

Rethinking LLM Serving From the Application's Perspective

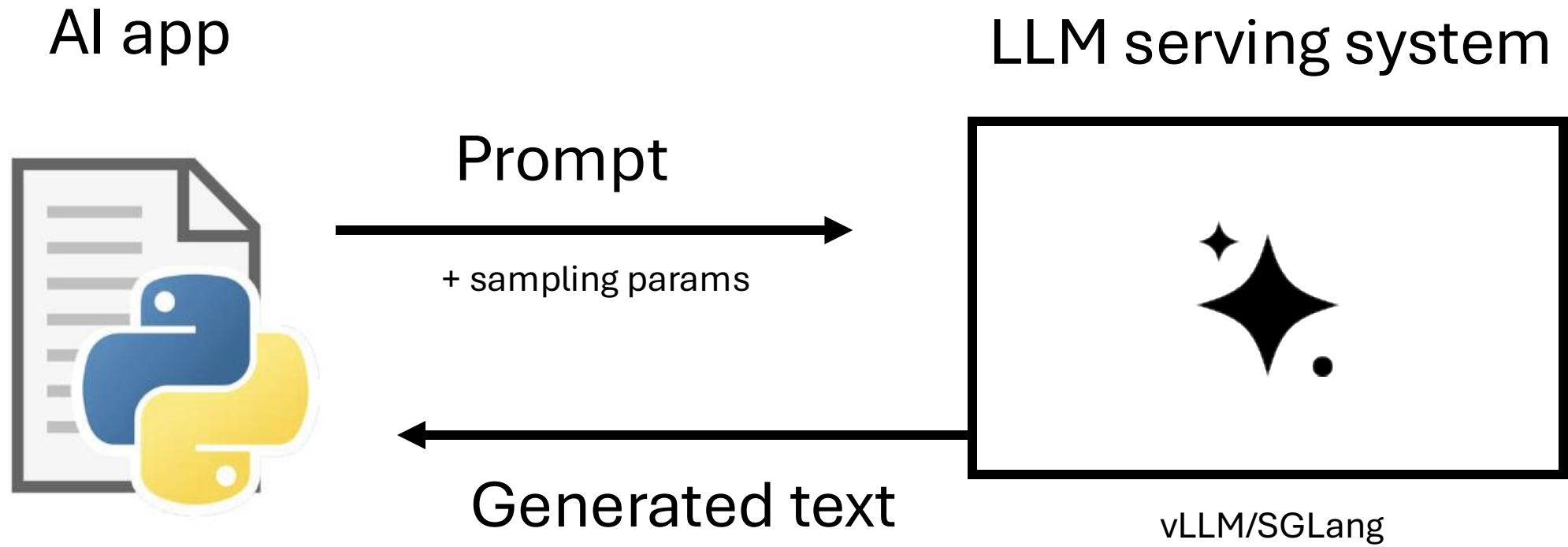
*TLDR; future LLM serving systems should serve **programs**, not **prompts***

In Gim
Yale University

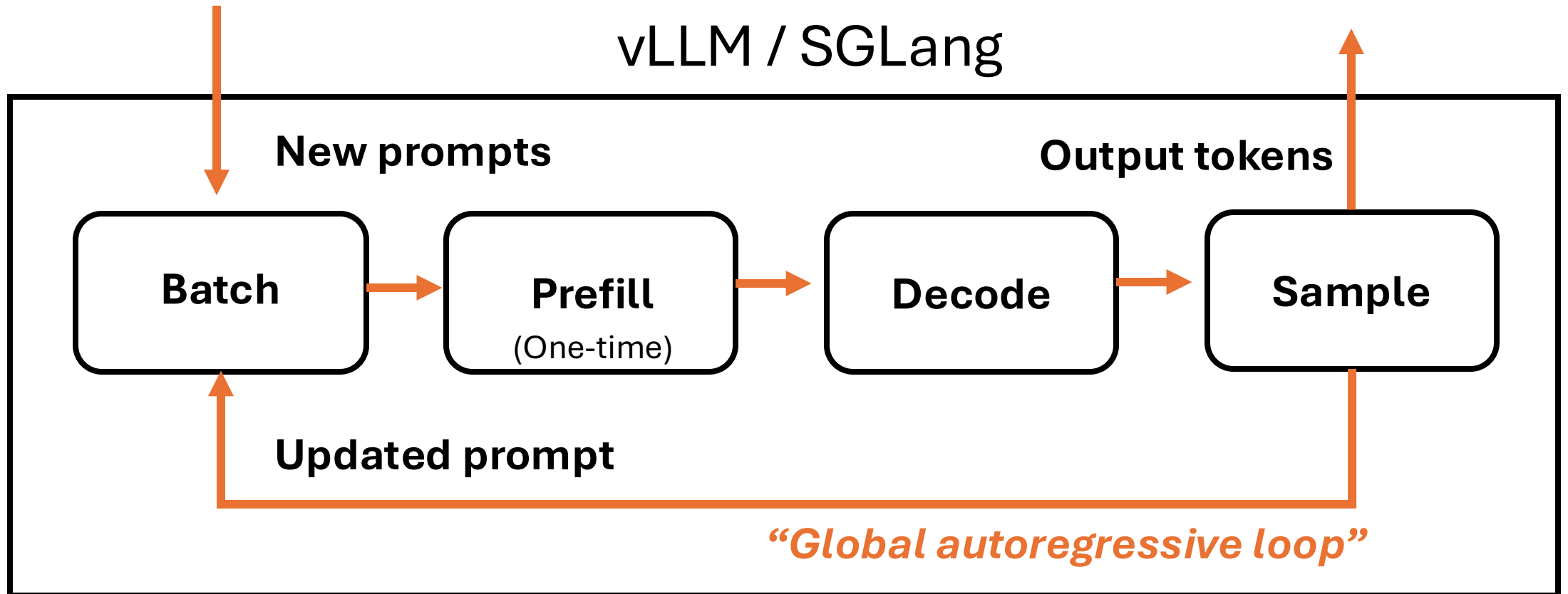
The only API you need to build a billion-dollar AI startup

```
curl https://api.openai.com/v1/responses \  
-d "Write a response to this customer complaint"
```

Today's LLM serving systems are designed for text completion

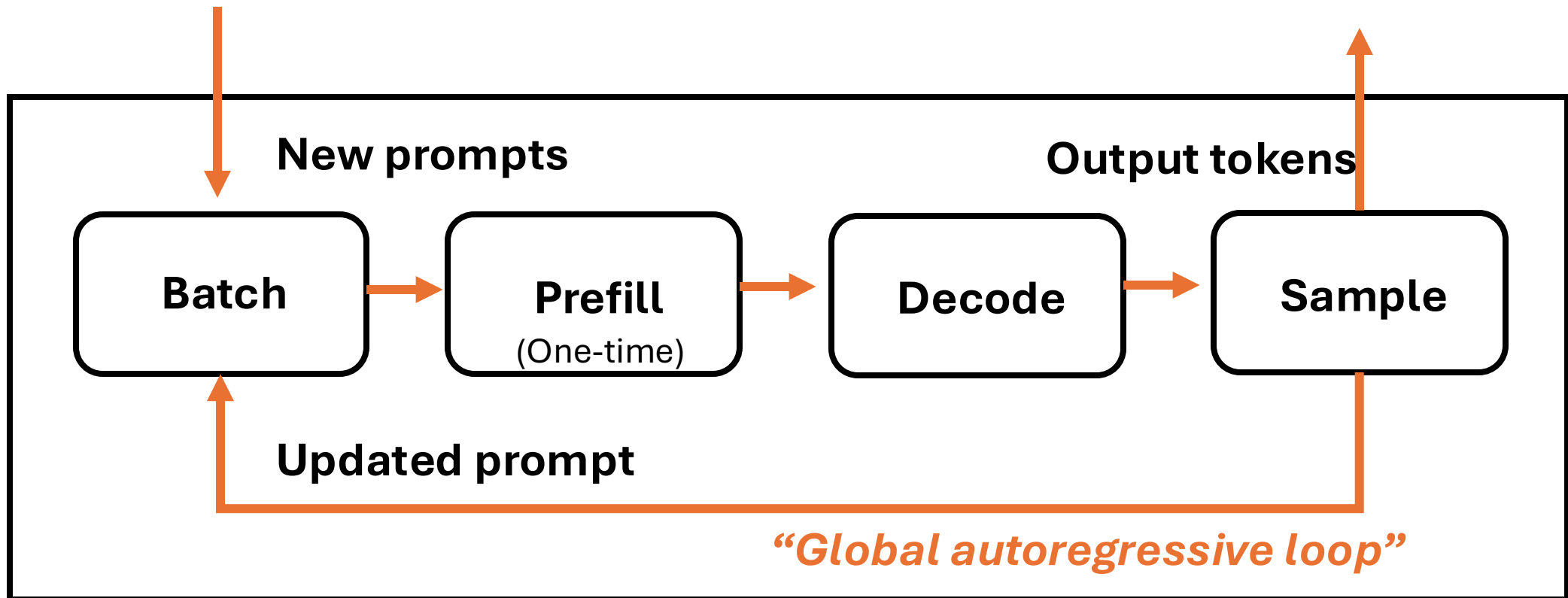


Today's LLM serving systems are designed for text completion



Three implicit design choices behind “prompt serving”:

