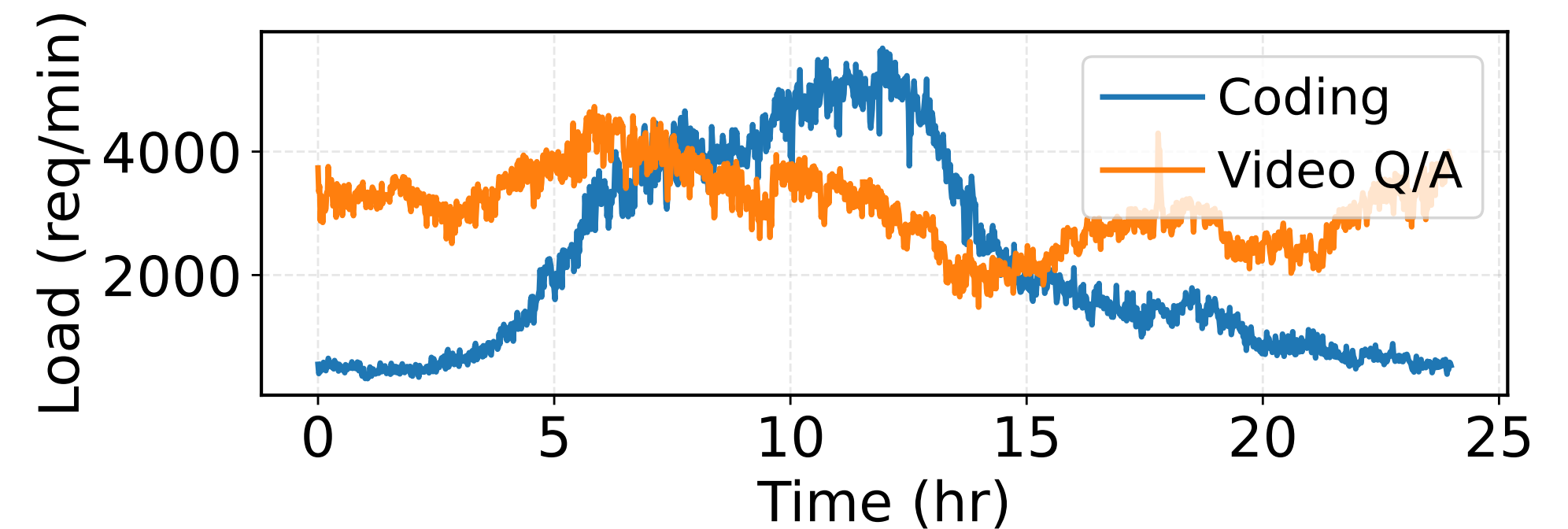
- 24 hour Azure Production LLM traces
- Video Q/A + Coding workflows
- SLO mix: 70% high-accuracy, 30% low-latency



Experiment Setup

Workload:

- 24 hour Azure Production LLM traces
- Video Q/A + Coding workflows
- SLO mix: 70% high-accuracy, 30% low-latency



Baselines:

