

XGrammar

Aditya Ramesh

Motivation



```
{  
  "name": "Billi",  
  "animal": "cat",  
  "age": 9  
}
```

Desired



```
{  
  "name": "Billi",  
  "animal": "cat",  
  "age": 9,  
  "mood": "angry"  
}
```

Added key



```
{  
  name: "Billi"  
  "animal": "cat",  
  "age": "9"  
}
```

Incorrect format

We want to specify a specific JSON format but LLMs are not guaranteed to do so correctly!

Improve Performance

1. Prompt engineering
 - a. In context learning, multi-stage prompting
2. Finetuning
3. **Constrained Decoding**

Finite State Machines

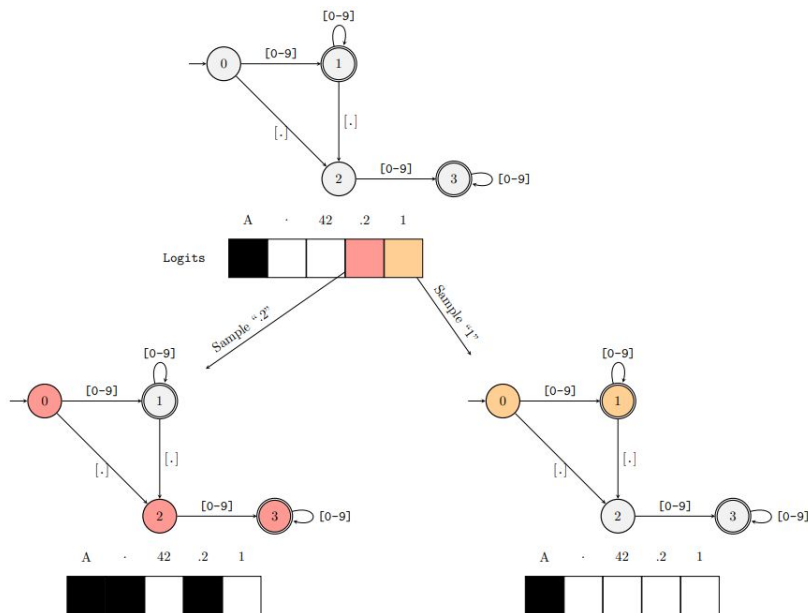


Figure 1: FSM masking for the regular expression $([0-9]^*)\.[0-9]^*$.

Use regex to construct a FSM that constrains the sampling of an LLM!

Finite State Machines



Use a mask to constrain the LLM

$$\begin{aligned}\alpha &= \text{LLM}(S_t, \theta) \\ s_{t+1} &\sim \text{Categorical}(\alpha)\end{aligned}$$

$$\begin{aligned}\alpha &= \text{LM}(\tilde{S}_t, \theta) \\ \tilde{\alpha} &= m(\tilde{S}_t) \odot \alpha \\ \tilde{s}_{t+1} &\sim \text{Categorical}(\tilde{\alpha})\end{aligned}$$

$\tilde{S}_t$: current sequence,
 $m(\tilde{S}_t)$: token mask,
 α : logits,
 $\text{Categorical}(\tilde{\alpha})$: sampling method.

Outlines (2023)



1. Define regex
2. Convert regex into FSM
3. Precompute mask for each state in FSM
4. Store mask as hashmap
5. Use mask to sample

Algorithm 2 LLM token sampling with masking

```
1: function SAMPLE_TOKENS( $L$ )
2:    $s \leftarrow ()$ 
3:   for  $i \leftarrow 1, L$  do
4:      $\alpha \leftarrow \text{LLM}(s, \theta)$ 
5:     Construct the mask  $m(s)$  ← Precomputed before sampling
6:      $\tilde{\alpha} \leftarrow m \odot \alpha$ 
7:     Sample  $\tilde{s} \sim \text{Categorical}(\tilde{\alpha})$ 
8:     if  $\tilde{s} = \text{EOS}$  then
9:       break
10:    end if
11:     $s \leftarrow \text{append}(s, \tilde{s})$ 
12:  end for
13:  return  $s$ 
14: end function
```

Interleave Optimization (2023)

