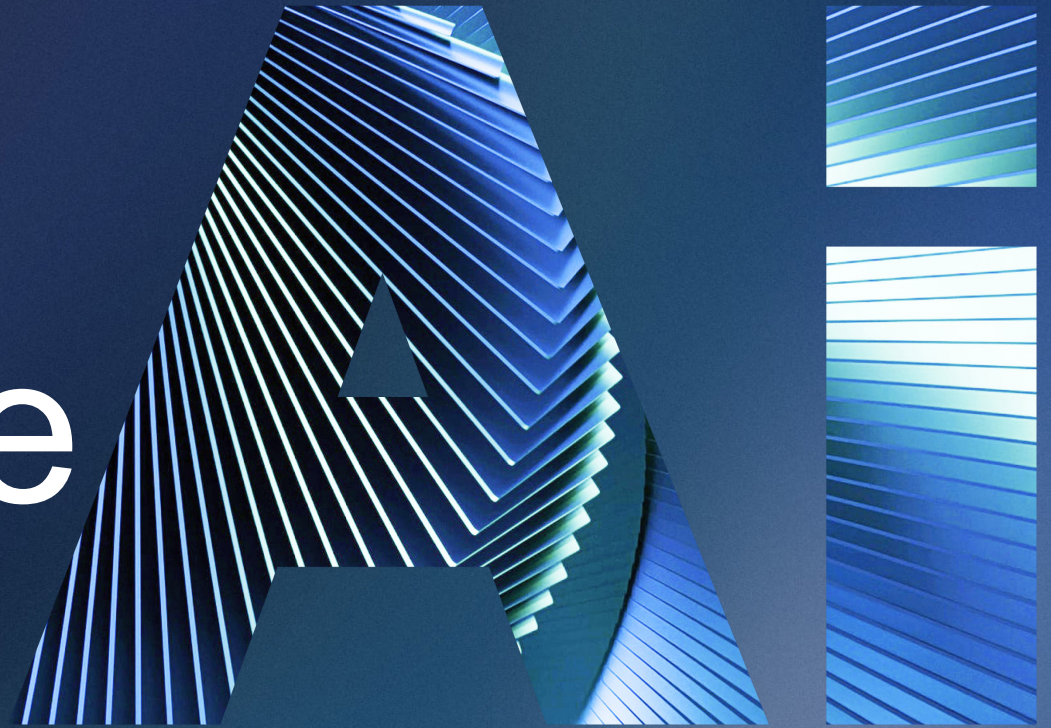


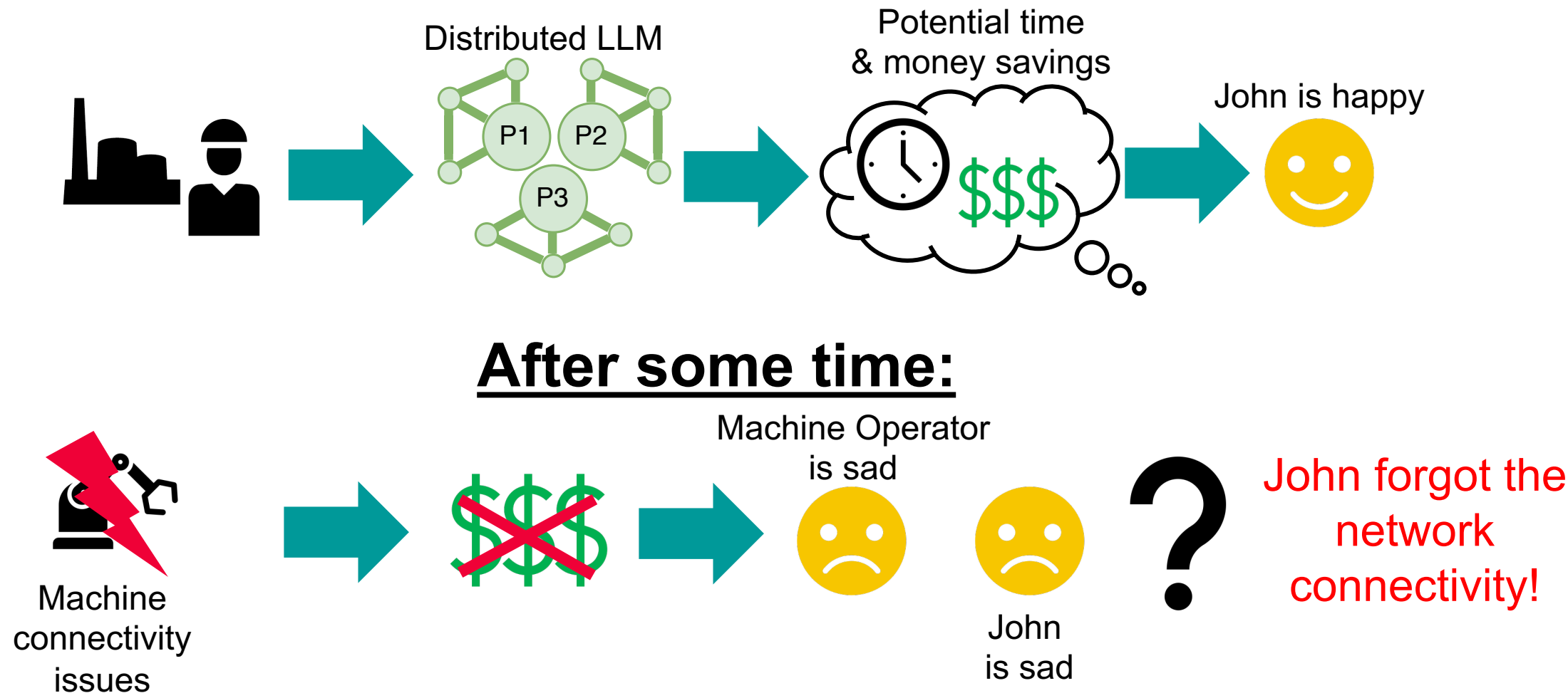
LLMs on Edge