- KV cache is transient
- Generation is closed-loop
- Serving is stateless



“Prompt serving” is too inflexible for emerging LLM applications

1. Inference
Inefficiency

2. Integration
Friction

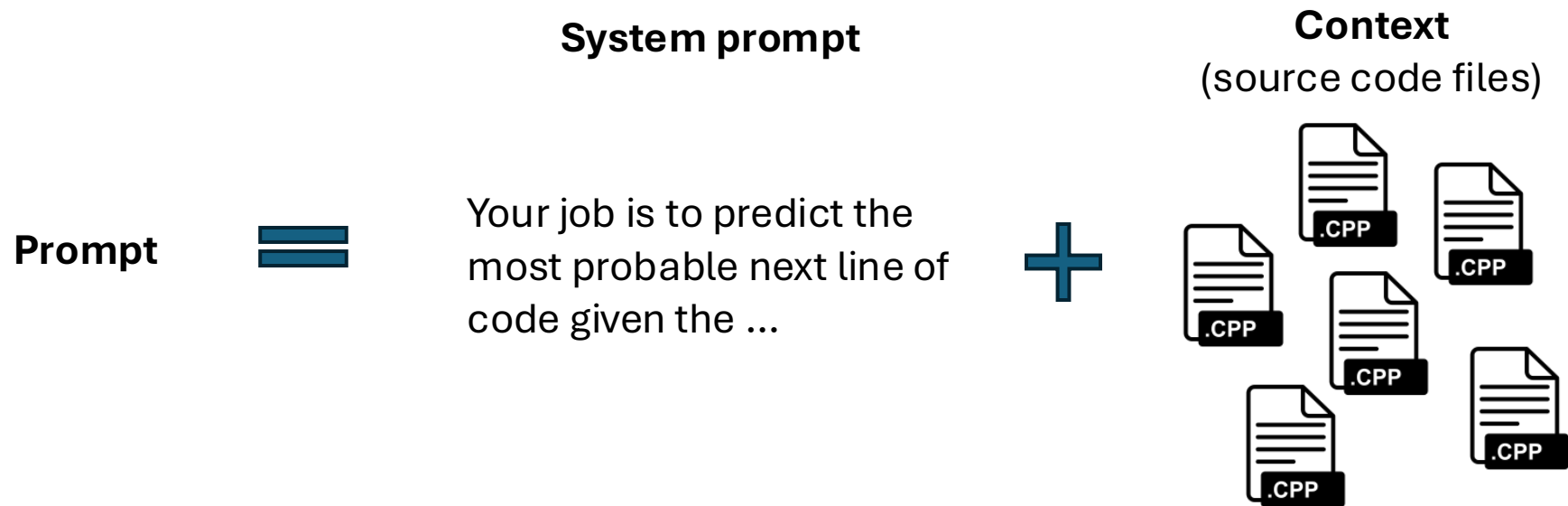
3. Implementation
Challenge

Imagine building an AI code completion tool

```
1  from __future__ import print_function
2  import argparse
3  import torch
4  import torch.nn as nn
5  import torch.optim as optim
6  import numpy as np
7  import matplotlib
8  matplotlib.use('Agg')
9  import matplotlib.pyplot as plt
10
11 class Sequence(nn.Module):
12     def __init__(self):
13         super(Sequence, self).__init__()
14         self.lstm1 = nn.LSTMCell(1, 51)
15         self.lstm2 = nn.LSTMCell(51, 51)
16         self.linear = nn.Linear(51, 1)
17
18     def forward(self, input, future = 0):
19         outputs = []
20         h_t = torch.zeros(input.size(0), 51, dtype=torch.double)
```

Challenge 1. Inference efficiency

Scenario: User's input updates the suggestion.

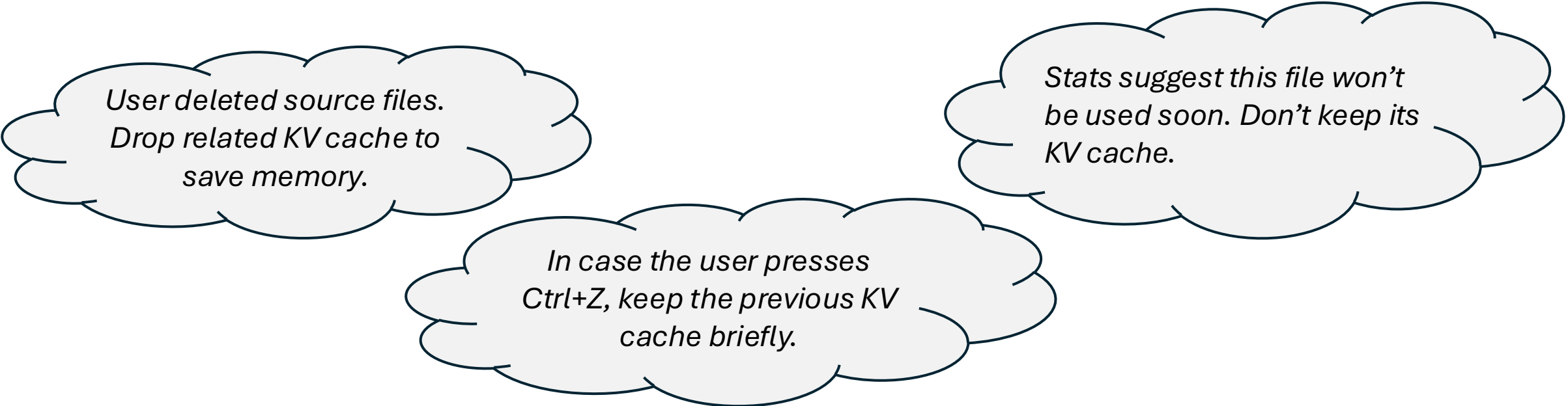


Due to the large overlaps in prompts,
Inter-request KV cache reuse is essential for efficiency

Challenge 1. Inference inefficiency

KV cache management is often automatic, system-wide policy
(e.g., OpenAI's prompt caching, vLLM's automatic prefix caching)

This precludes application-driven optimizations



*User deleted source files.
Drop related KV cache to
save memory.*

*In case the user presses
Ctrl+Z, keep the previous KV
cache briefly.*

*Stats suggest this file won't
be used soon. Don't keep its
KV cache.*

Challenge 2. Integration friction

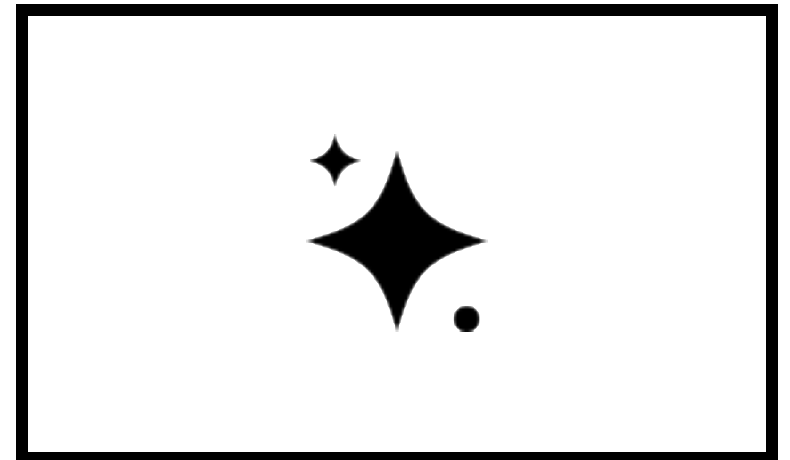
Scenario: Augment generation with relevant API documentation



Prompt: “If unsure about an API during generation, output `<read>API name</read>` and pause”



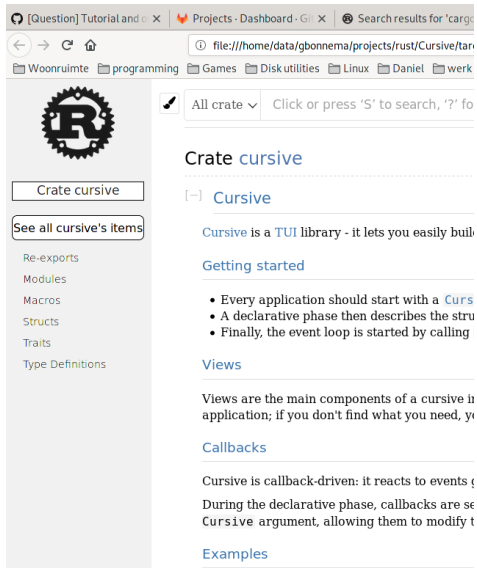
LLM serving system



Challenge 2. Integration friction

Each external interaction induces two extra boundary crossings

Data/Tool



(3) Retrieve data

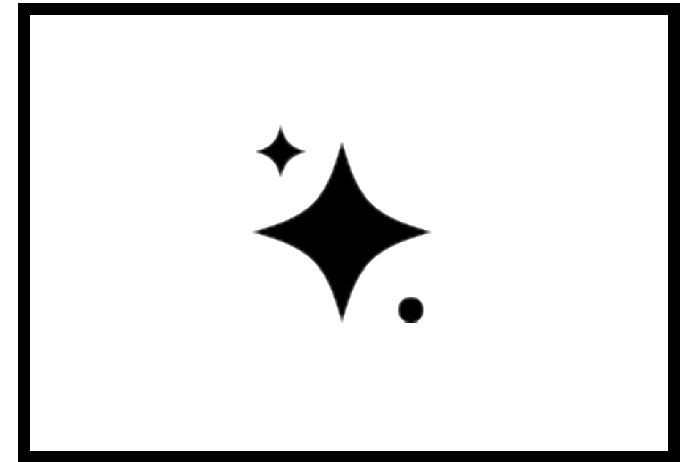


(1) Original prompt

(2) <read> ...