{guidance}

```
@guidance
def character_maker(lm, name):
    lm += f"""\
{name} is a character in Harry Potter. Please fill in the following information about this character.
{{
  "name": "{guidance.gen("name", regex=r"[\w\d\s]+")}",
  "age": "{guidance.gen("age", regex=r"[0-9]+")}",
  "house": "{guidance.select(options=["Gryffindor", "Slytherin", "Ravenclaw", "Hufflepuff"], name="house")}"
}}
"""
    return lm
```

Interleaved Syntax

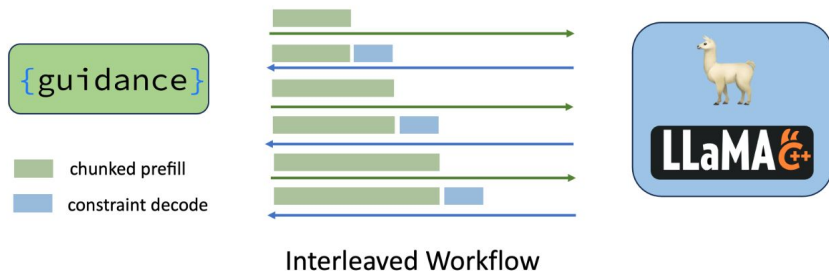


Figure from [LMSYSORG](https://lmsys.org)

Jump-Forward Optimization (2023)



```
Please fill in the following  
information about Harry Potter.  
{  
  "name": "Harry",  
  "age": 15,  
  "house": "Gryffindor"  
}
```

```
Please fill in the following  
information about Harry Potter.  
{  
  "name": "Harry",  
  "age": 15,  
  "house": "Gryffindor"  
}
```

Figure from [LMSYSORG](https://www.lmsys.org/)

Is regex enough? No!

JSON and SQL have recursive structures

regex can not model **recursive structures** (Chomsky 1956)

JSON and SQL need to be modeled with **Context-free Grammar**
(Chomsky 1956)

Context-Free Grammar (CFG)

the man	took	the book
NP	Verb	NP
	VP	
Sentence		

(20) $\Sigma : \# \wedge \text{Sentence} \wedge \#$
 $F: \text{Sentence} \rightarrow \text{NP} \wedge \text{VP}$
 $\text{VP} \rightarrow \text{Verb} \wedge \text{NP}$
 $\text{NP} \rightarrow \text{the} \wedge \text{man}, \text{the} \wedge \text{book}$
 $\text{Verb} \rightarrow \text{took}$

(21) $D_1: \# \wedge \text{Sentence} \wedge \#$
 $\# \wedge \text{NP} \wedge \text{VP} \wedge \#$
 $\# \wedge \text{NP} \wedge \text{Verb} \wedge \text{NP} \wedge \#$
 $\# \wedge \text{the} \wedge \text{man} \wedge \text{Verb} \wedge \text{NP} \wedge \#$
 $\# \wedge \text{the} \wedge \text{man} \wedge \text{Verb} \wedge \text{the} \wedge \text{book} \wedge \#$
 $\# \wedge \text{the} \wedge \text{man} \wedge \text{took} \wedge \text{the} \wedge \text{book} \wedge \#$

$D_2: \# \wedge \text{Sentence} \wedge \#$
 $\# \wedge \text{NP} \wedge \text{VP} \wedge \#$
 $\# \wedge \text{the} \wedge \text{man} \wedge \text{VP} \wedge \#$
 $\# \wedge \text{the} \wedge \text{man} \wedge \text{Verb} \wedge \text{NP} \wedge \#$
 $\# \wedge \text{the} \wedge \text{man} \wedge \text{took} \wedge \text{NP} \wedge \#$
 $\# \wedge \text{the} \wedge \text{man} \wedge \text{took} \wedge \text{the} \wedge \text{book} \wedge \#$

Introduced as **Phrase Structure** (Chomsky 1956)

JSON as CFG

JSON-text = ws value ws

These are the six structural characters:

begin-array = ws %5B ws ; [left square bracket

begin-object = ws %7B ws ; { left curly bracket

end-array = ws %5D ws ;] right square bracket

end-object = ws %7D ws ; } right curly bracket

name-separator = ws %3A ws ; : colon

value-separator = ws %2C ws ; , comma

```
ws = *(
    %x20 /           ; Space
    %x09 /           ; Horizontal tab
    %x0A /           ; Line feed or New line
    %x0D )           ; Carriage return
```

value = false / null / true / object / array / number / string

false = %x66.61.6c.73.65 ; false

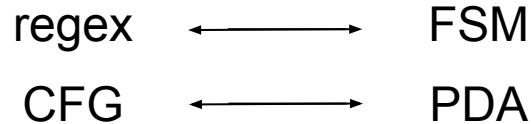
null = %x6e.75.6c.6c ; null

true = %x74.72.75.65 ; true

JSON is defined with Context-Free Grammar!