Network Traffic Characteristics of Distributed Inference under the Loupe

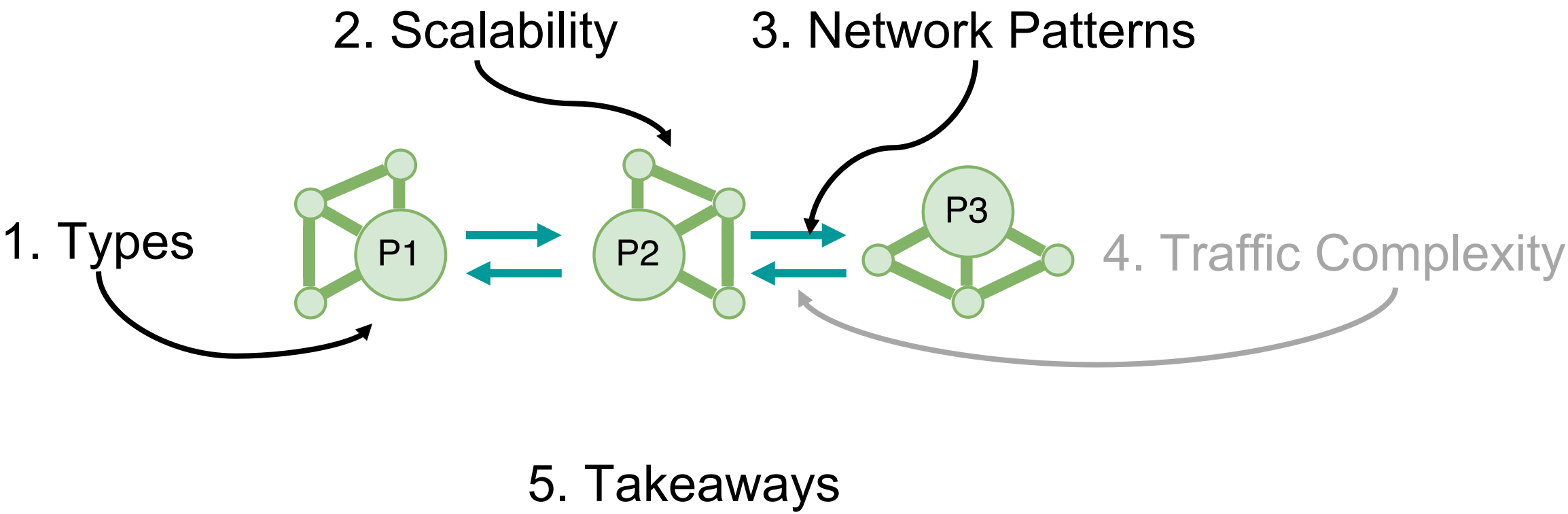