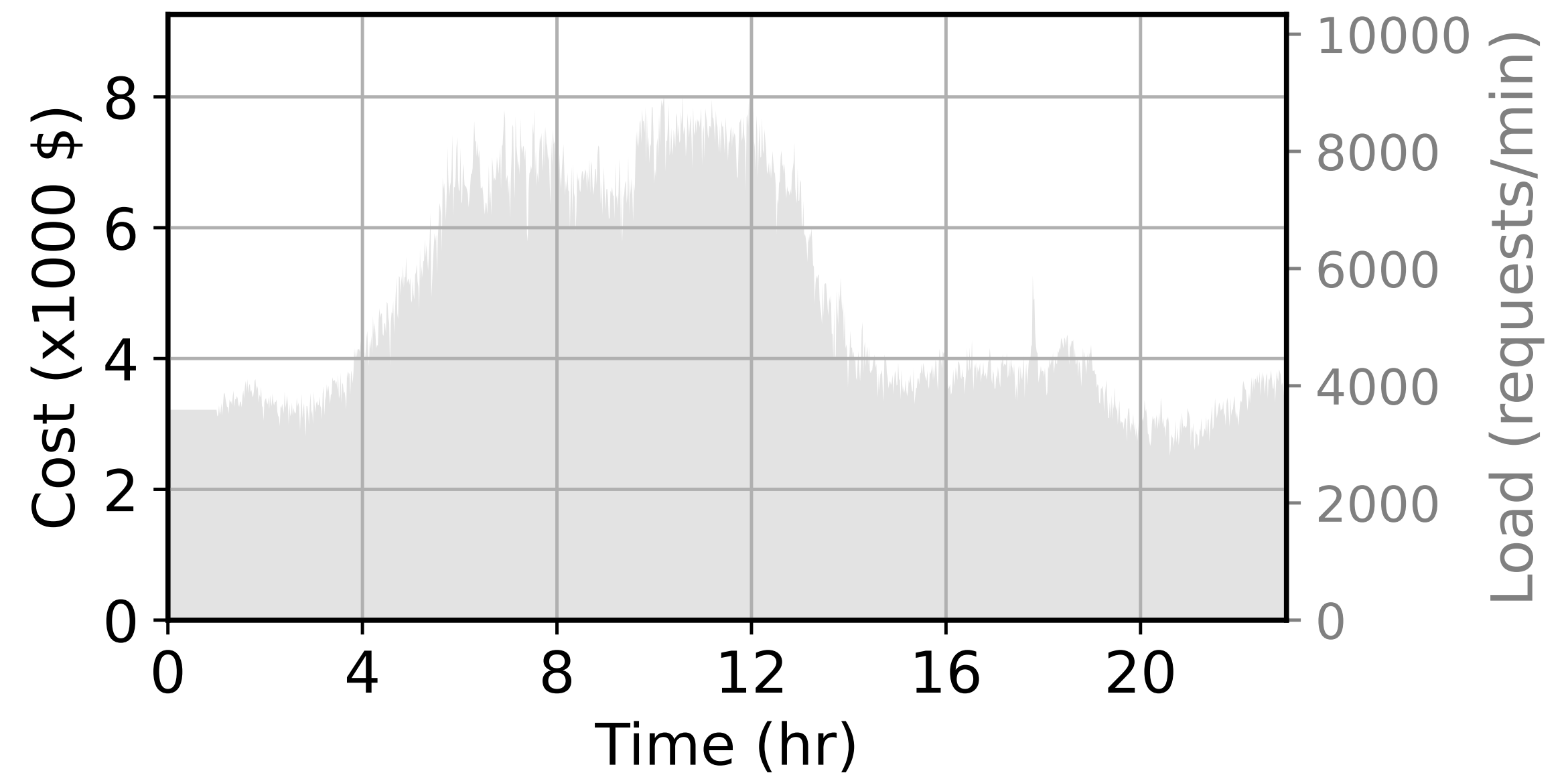
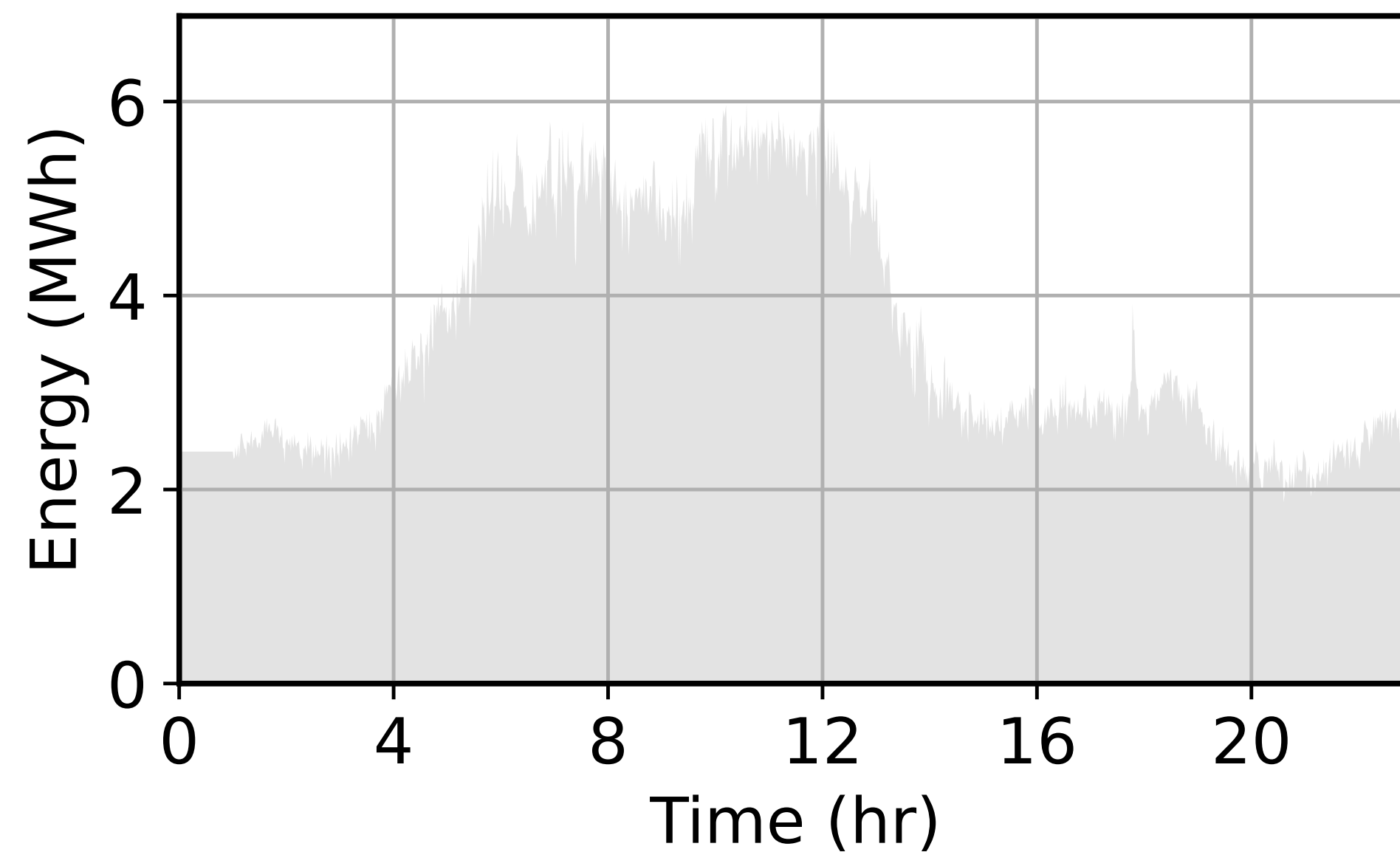
- LangGraph (*hand-tuned; balancing latency/accuracy; no workflow visibility, static allocation*)
- LangGraph + Auto-Scaling (*same auto-scaler as Murakkab, for fairness*)

Video Q/A + Coding

SLO mix: 70% high-accuracy, 30% low-latency

—•— LangGraph —◆— LangGraph + AutoScaling —■— Murakkab —▲— Murakkab + Multiplexing

■ Load (requests/min)