<https://datatracker.ietf.org/doc/html/rfc8259#section-2>

Pushdown Automata (PDA)



Finite state machines can not model recursion!
(Chomsky 1956)

Pushdown (or stack) is needed to analyze nested formal languages left-to-right!
(Oettinger 1961)

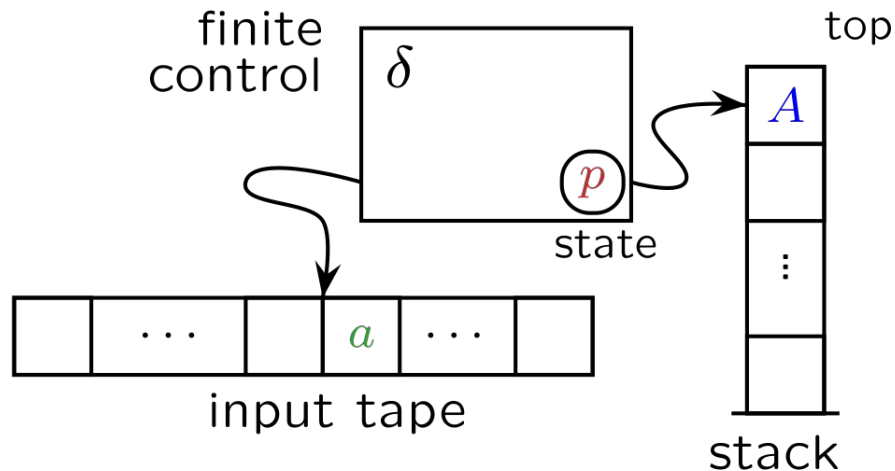
Pushdown Automata creates a context-free language!
(Schützenberger 1963)

Pushdown Automata (Informal)

A **pushdown automaton** is a finite-state machine plus a **stack-based memory**.

On transition we can

- a) Transition along the FSM
- b) Manipulate the stack



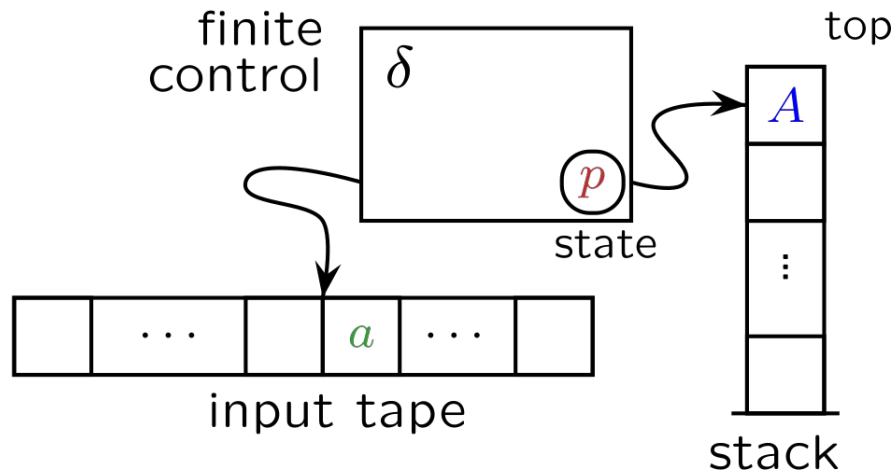
From [Wikipedia](https://en.wikipedia.org/wiki/Pushdown_automaton)

Pushdown Automata

Deterministic PDA (PDA) - For every input and stack symbol, there is only one transition

Non-deterministic PDA (NPDA) - if several actions are possible at a input and stack symbol

NPDA will maintain multiple stacks in parallel for each path



From [Wikipedia](https://en.wikipedia.org/wiki/Pushdown_automaton)

Pushdown Automata Example

$$L = \{a^n b^n \mid n \geq 0\}$$

Examples

aabb
aaabbb
ab

Non-Examples

abb
b
bbaa
ababa

FSM alone can not keep track of how many 'a' are present as we parse left to right. A stack is needed!