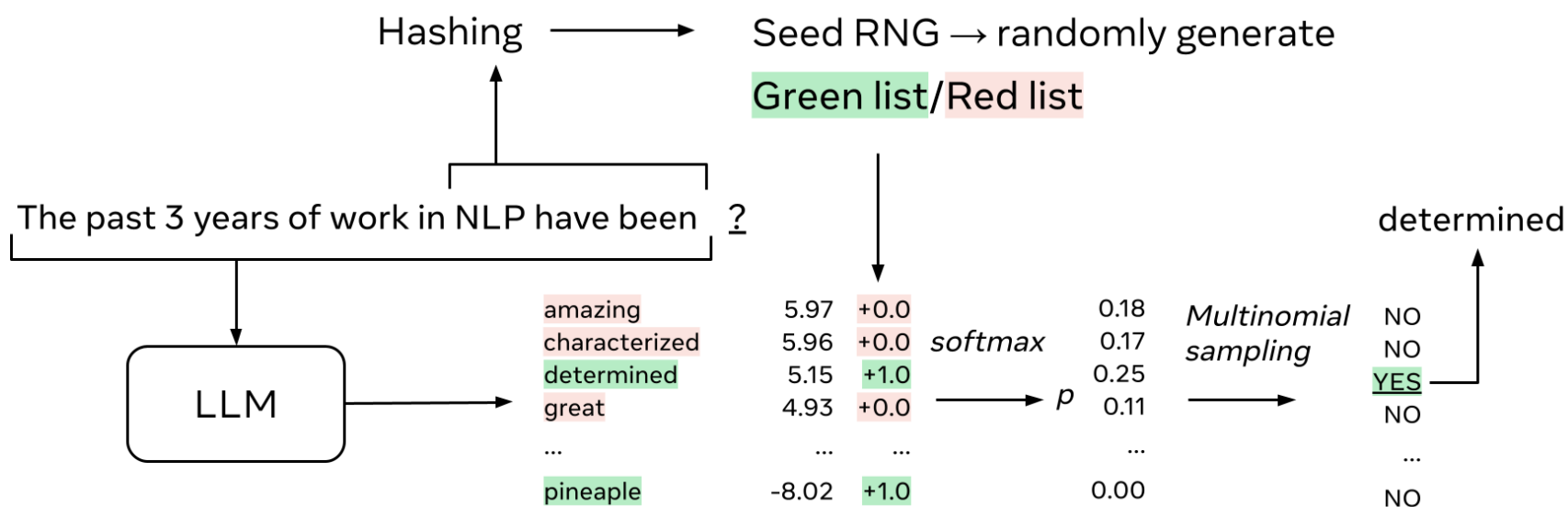
(4) updated prompt

LLM serving system



Challenge 3. Implementation

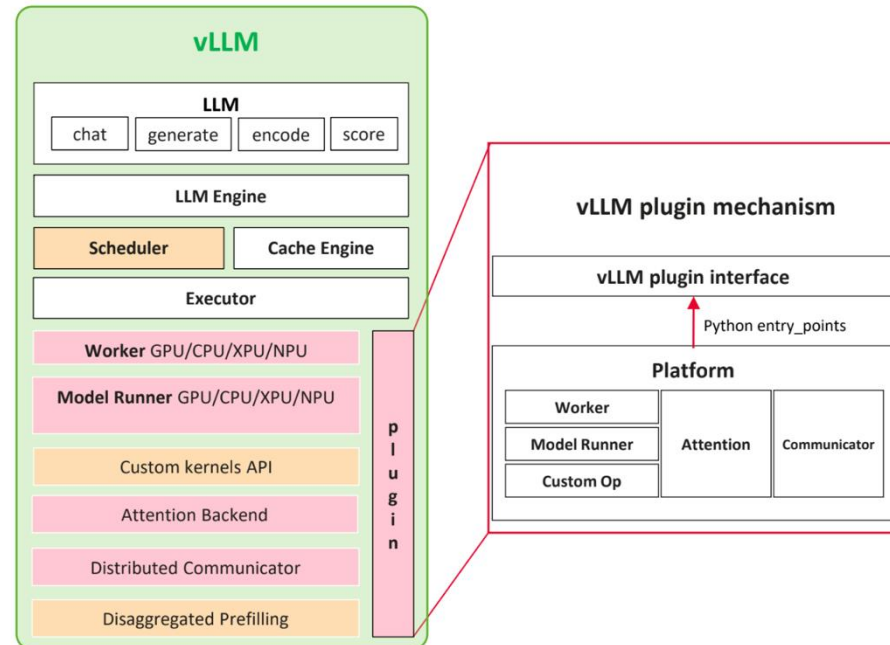
Scenario: Watermark output suggestions for IP protection



Challenge 3. Implementation

customizing existing serving systems?

vLLM plugin system



Challenge 3. Implementation

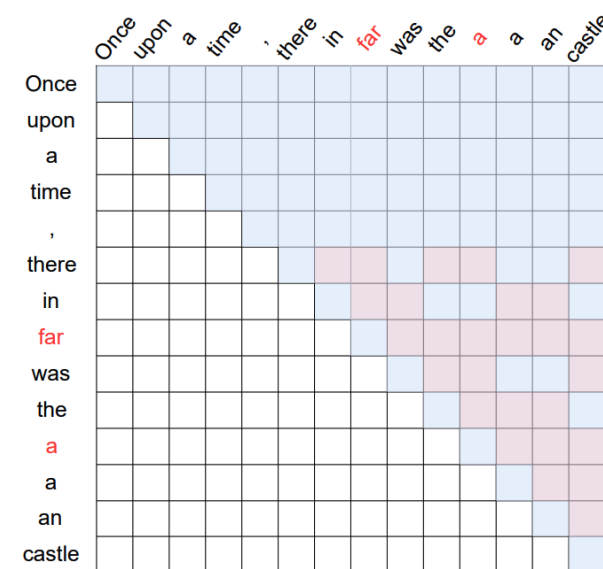
It imposes a non-trivial engineering cost due to:

- (1) Monolithic generation pipeline
- (2) Per-request policy control
- (3) Lack of extensible interface

Challenge 3. Implementation

Monolithic token generation loops make it hard to implement custom decoding strategies

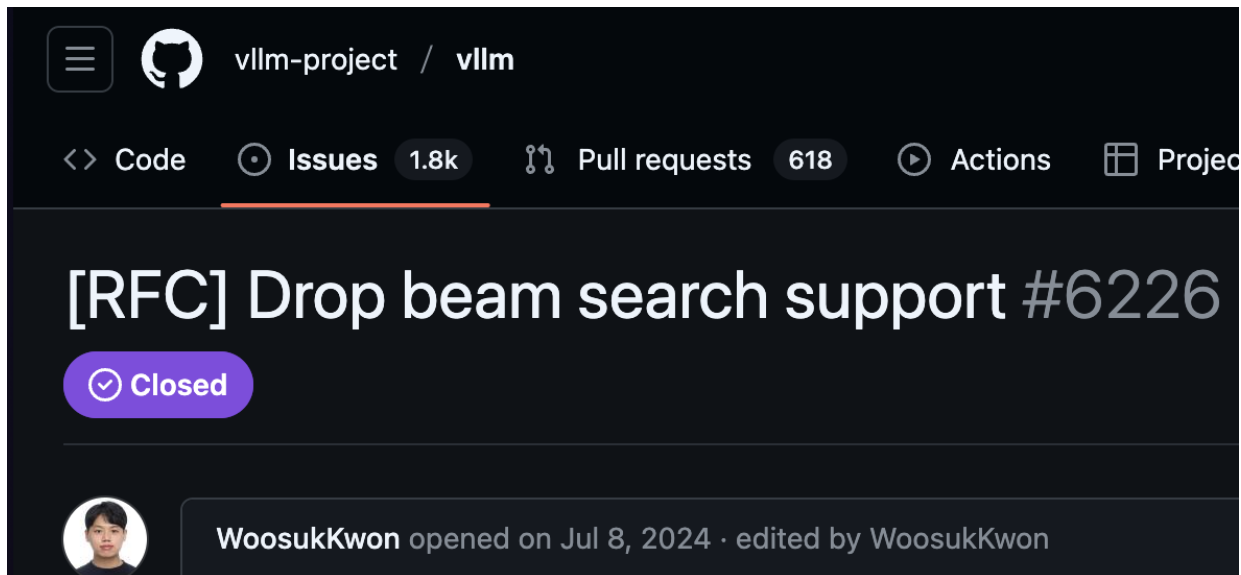
e.g., beam search



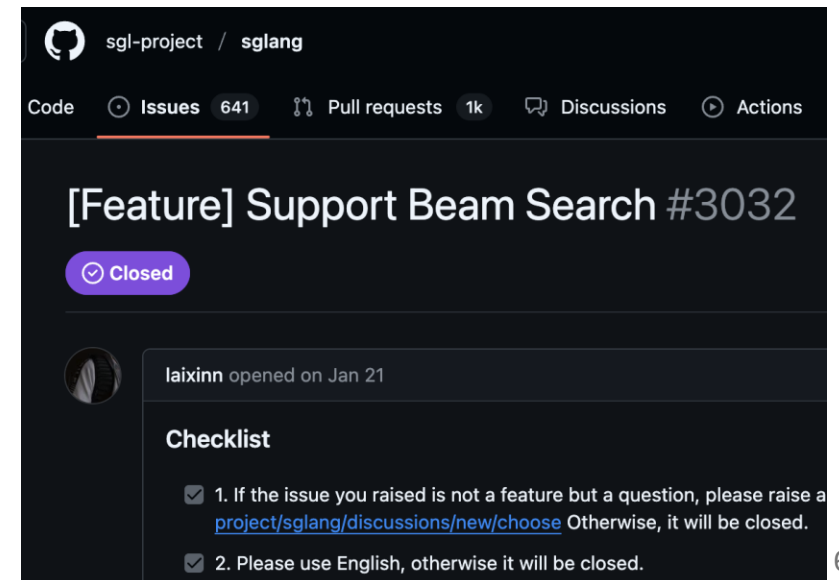
Challenge 3. Implementation

Monolithic execution loops entangle model logic with generation control, complicating non-standard decoding methods.