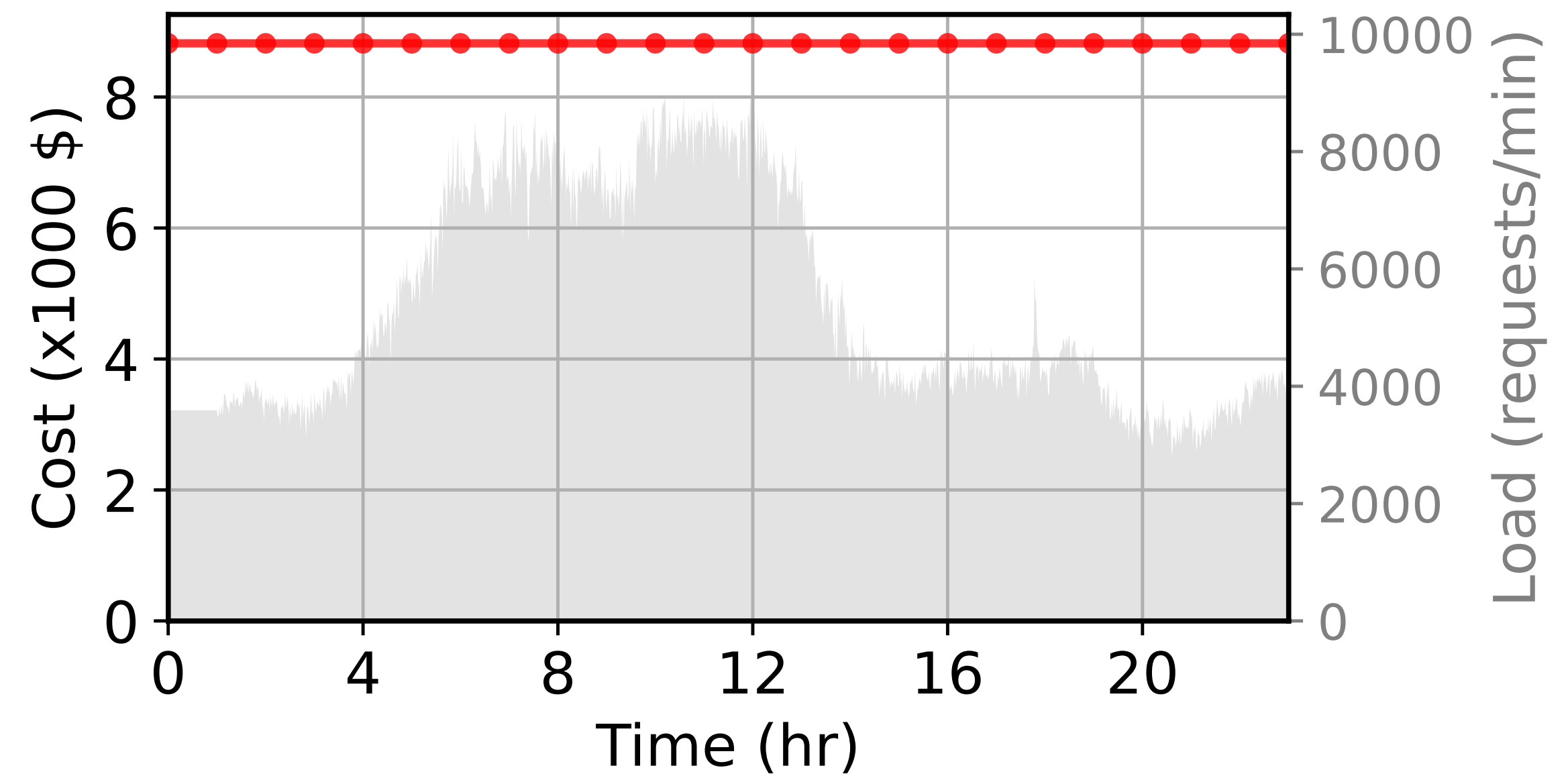
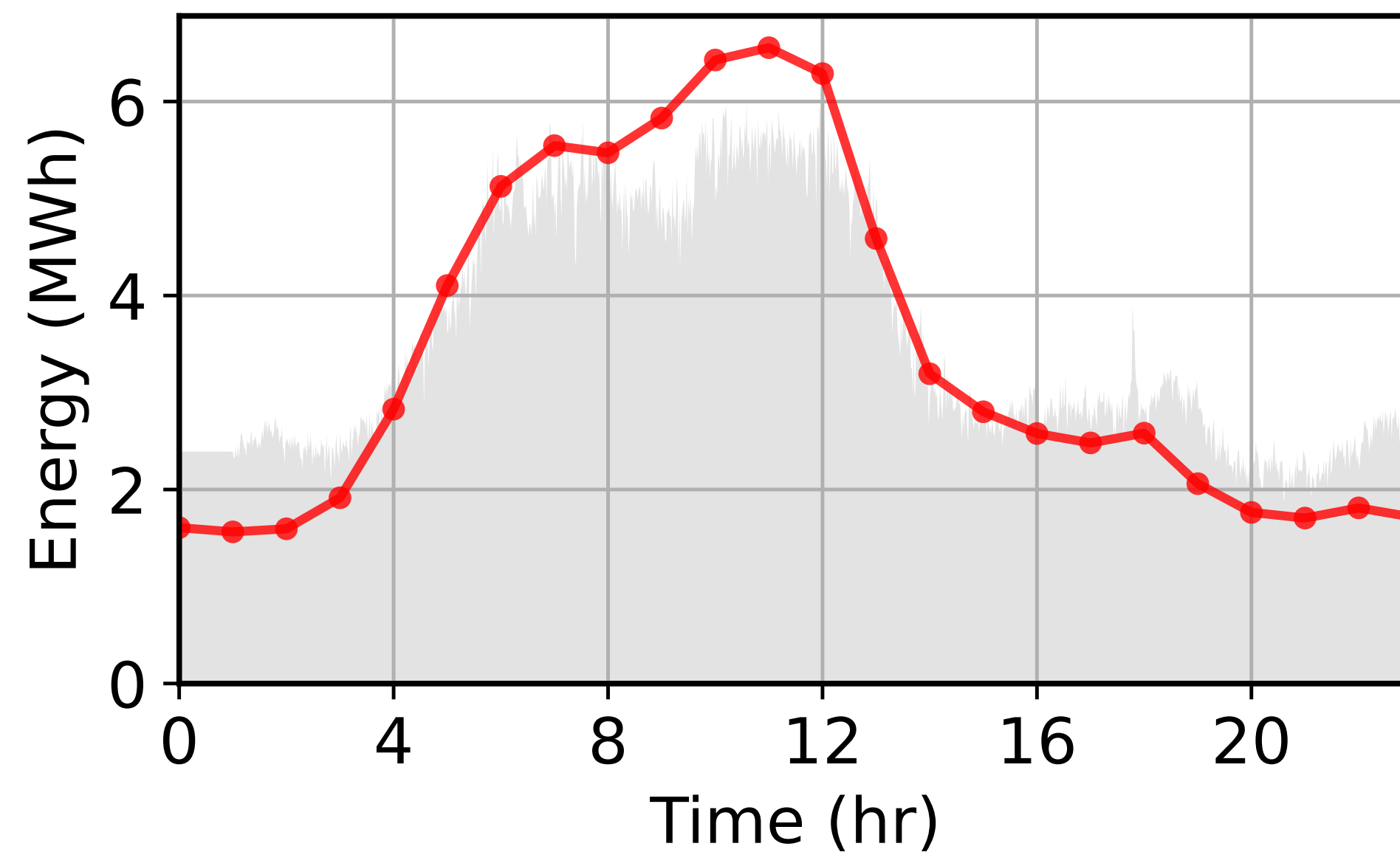