Pushdown Automata Example

$$L = \{a^n b^n \mid n \geq 0\}$$

Intuition

Read a: push A

Read b: pop A

Accept if input is finished and stack is empty

Input left	Stack
aaabbb	\$
aabbb	A\$
abbb	AA\$
bbb	AAA\$
bb	AA\$
b	A\$
ϵ	\$
accept	

JSON

```
{
  "name": "Billi",
  "friends": [
    { "name": "Lea" },
    { "name": "Summer" },
    { "name": "Bunny" }
  ]
}
```

Input token	Action	Stack
{	push OBJECT	OBJECT
"name"	read object key	OBJECT
:	read colon	OBJECT
"Billi"	read string value	OBJECT
,	expect another pair	OBJECT
"friends"	read object key	OBJECT
:	read colon	OBJECT
[	push ARRAY	OBJECT ARRAY
{	push OBJECT	OBJECT ARRAY OBJECT
"name"	read object key	OBJECT ARRAY OBJECT
:	read colon	OBJECT ARRAY OBJECT
"Lea"	read string value	OBJECT ARRAY OBJECT
}	pop OBJECT	OBJECT ARRAY
,	expect another element	OBJECT ARRAY
{	push OBJECT	OBJECT ARRAY OBJECT
"name"	read object key	OBJECT ARRAY OBJECT
:	read colon	OBJECT ARRAY OBJECT
"Summer"	read string value	OBJECT ARRAY OBJECT
}	pop OBJECT	OBJECT ARRAY
,	expect another element	OBJECT ARRAY
{	push OBJECT	OBJECT ARRAY OBJECT
"name"	read object key	OBJECT ARRAY OBJECT
:	read colon	OBJECT ARRAY OBJECT
"Bunny"	read string value	OBJECT ARRAY OBJECT
}	pop OBJECT	OBJECT ARRAY
]	pop ARRAY	OBJECT
}	pop OBJECT	empty
end	accept	empty

XGrammar

Problem

PDA's unbounded stack length results in infinite number of states

Can't precompute token masks for all scenarios

Must compute mask during runtime

XGrammar

Preprocessing phase:

Generate adaptive token mask cache for ***context-independent*** tokens

Runtime phase:

Process ***context-dependent*** tokens with persistent execution stack

XGrammar - Token Context

During transition, we either

- 1) Expand into child rule (push onto stack)
- 2) Advance within current rule (FSM)
- 3) Return to parent rule (pop from stack)

For case (1) and (2), we only need the stack's top node - ***Context-independent!***

For case (3), we inspect the stack - ***Context-dependent!***

Experiment show >1% of tokens are context dependent!
(Llama-3.1 on JSON grammar)

XGrammar - Token Mask Cache

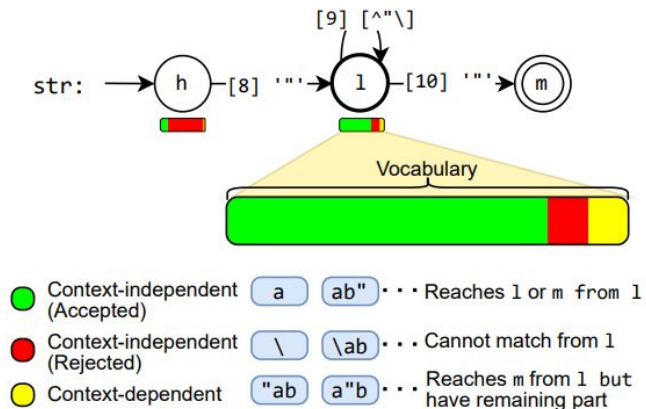


Figure 4. An example for the token mask cache. Tokens are categorized into three types: context-independent (accepted), context-independent (rejected), and context-dependent. The first two types can be directly determined for mask generation at runtime.