Philippe Buschmann, Arne Bröring, Georg Carle, Andreas Blenk

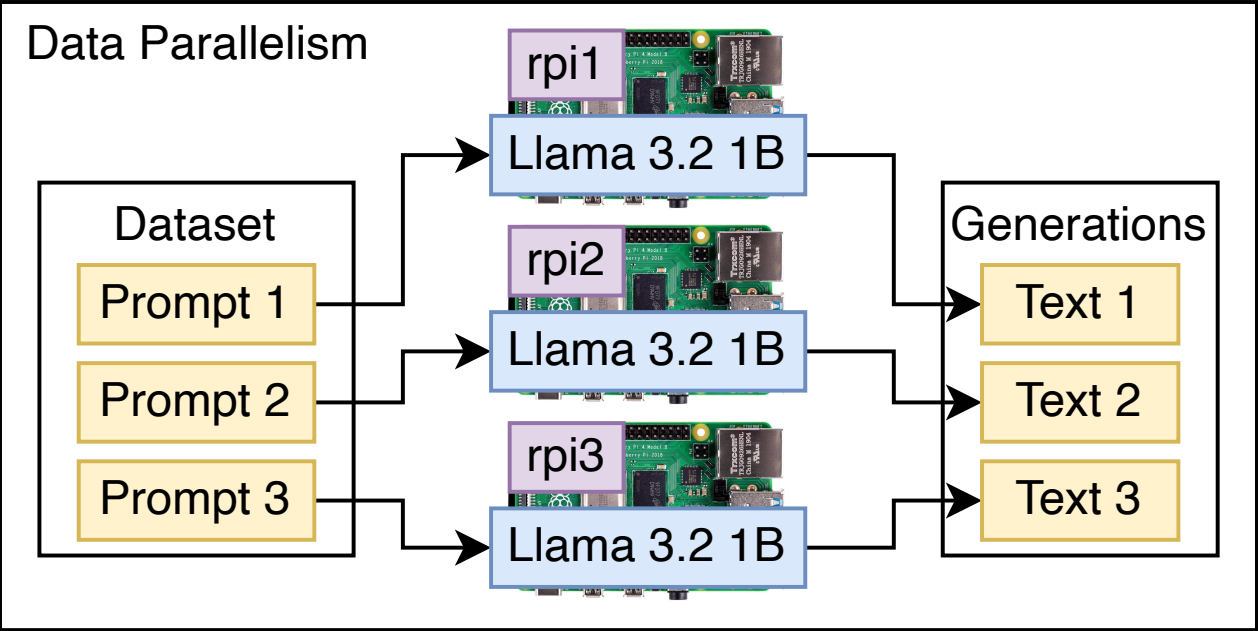
Motivation