Video Q/A + Coding

SLO mix: 70% high-accuracy, 30% low-latency

LangGraph LangGraph + AutoScaling Murakkab Murakkab + Multiplexing

Load (requests/min)

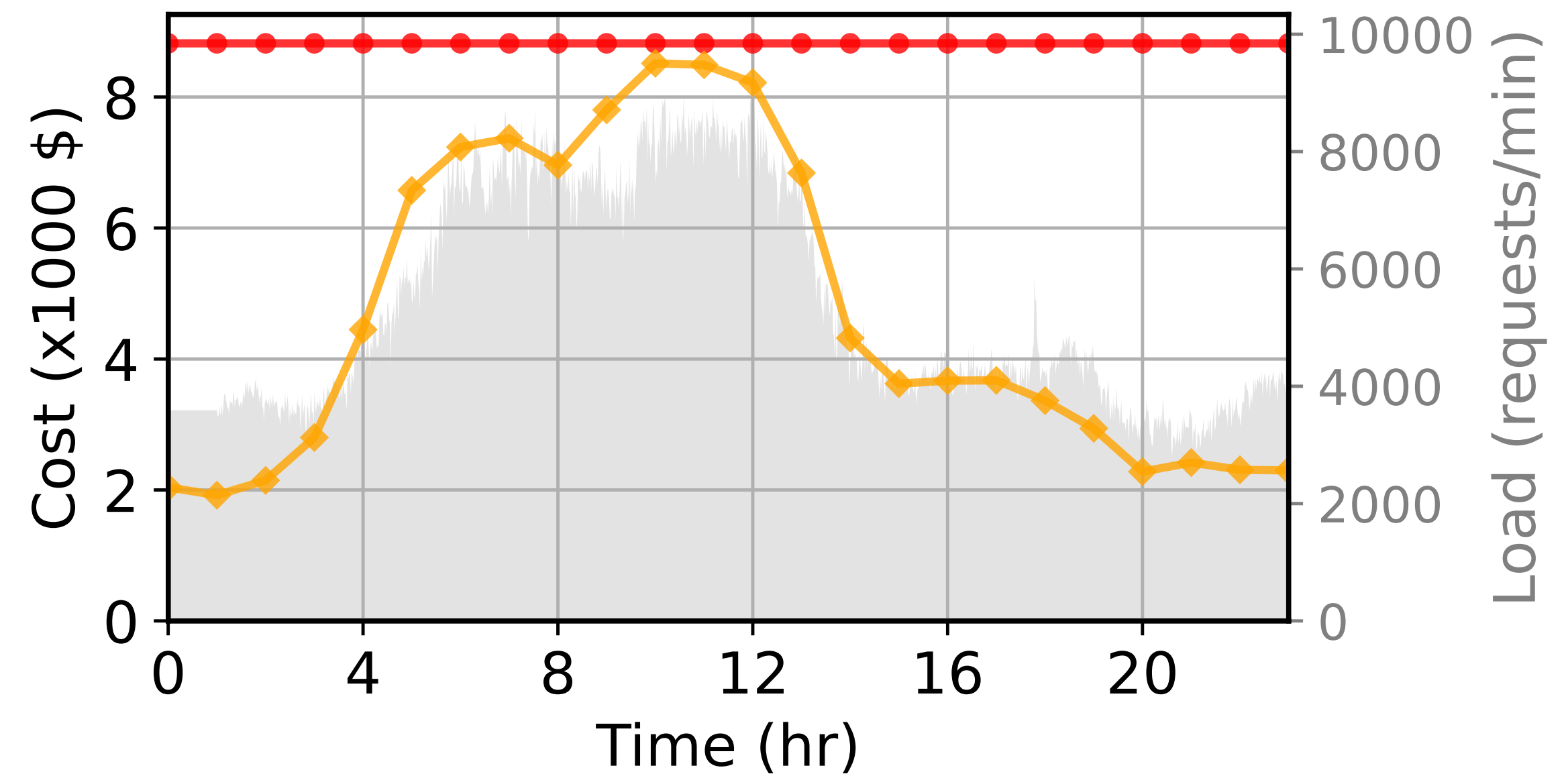
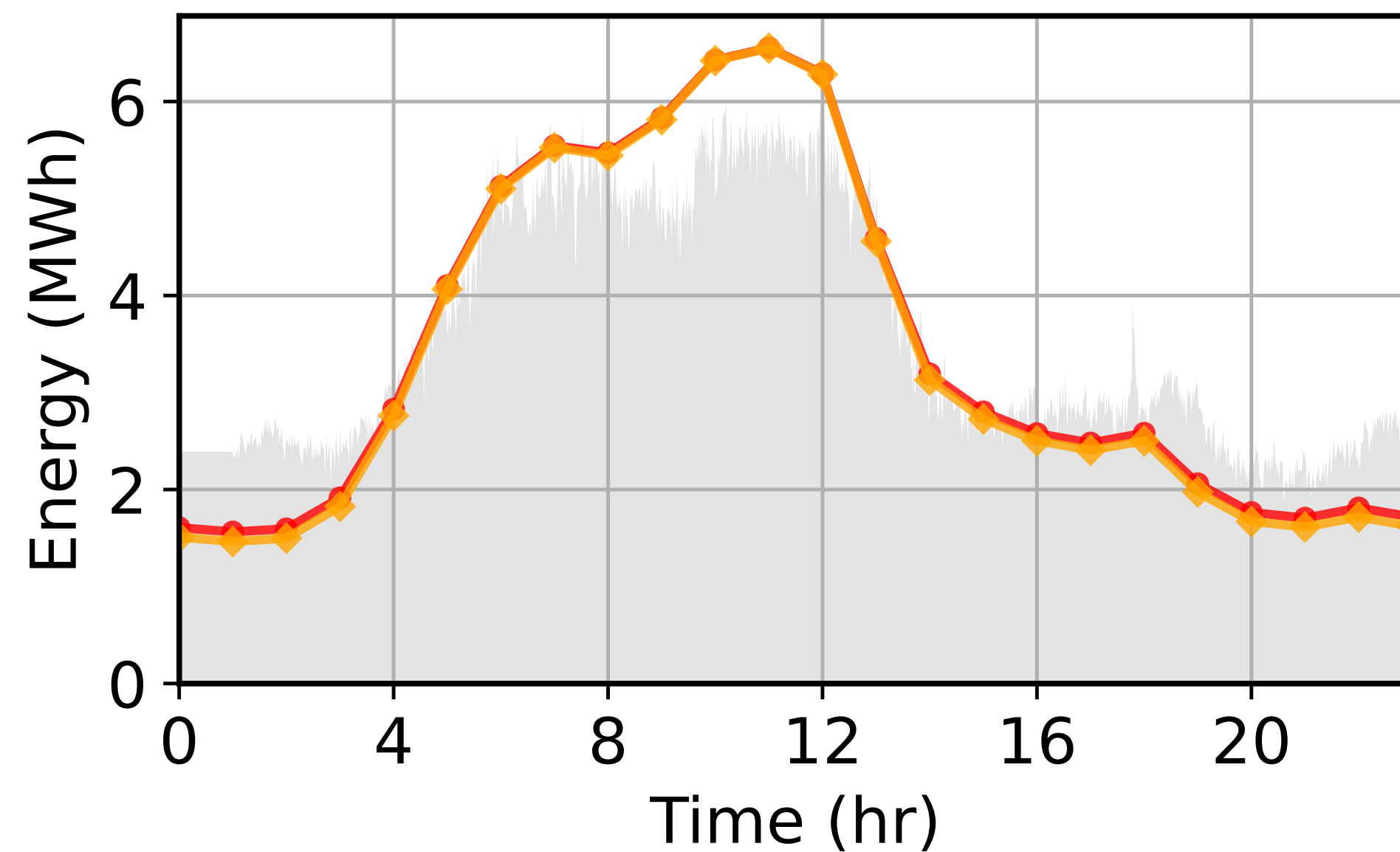