Figure from [XGrammar paper](#)

Each node has a set of context-independent and context-dependent tokens

XGrammar - Adaptive Storage Format



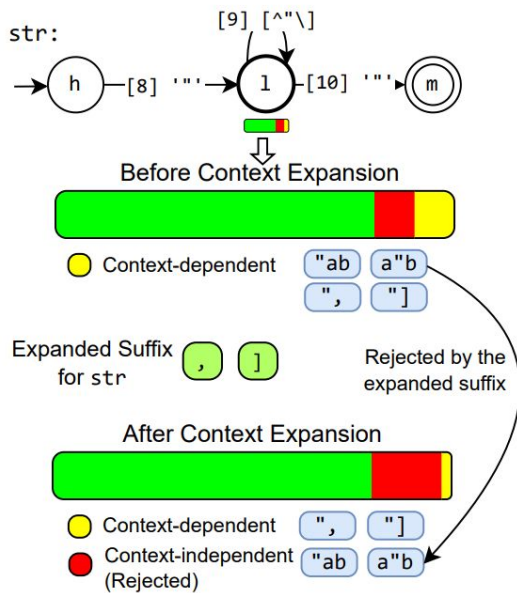
Figure from [XGrammar paper](#)

Figure 5. The adaptive storage format. In accept-heavy cases, we store the rejected tokens and context-dependent tokens. In reject-heavy cases, we store the accepted tokens and context-dependent tokens. In rare cases where two kinds of tokens are equal, we compress the accepted and rejected tokens into a bitset of the vocabulary size.

For Llama-3.1 on JSON grammar, reduced total memory from 160 MB to 0.46 MB!

XGrammar - Context Expansion

Checking context-dependent tokens still a bottleneck



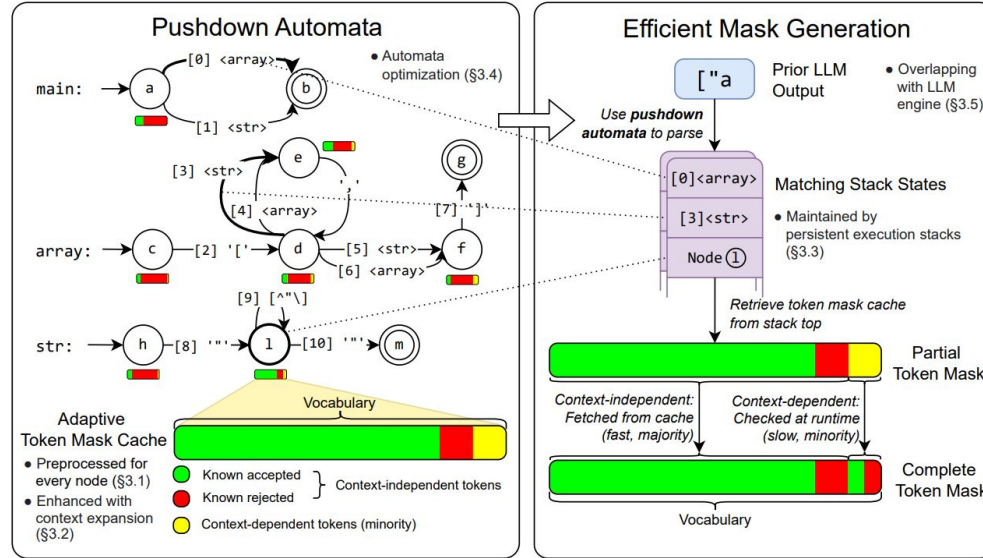
rule: str

possible parent continuations:

`" , "` ...
`" ] "` ...
`" } "` ...
`" : "` ...

Can be precomputed!

XGrammar - Persistent Execution Stack



The may be ambiguity in what the stack will be as we parse left-to-right

XGrammar - Persistent Execution Stack

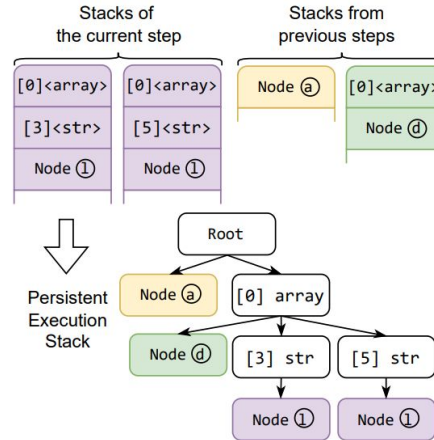


Figure 7. The persistent stack organizes multiple matching stacks from the current step, as well as stacks from previous steps, into a single tree. It reduces memory consumption and supports rolling the state back to previous steps.