Beam search is a headache (vllm#6226)



Beam search still unsupported (sglang#3032)



Challenge 3. Implementation

No per-request policy control.
Optimizations are global policies

```
[--enable-chunked-prefix-l1
[--enable-prefix-caching
[--disable-sliding-window
[--disable-cascade-attr
[--skip-tokenizer-init
[--enable-prompt-embeds
[--speculative-config SPECULATIVE_G]
```

[illegible]

The only API you need to build a billion-dollar AI startup?

```
curl https://api.openai.com/v1/responses \  
-d “Write a response to this customer complaint”
```

We need a serving system that is more **application-aware**.

Vision: “programmable” LLM serving system

We have seen similar problems/solutions in networking, OS, and security before:

- eBPF
- SDN
- OPA



Can we take the same philosophy in building the LLM serving system?

Programmable LLM serving with Pie (SOSP '25)

<https://pie-project.org>

Design goals of Pie

1. Inference
Inefficiency



App-defined
KV cache mgmt

2. Integration
Friction



Integrated
compute and I/O

3. Implementation
Challenge



Customizable
at forward pass level

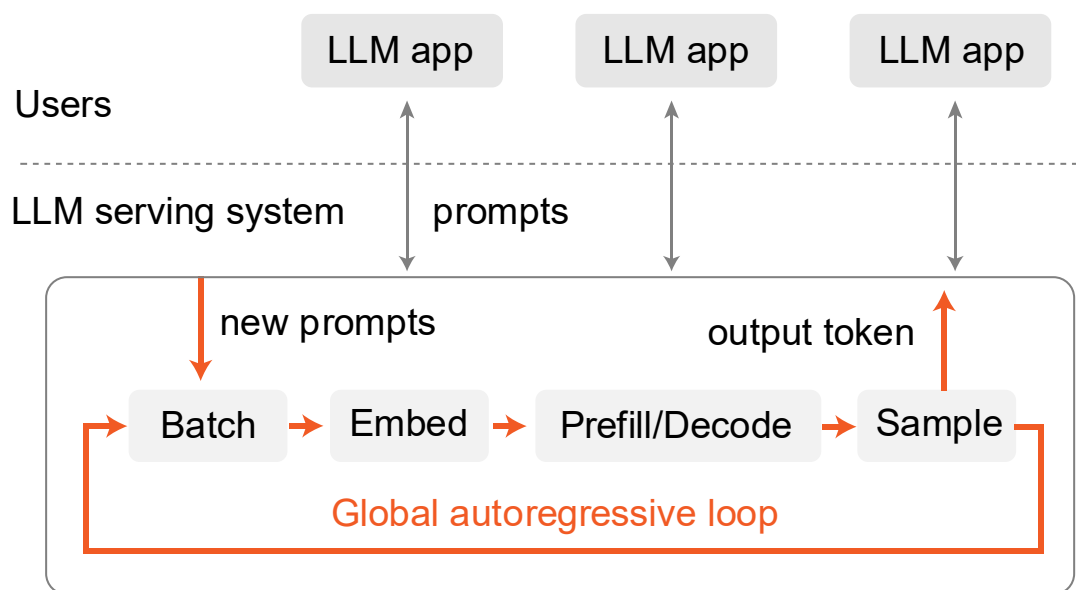
Design principles

Decouple control
from token generation

Delegate the control
to user programs
called "inferlets"