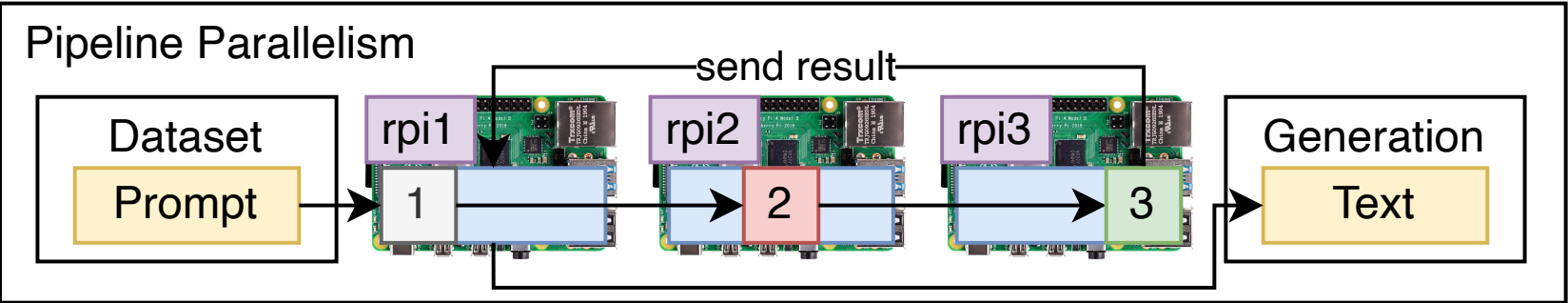
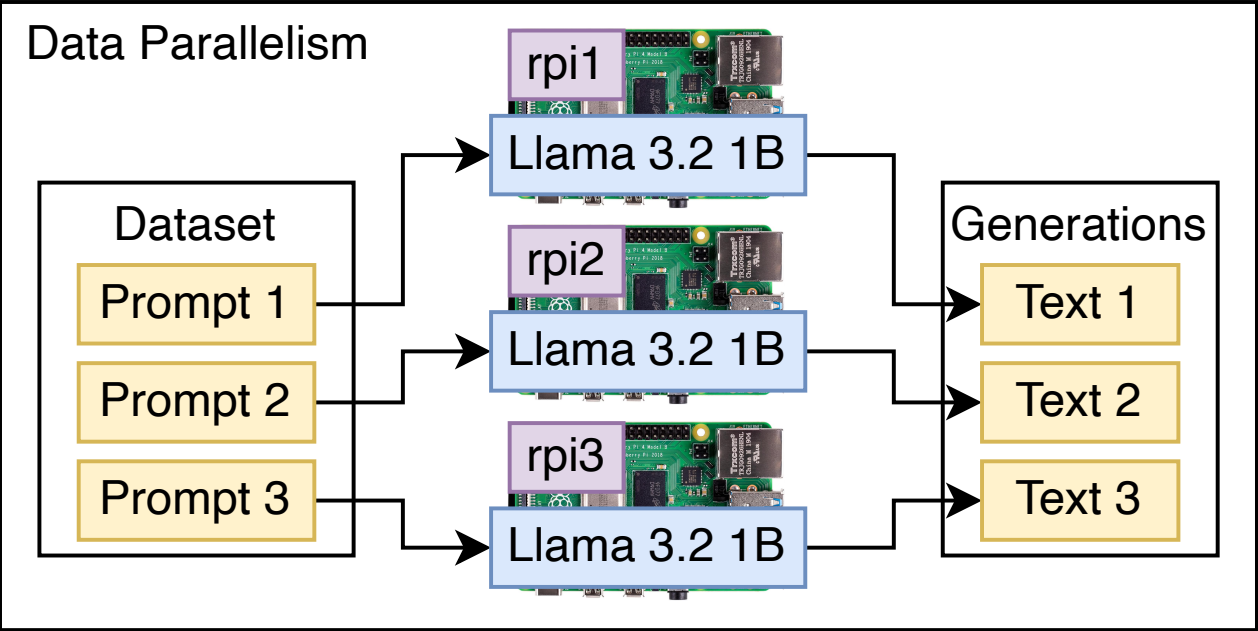
Distributed Large Language Models