Video Q/A + Coding

SLO mix: 70% high-accuracy, 30% low-latency

LangGraph LangGraph + AutoScaling Murakkab Murakkab + Multiplexing

Load (requests/min)

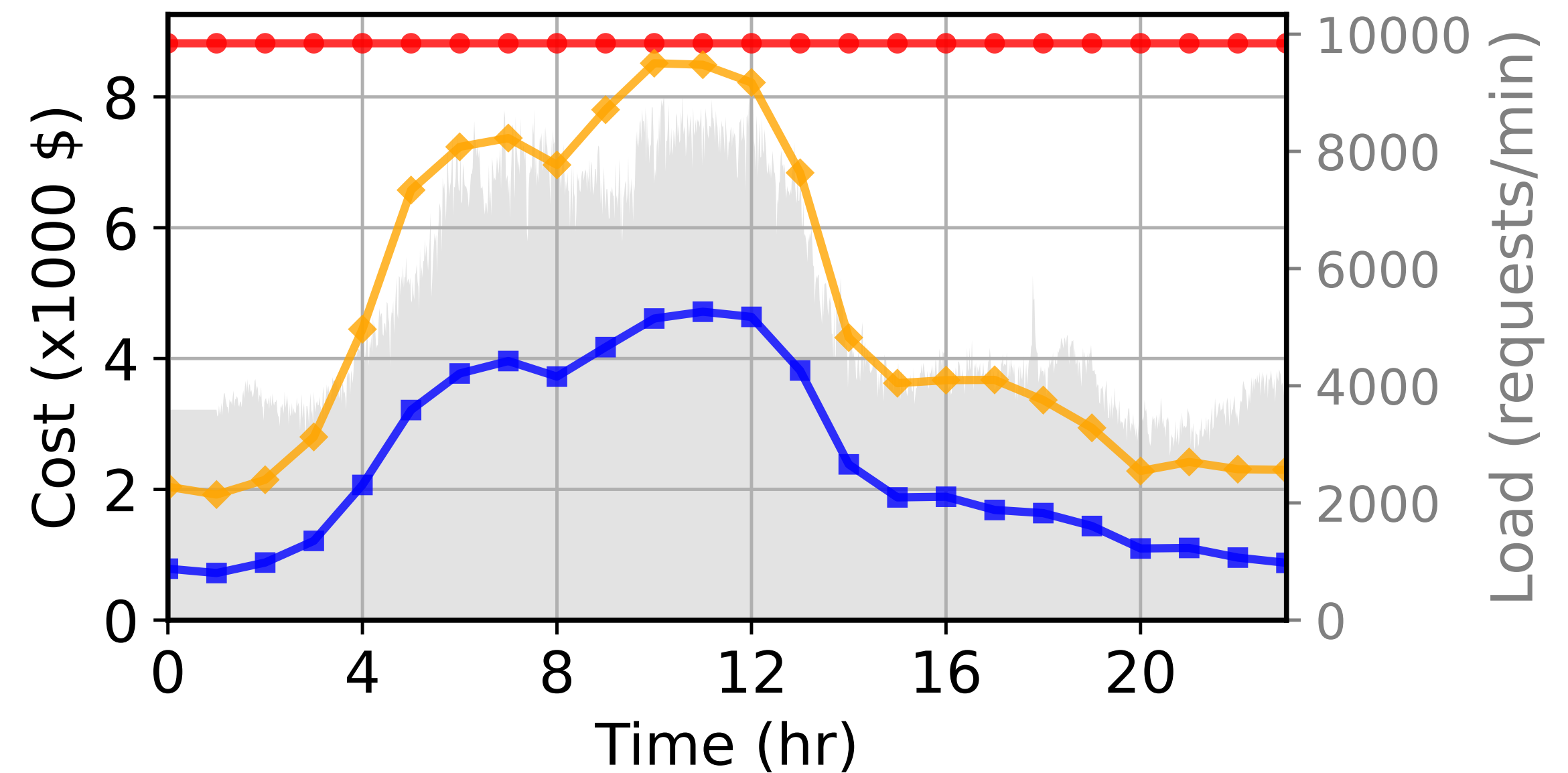
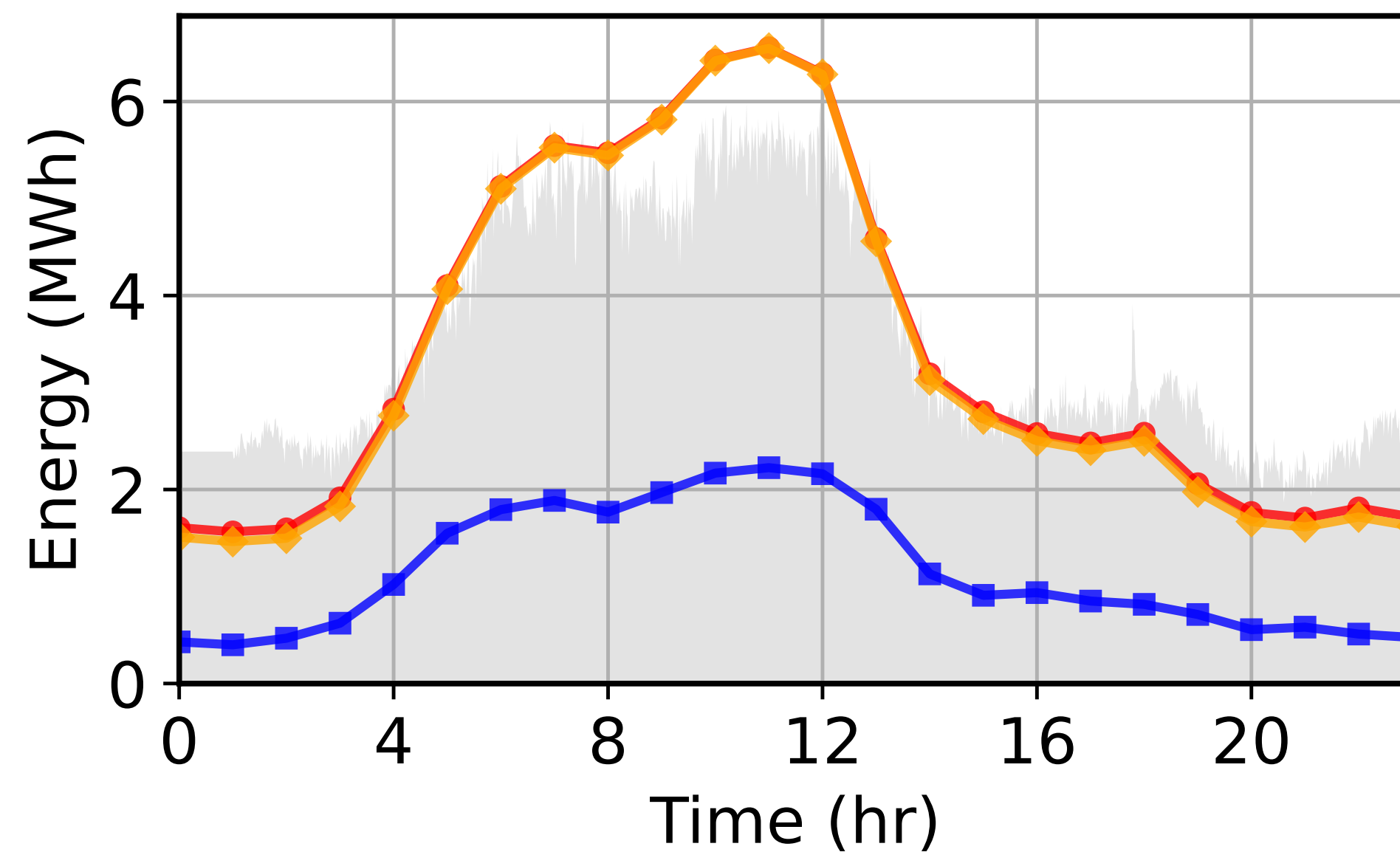