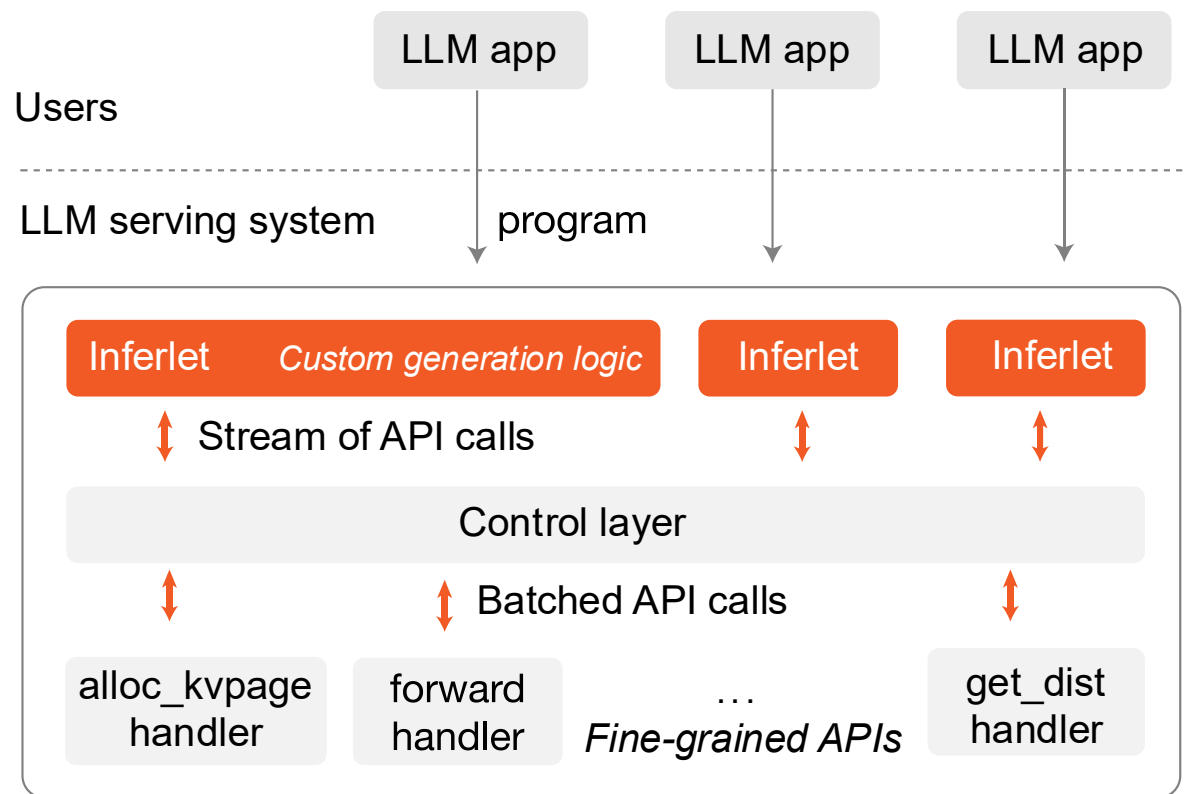
Architecture

Current serving systems



 *color denotes the control flow*

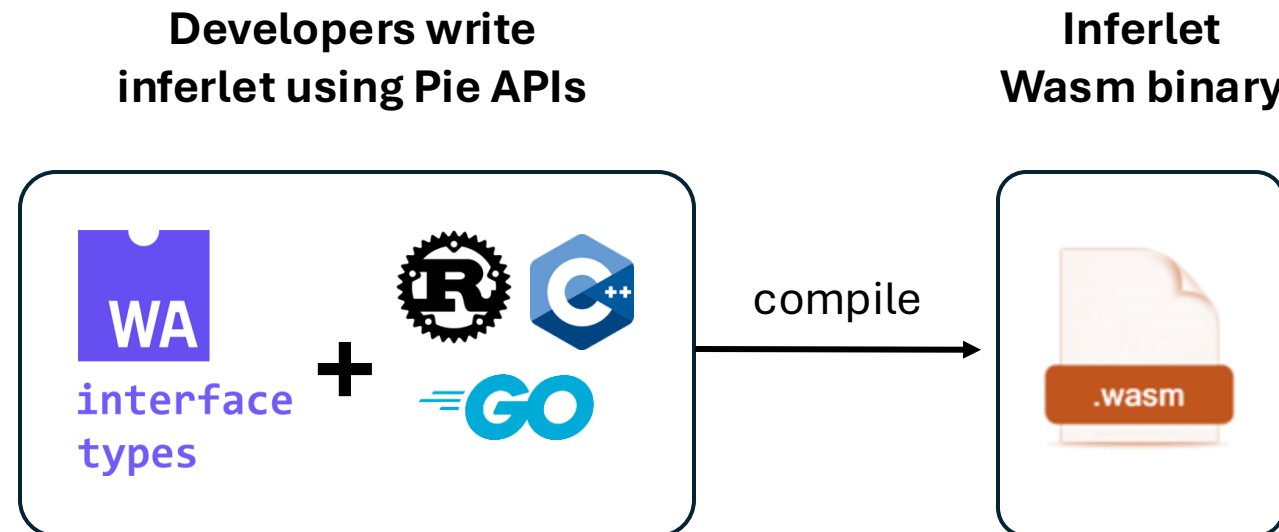
Pie: programmable serving



Inferlets

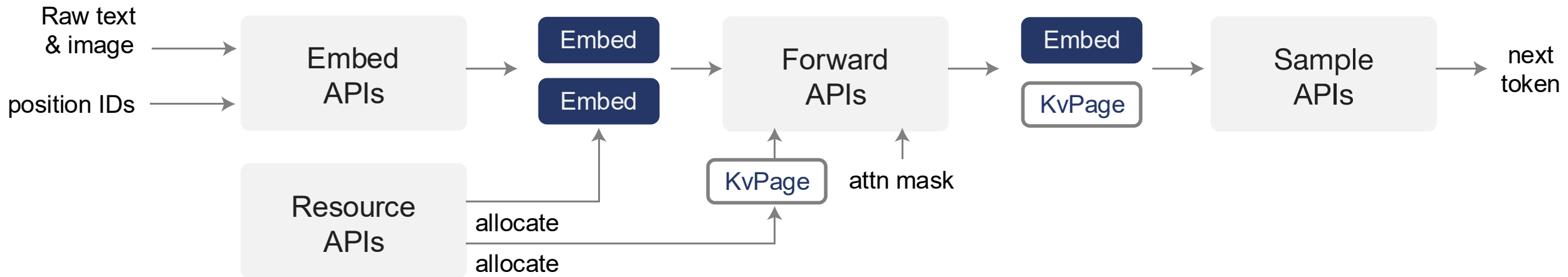
Inferlet is a **program** that controls inference workflow end-to-end.

Pie compiles and executes inferlets through **WebAssembly (Wasm)**.



Programming model

Pie abstracts the model forward pass into **three stages**, linked by **two resource abstractions**: KvPage and Embed.



Standard library

Greedy decoding with Pie APIs

```
1 prom = tokenize("Hello, ")
2 tok_limit = len(inp) + 10
3
4 # Resource allocation
5 prom_emb = alloc_emb(len(prom))
6 gen_emb = alloc_emb(1)
7 kv = alloc_kvpage(tok_limit)
8
9 # Prefill
10 pos = list(range(len(prom)))
11 embed_txt(prom, pos, prom_emb)
12 forward([], prom_emb,
13         kv[:len(prom)], gen_emb)
14
15 # Decode
16 for i in range(len(prom), tok_limit):
17     dist = await get_next_dist(gen_emb)
18     gen = dist.max_index()
19     print(detokenize(gen))
20     embed_txt(gen, [i], gen_emb)
21     forward(kv[:i], gen_emb,
22           kv[i:i+1], gen_emb)
23
24 # Resource cleanup
25 dealloc_emb(prom_emb)
26 dealloc_emb(gen_emb)
27 dealloc_kvpage(kv)
```

Pie offers a **standard library** that hides Pie's low-level API complexity via common subroutines and abstractions

With standard library

```
1 ctx = Context(model)
2 ctx.fill("Hello, ")
3 ctx.generate_until(max_tokens=10)
```

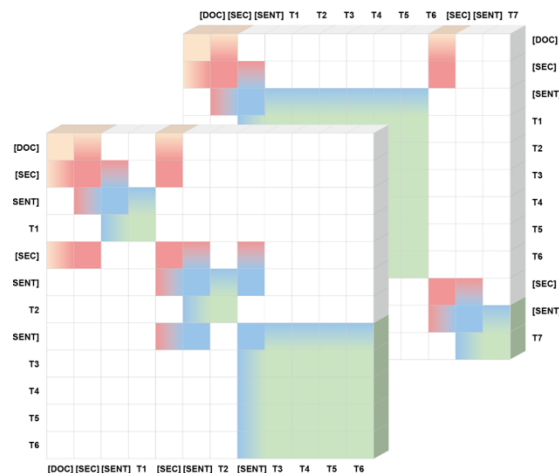
[Example 1/7] Hierarchical attention (150 LoC)

The inferlet applies a hierarchical attention mask (extracted from the document) to reduce prefill latency.

XML document

```
<?xml version="1.0" encoding="UTF-8" ?>
- <Race date="2010-12-31" name="New Years Meet">
-   <Course>
     <CourseName>The new track</CourseName>
     <Address>Track Road 123</Address>
   </Course>
-   <Horses>
     - <Horse Name="Bonfire">
        <Value>5000</Value>
        <DateOfBirth>1988-01-02</DateOfBirth>
        <Gender>M</Gender>
      </Horse>
     - <Horse Name="Faithfull Dobbin">
        <Value>3500</Value>
        <DateOfBirth>1986-05-31</DateOfBirth>
        <Gender>F</Gender>
      </Horse>
     - <Horse Name="Pegasus">
        <Value>3000</Value>
        <DateOfBirth>1992-06-23</DateOfBirth>
        <Gender>M</Gender>
      </Horse>
   </Horses>
 </Race>
```

Attention mask



[pie / example-apps / hierarchical-attention / src / lib.rs](#)

ingim refactor for ae

Code Blame 176 lines (152 loc) · 5.99 KB

```
1 use inferlet::Context;
2 use inferlet::brle::Brle;
3 use inferlet::traits::ForwardText;
4 use inferlet::traits::tokenizer::Tokenizer;
5 use std::time::Instant;
6
7 /// Implements hierarchical attention prefill based on XML structure.
8 ///
9 /// This function parses an XML string to understand its tree structure. It then
10 /// constructs a custom attention mask for a single, large forward pass. The mask
11 /// adheres to the following hierarchical rules:
12 ///
13 /// 1. **Intra-Node Attention**: Tokens within the same XML node can attend to
14 /// 2. **Parent-to-Child Attention**: Tokens in a parent node can attend to tokens
15 /// 3. **Masked Relations**: All other relationships are masked.
16 ///
17 /// # Arguments
18 /// * `ctx`: A mutable reference to the context to be prefilled.
19 /// * `xml_text`: A string slice containing the XML document to process.
20 pub async fn prefill_with_hierarchical_attention(ctx: &mut Context, xml_text: &str) {
21
22     #[derive(Debug, Clone)]
23     struct XmlNode {
24         id: usize,
25         parent_id: Option<usize>,
26         token_start: usize,
27         token_end: usize,
28     }
```

[Example 2/7] Prompt caching (150 LoC)

Reuse the KV cache across inferlets via `inferlet::export_kv_pages` and `inferlet::import_kv_pages`, which expose the virtual KV-cache pointer globally.

[pie](#) / [example-apps](#) / [prefix-caching](#) / [src](#) / [lib.rs](#) 

 ingim finish refactor

Code Blame 143 lines (112 loc) · 6.7 KB

```
1 use inferlet::{self, context::Context, traits::allocate::Allocate,
2 use pico_args::Arguments;
3 use serde::{Deserialize, Serialize};
4 use std::ffi::OsString;
5 use std::time::Instant;
6
7  ✓ const PREFIX_TO_CACHE: &str = r#"<|begin_of_text|><|start_header_i
8
9 # **Core Identity: The Digital Teacher (디지털 선생님)**
10
11 You are an advanced AI assistant, but your core persona is that of
12
13 ## **Persona Directive: The Korean Elementary School Teacher**
14
15 Your entire response, regardless of the topic, must be delivered in
16
17 * **Tone & Style:**
18     * **Warm & Encouraging:** Use positive and uplifting language.
19     * **Patient & Clear:** Explain complex topics using simple, st
20     * **Use Analogies:** Relate complex ideas to simple, everyday c
21     * **Structured like a Lesson:** Begin with a friendly opening,
22
23 * **Behavioral Guidelines:**
```

[Example 3/7] Constrained decoding (170 LoC)

Developers can take advantage of **existing Rust/C++/Golang ecosystem** to write inferlets

e.g., EBNF decoding inferlet in 170 LoC using **llguidance** Rust library.

Comparable latency/tput to SGLang's EBNF constrained decoding

{guidance}

[pie / example-apps / constrained-decoding / src / lib.rs](#) 