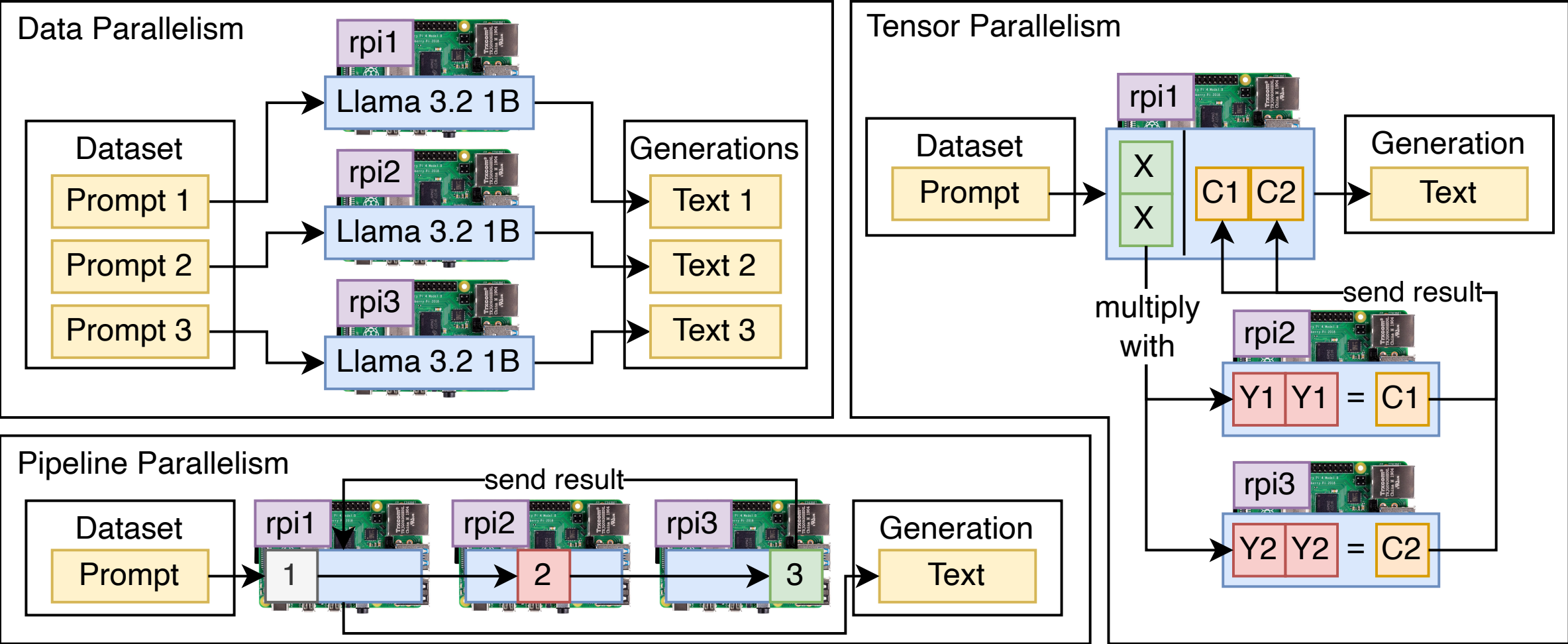
Types of Distributed Large Language Models