Video Q/A + Coding

SLO mix: 70% high-accuracy, 30% low-latency

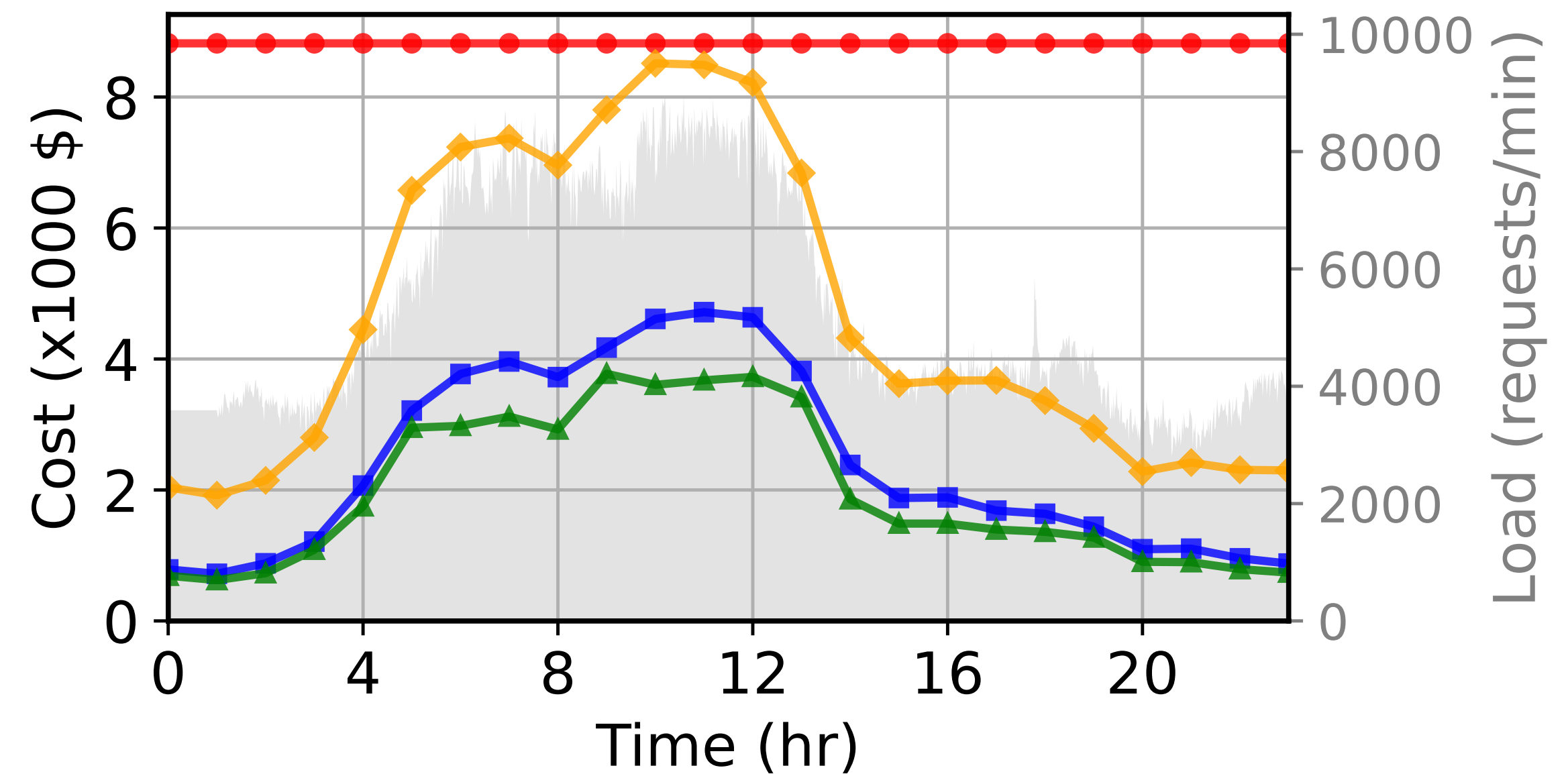
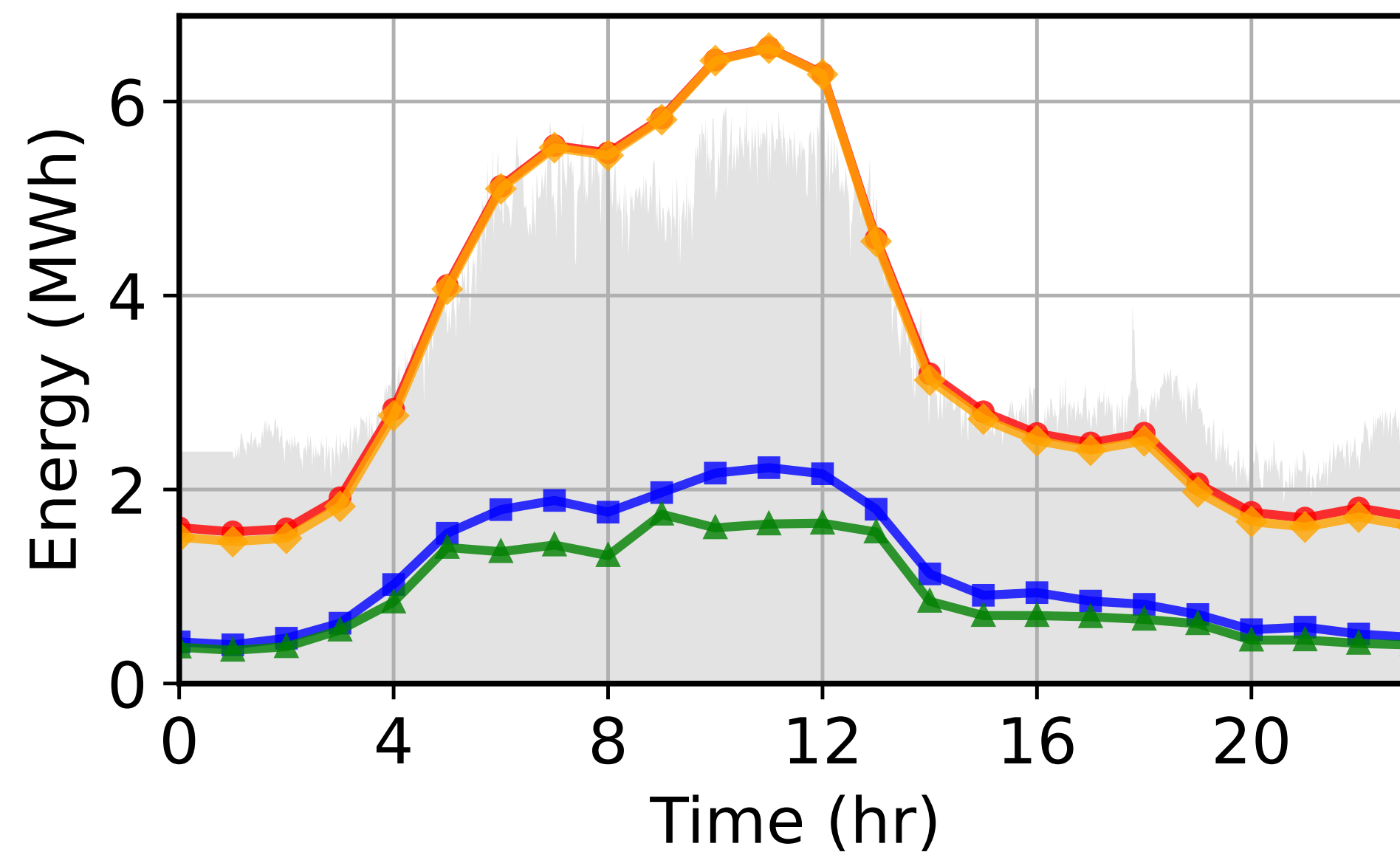
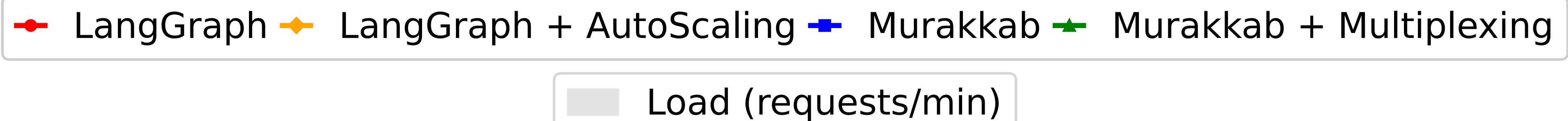
LangGraph LangGraph + AutoScaling Murakkab Murakkab + Multiplexing