Persistent Execution Stack for efficient state branching!

Allows for efficient rollback to check for valid states

XGrammar - PDA optimization

Rule inlining

Before

```
array -> "[" elements? "]"
elements -> value ("," value)*
```

After

```
array -> "[" value ("," value)* "]"
```

XGrammar - PDA optimization

Node Merging

```
state S
├─ 'a' -> node X
└─ 'a' -> node Y
```

```
state S
└─ 'a' -> merged node XY
```

We can merge when all ways to reach the node are equivalent

XGrammar - Overlap

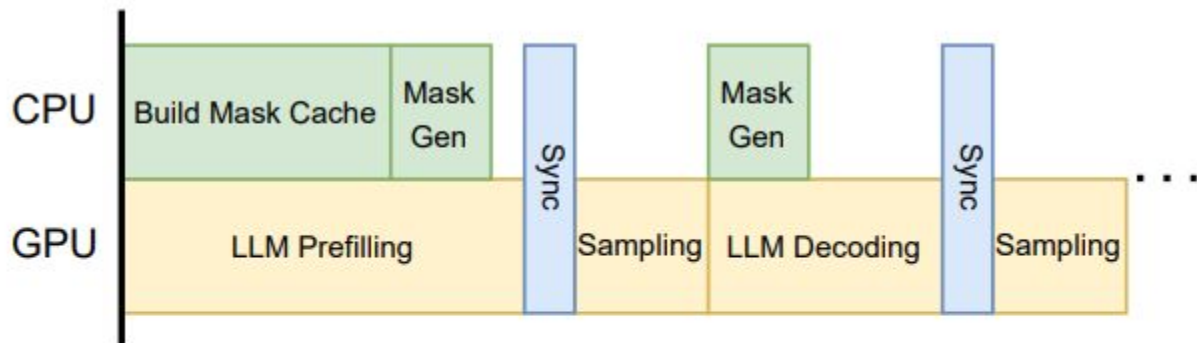


Figure 8. Overlapping building the mask cache with LLM prefilling, and mask generation with LLM decoding to minimize the overhead.

XGrammar - Results

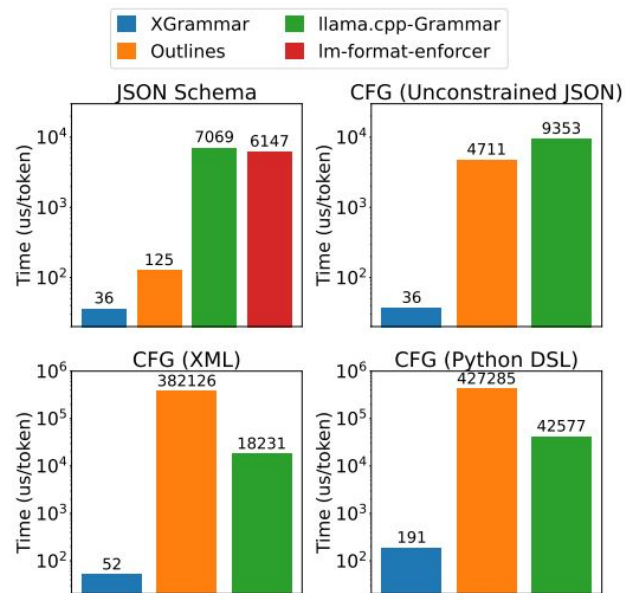


Figure 9. Per token mask generation latency. XGrammar consistently outperforms existing constrained decoding libraries.

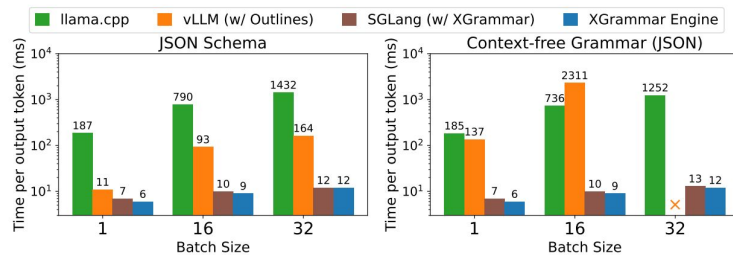


Figure 10. End-to-end evaluation on Llama 3.1 inference with structured constraints. Some results with a batch size of 32 are not reported because their API call time exceeded the API timeout limit of 600 seconds.