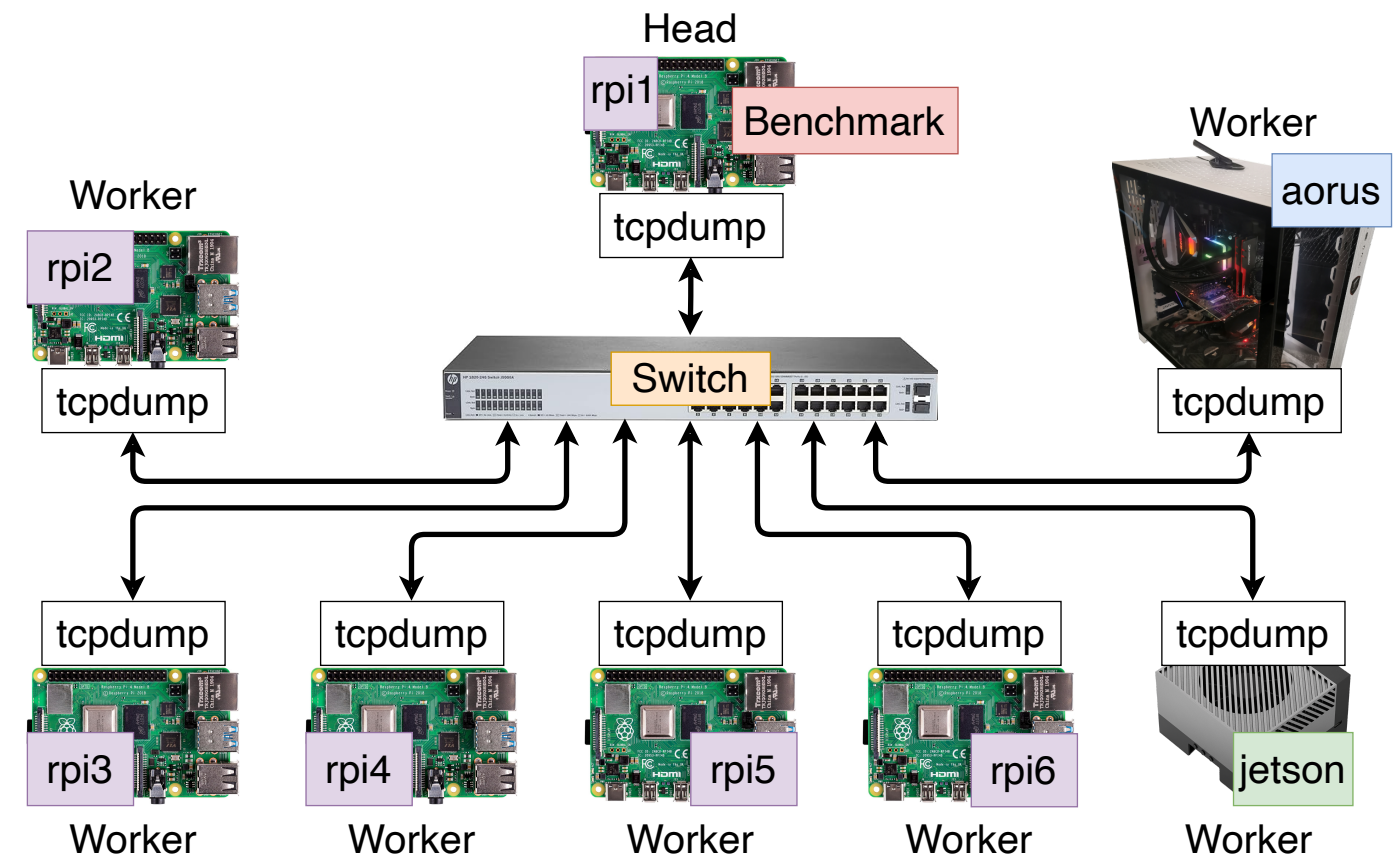
Types of Distributed Large Language Models