Load (requests/min)



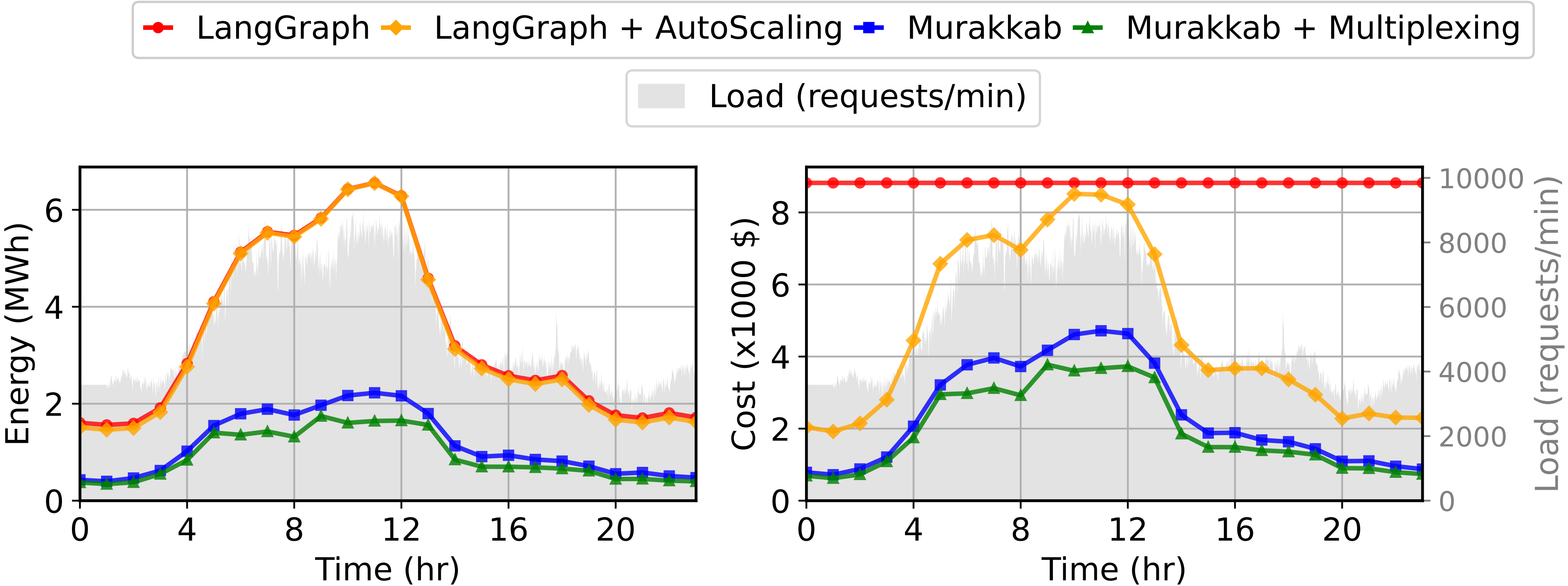
Video Q/A + Coding

SLO mix: 70% high-accuracy, 30% low-latency



Murakkab unlocks higher resource-efficiency

Up to 3.7× less energy, 4.3× lower cost, 2.8× fewer GPUs



Murakkab

Summary

- Today's agentic workflow execution is **siloed** and loses out on the opportunity for unified, end-to-end optimization across the stack.

Murakkab

Summary

- Today's agentic workflow execution is **siloed** and loses out on the opportunity for unified, end-to-end optimization across the stack.
- **Decoupling** workflow logic from execution details opens up a vast space of automated **runtime optimizations** while satisfying SLOs.

Murakkab

Summary

- Today's agentic workflow execution is **siloed** and loses out on the opportunity for unified, end-to-end optimization across the stack.
- **Decoupling** workflow logic from execution details opens up a vast space of automated **runtime optimizations** while satisfying SLOs.
- Visibility into all layers of the stack allows **joint optimization** across different **multi-tenant workflows** in cloud platforms for greater **resource-efficiency**.