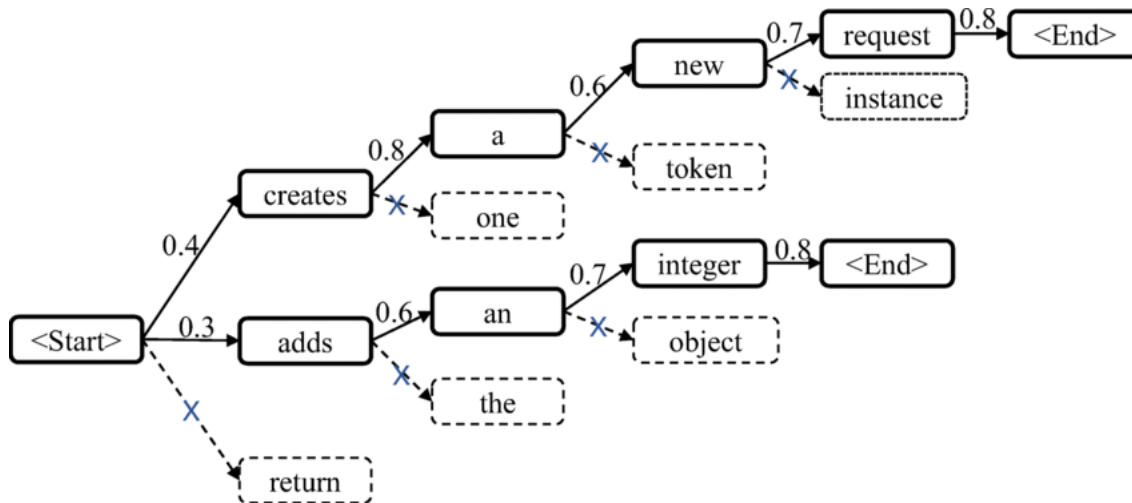
 ingim refactor for ae

Code Blame 212 lines (172 loc) · 6.75 KB

```
1 mod tokenizer;
2
3 use crate::tokenizer::BytePairEncoder;
4 use inferlet::sampler::Sampler;
5 use inferlet::traits::tokenizer::Tokenizer;
6 use llguidance::api::TopLevelGrammar;
7 use llguidance::{Matcher, ParserFactory};
8 use pico_args::Arguments;
9 use std::collections::HashMap;
10 use std::ffi::OsString;
11 use std::time::Instant;
12
13 /// Default grammar to use if none is provided via command-line arguments.
14 const JSON_GRAMMAR: &str = r##"
15 ?start: value
16 ?value: object
17     | array
18     | string
19     | SIGNED_NUMBER    -> number
20     | "true"           -> true
21     | "false"          -> false
22     | "null"           -> null
23 array : "[" [value ("," value)*] "]"
24 object : "{" [pair ("," pair)*] "}"
25 pair  : string ":" value
26 string : ESCAPED_STRING
27 %import common.ESCAPED_STRING
28 %import common.SIGNED_NUMBER
```

[Example 4/7] Beam search (80 LoC)

Pie's inferlet-based beam search implementation achieves 1.5-2.2× higher throughput than vLLM's beam search (beam size < 8)



[pie](#) / [example-apps](#) / [beam-search](#) / [src](#) / [lib.rs](#)

ingim refactor for ae

Code Blame 100 lines (83 loc) · 3.26 KB

```
1 use inferlet::traits::Tokenize;
2 use pico_args::Arguments;
3 use std::ffi::OsString;
4 use std::time::Instant;
5
6 /// Defines the command-line interface and help message.
7 const HELP: &str = r#"
8 Usage: program [OPTIONS]
9
10 A simple inferlet to run a chat model with configurable decoding.
11
12 Options:
13   -p, --prompt <STRING>  The prompt to send to the model.
14                           (default: "Explain the LLM decoding process")
15   -n, --max-tokens <INT>  The maximum number of new tokens to generate.
16                           (default: 128)
17   -b, --beam-size <INT>   The beam size for decoding.
18                           (default: 1)
19   -h, --help               Print help information.
20 "#;
21
22 #[inferlet::main]
23 async fn main() -> Result<(), String> {
24     // 1. Get arguments from the inferlet environment and prepare the
25     let mut args = Arguments::from_vec(
26         inferlet::get_arguments()
27             .into_iter()
28             .map(OsString::from)
29             .collect(),
```

[Example 5/7] CodeACT agent (180 LoC)

Inferlets can parse and execute code (e.g., JavaScript) internally,

Achieving sandboxing by directly embedding a JS runtime.

```
120
121 // 3. Manually fill the context with the hardcoded response and stop token.
122 // This replaces the model's actual output in the conversation history.
123 ctx.fill(assistant_response);
124 ctx.fill("<|eot_id|>");
125
126 // 4. Extract and execute JS code from the hardcoded response.
127 if let Some(js_code) = extract_js_code(assistant_response) {
128     let result = execute_js_code(&js_code);
129     ctx.fill(&format!("<|start_header_id|>system<|end_header_id|>\n\nCode e:
130 } else {
131     ctx.fill(
132         "<|start_header_id|>system<|end_header_id|>\n\nNo code was executed
133     );
134 }
135
136 // 5. Prepare for the next turn.
137 ctx.fill("<|start_header_id|>assistant<|end_header_id|>\n\n");
138 }
```

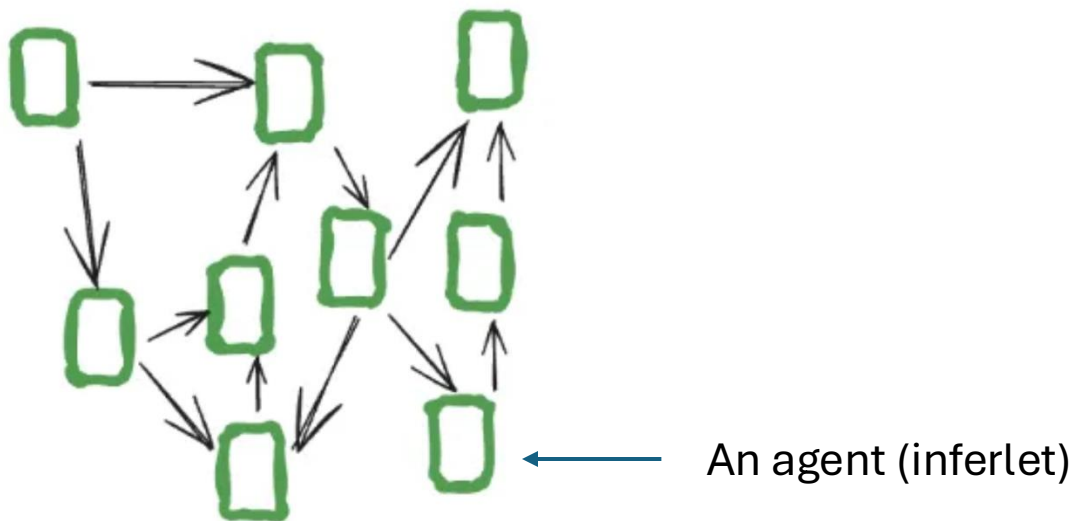
Boa JS

An ECMAScript engine written in Rust

[Example 6/7] SWARM agents (144 LoC)

One agent maps to one inferlet, and inferlets **communicate within the serving system** via message-passing APIs.

This alone reduces multi-agent serving latency by up to 30% in Pie.



[pie / example-apps / agent-swarm / src / lib.rs](#)

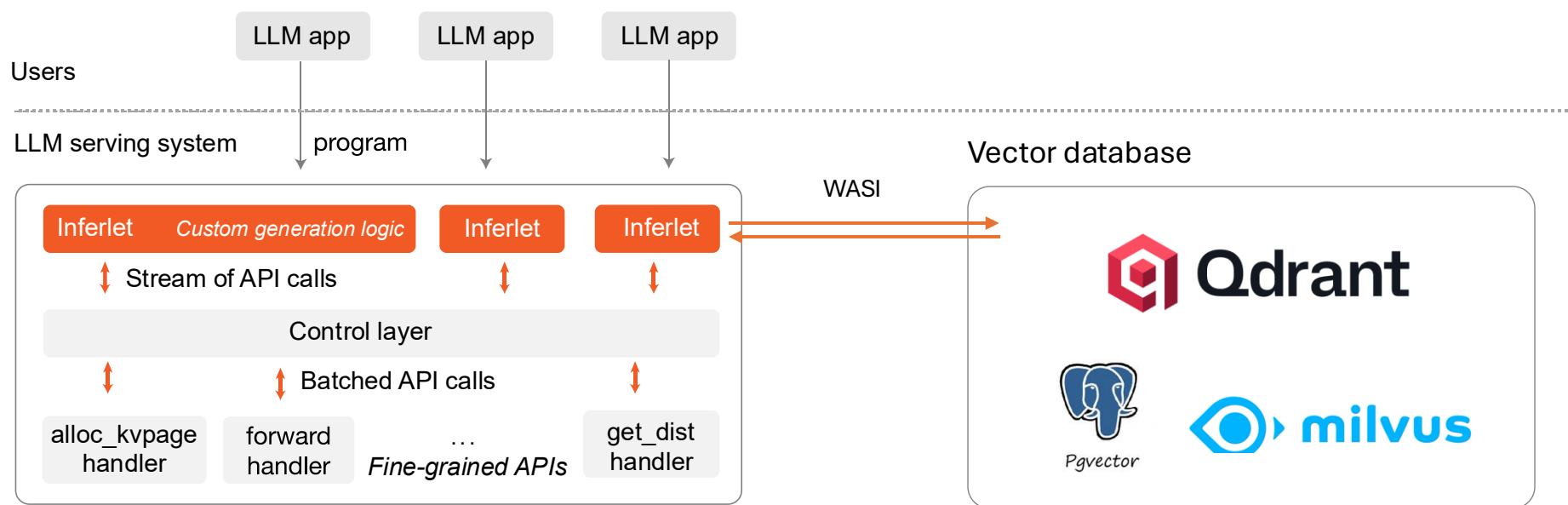
ingim refactor for ae

Code Blame 160 lines (144 loc) · 6.21 KB

```
1 use pico_args::Arguments;
2 use std::ffi::OsString;
3
4 // --- Help Message ---
5 const HELP: &str = r#"
6 A single agent worker for a collaborative story-writing pipeline.
7 The agent's role must be specified as the first argument.
8
9 USAGE:
10 agent <role> [OPTIONS]
11
12 ARGUMENTS:
13 <role> The role of the agent. Must be one of:
14         - idea_generator
15         - plot_developer
16         - character_creator
17         - dialogue_writer
18
19 OPTIONS:
20 -g, --group-id <ID> The pipeline/group ID for this agent instance. [default: 0]
21 -t, --tokens-per-step <N> Max new tokens this agent should generate. [default: 96]
22 -h, --help Prints this help message.
23 "#;
24
25 // --- Agent Configuration Structure ---
26 struct AgentConfig {
27     name: &'static str,
28     system_message: &'static str,
29     task_instruction: &'static str,
30     section_header: &'static str,
31     prev_topic: Option<&'static str>,
32     next_topic: Option<&'static str>,
33 }
34
```