XGrammar 2

TagDispatch

Problem: Agent generates free form text then emit a tag (<function=get_weather>), then output follows JSON schema for function. Encoding as one EBNF grammar file is inefficient

Solution: TagDispatch compiles all candidate tag as on DFA and switches between dispatching and dispatched mode

Cross-Grammar Cache

Problem: Multiple EBNF structures may contain the same sub-FSM structure leading to redundant computing of mask

Solution: Cross Grammar cache reuse.

XGrammar 2

Earley-based Adaptive Token Mask Cache

Problem: NPDA can blow up, growing exponentially.

Solution: Adapt the token mask to use Earley states

JIT Compilation of Token Mask Cache

Problem: In dynamic agents, a request may have hundreds of unique tool calls and grammar, but only a few is used. Compiling all increases TTFT.

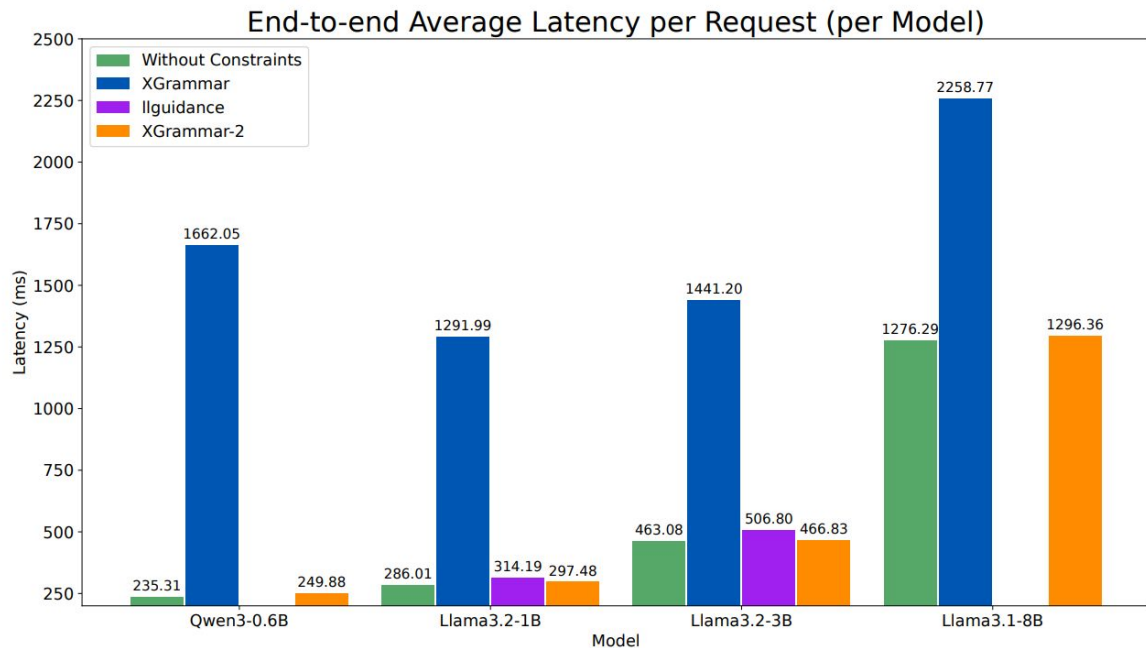
Solution: Only compile when used (JIT)

Repetition State Compression

Problem: Grammar might have bounded repetition, expanding this naively would lead to redundant computation of token-mask caches.

Solution: Introduce repetition construct

XGrammar 2 Performance





- [Fast JSON Decoding with FSM blog](#)
- [Outlines](#)
- [Coalescence](#)
- [Three Models for the Description of Language](#)
- [On Context-Free Languages and Push-Down Automata](#)
- [XGrammar](#)
- [XGrammar-2](#)
- [LLGuidance: Making Structured Outputs Go Brrr](#)