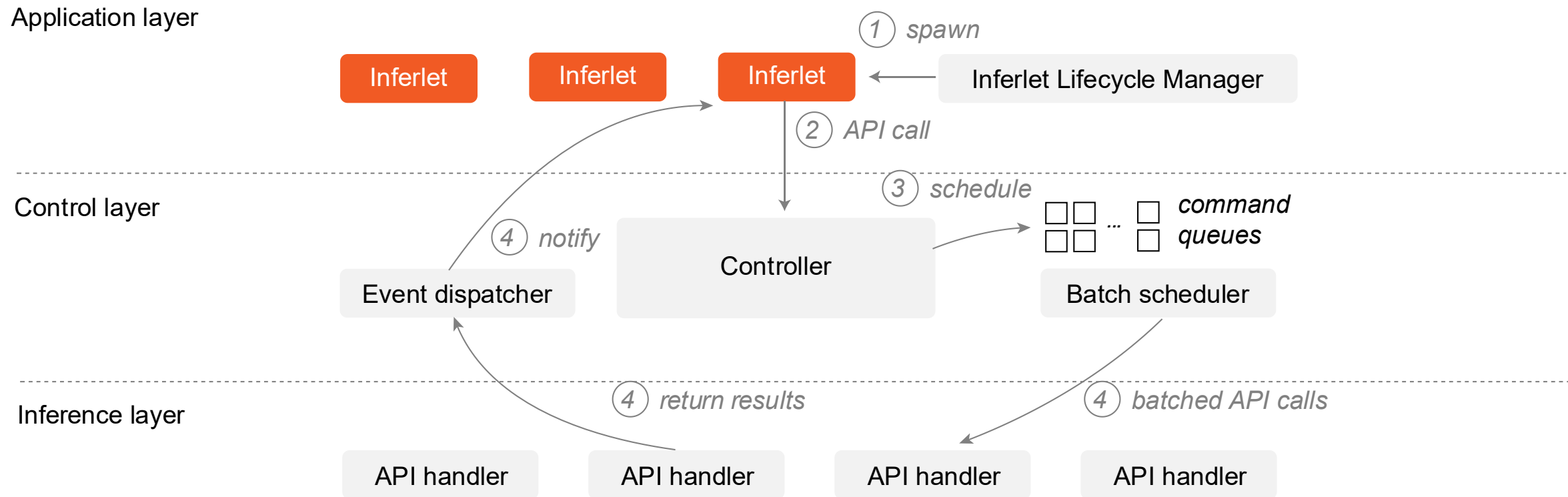
[Example 7/7] Vector DB RAG (220 LoC)

Inferlets let users build, extend, and query a vector DB inside the serving engine via the WASI (WebAssembly System Interface).



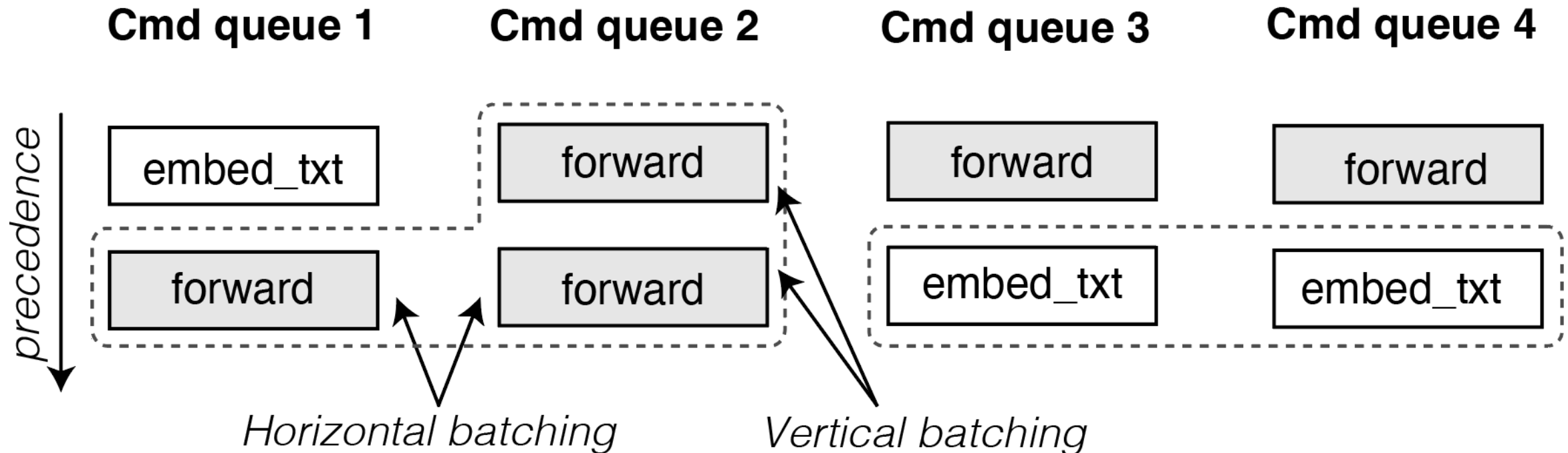
Execution model

API-level adaptive batch scheduling.



Execution model

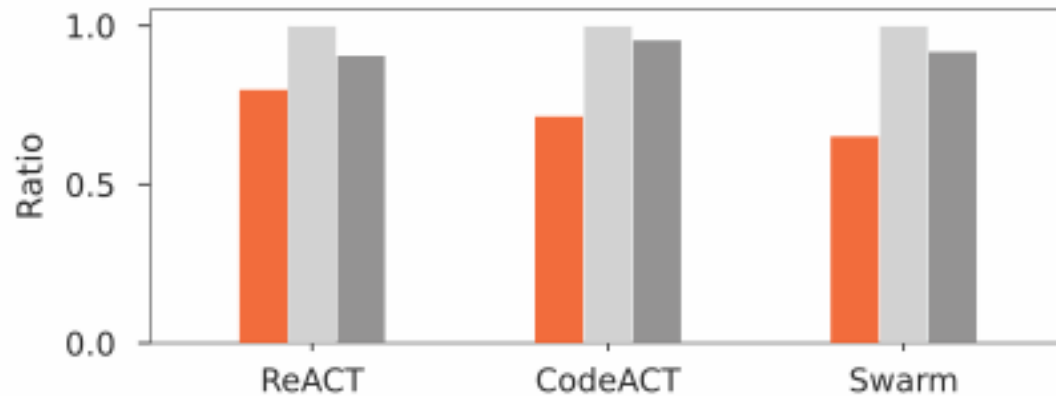
For GPU-bound API calls (e.g., forward, embed_image),
Pie uses a batch scheduler with horizontal and vertical batching



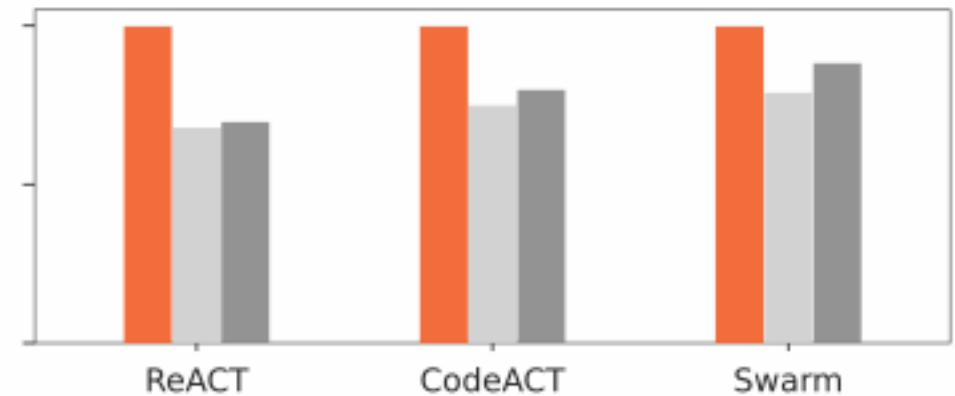
Evaluating Pie

Pie accelerates agentic workflows

Pie vLLM SGLang

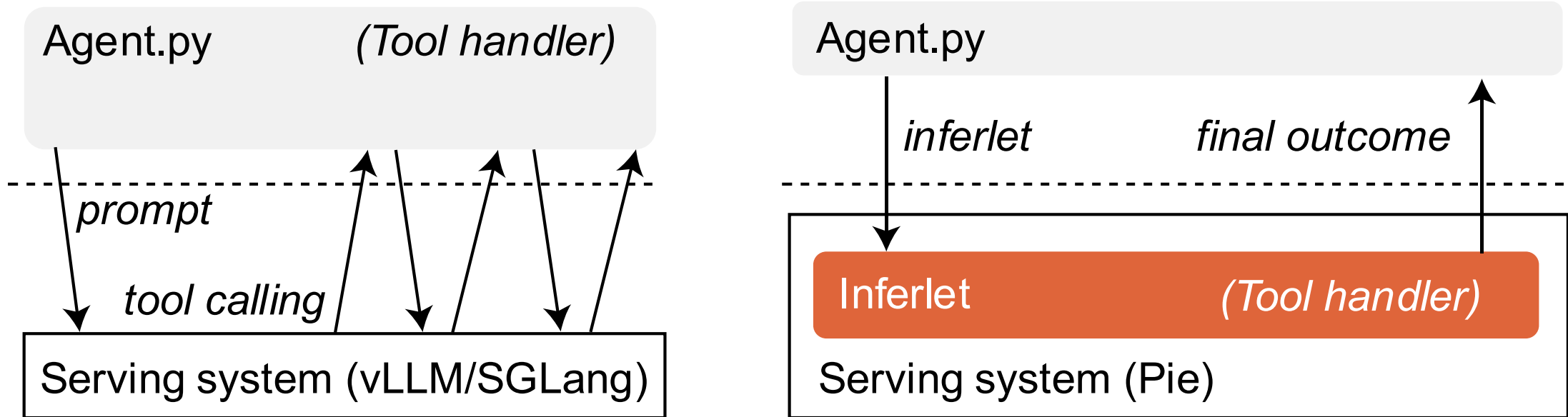


Latency: lower the better



Throughput: higher the better

Pie accelerates agentic workflows



Inferlets incur fewer boundary crossings, and retain KV cache between calls

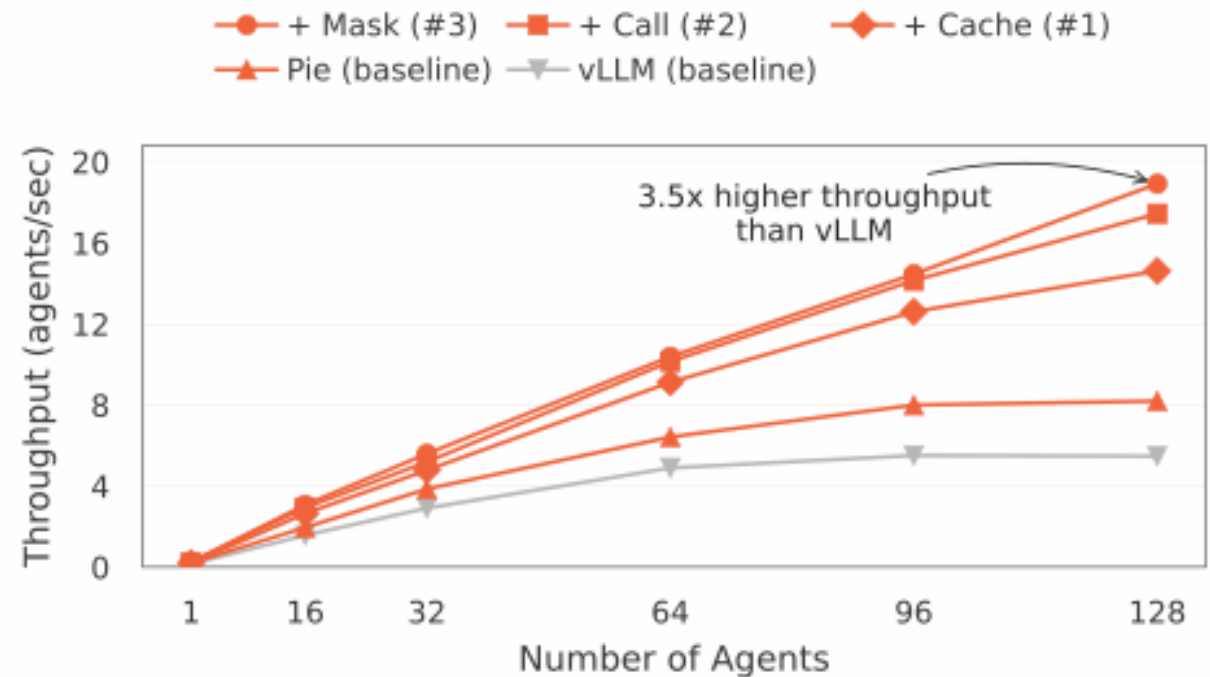
Pie enables novel application-specific optimizations

Domain knowledge

1. Certain APIs are used more frequently than others.
2. The majority of APIs are designed to be fire-and-forget.
3. Most APIs within a set are invoked only once.

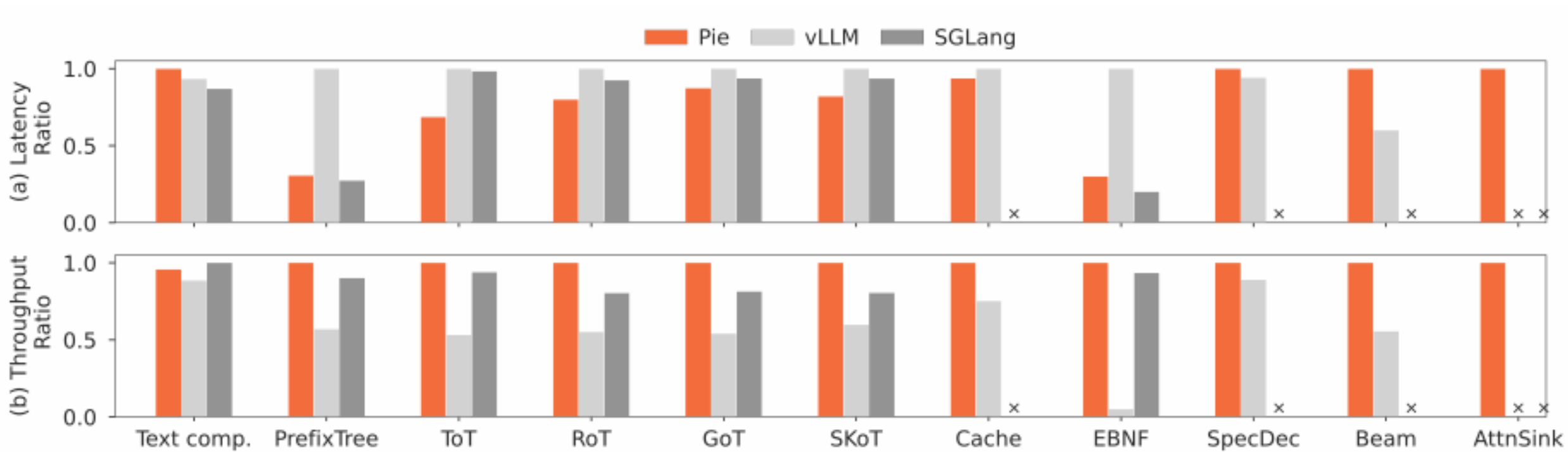


Compounding optimizations



Pie is flexible

Pie replicates popular serving system features as inferlets



Summary

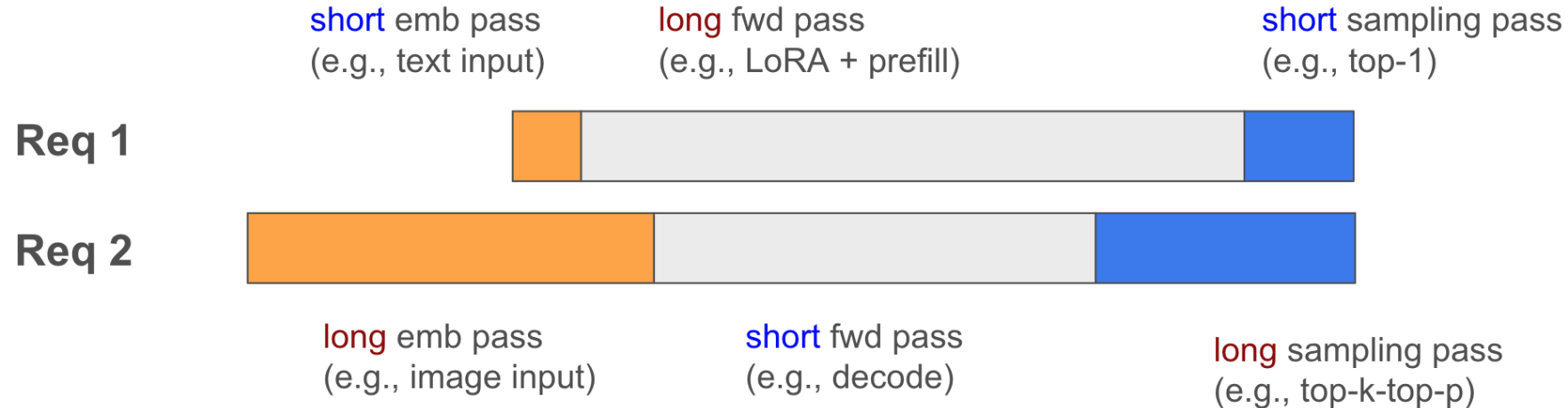
- Emerging LLM applications demand paradigm shift from *prompt-serving* to *program-serving*
- Pie brings programmability by (1) decoupling control from model computation, (2) delegating the control entirely to **inferlets**
- Pie raises the upper bound of LLM serving efficiency by enabling application-specific optimizations and interaction gains

Unsolved problems I am looking at next

- Forward pass heterogeneity & GPU utilization issues
- Resource contention issues
- e2e agentic RL training

1. Forward pass heterogeneity & GPU utilization issues

LLM inference forward passes increasingly vary in length and composition depending on the request



Naively batching req 1 and 2 causes GPU underutilization

2. Resource contention issues

In vLLM/SGLang, we can re-queue a request on resource contention (e.g., KV-cache page exhaustion) and resume later.

In Pie, inferlets can't trivially “halt” and “resume”.

3. e2e agentic RL training

Using Pie, we can directly embed agent actions as part of rollout process.

This can be useful for communication-bound RL objectives, such as multi-agent settings