Types of Distributed Large Language Models



Testbed for Distributed LLMs



Pipeline Parallel Framework

llama.cpp

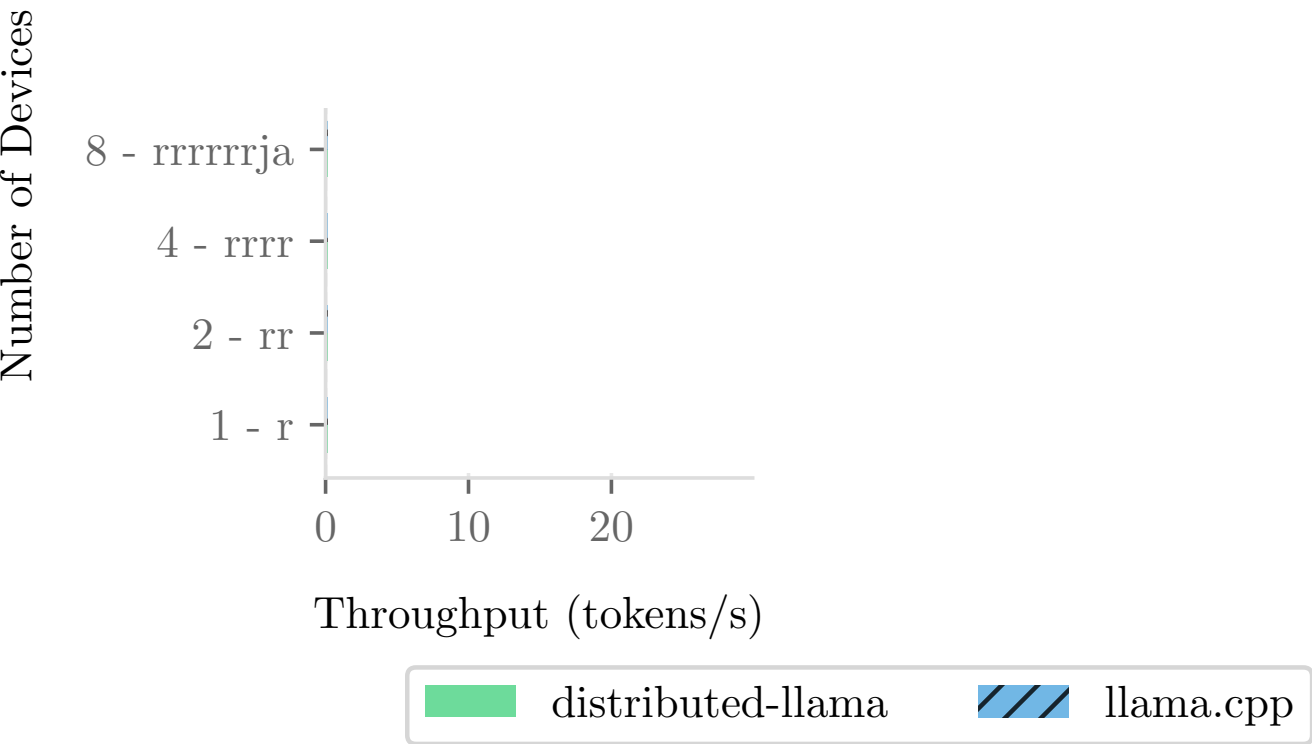
Tensor Parallel Framework

distributed-llama

Results – Scalability

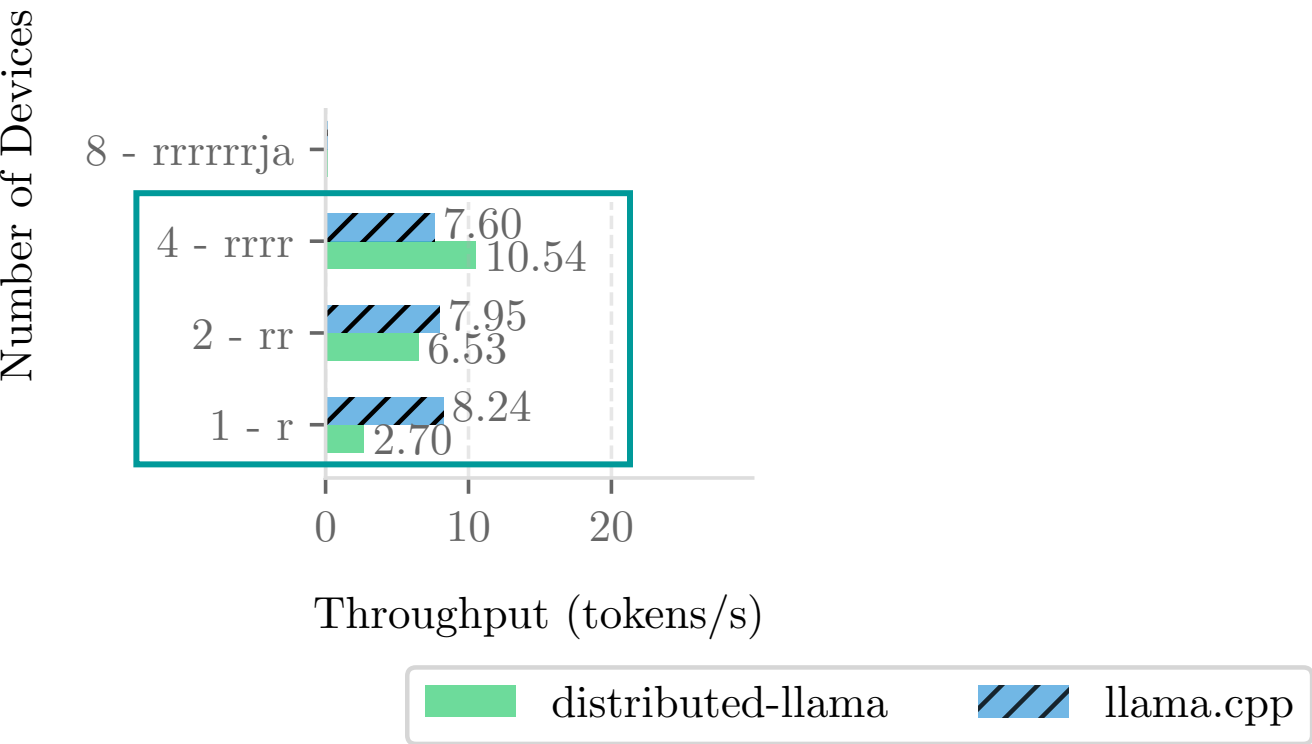
(a) Llama 3.2 1B

Results:



Results – Scalability

(a) Llama 3.2 1B

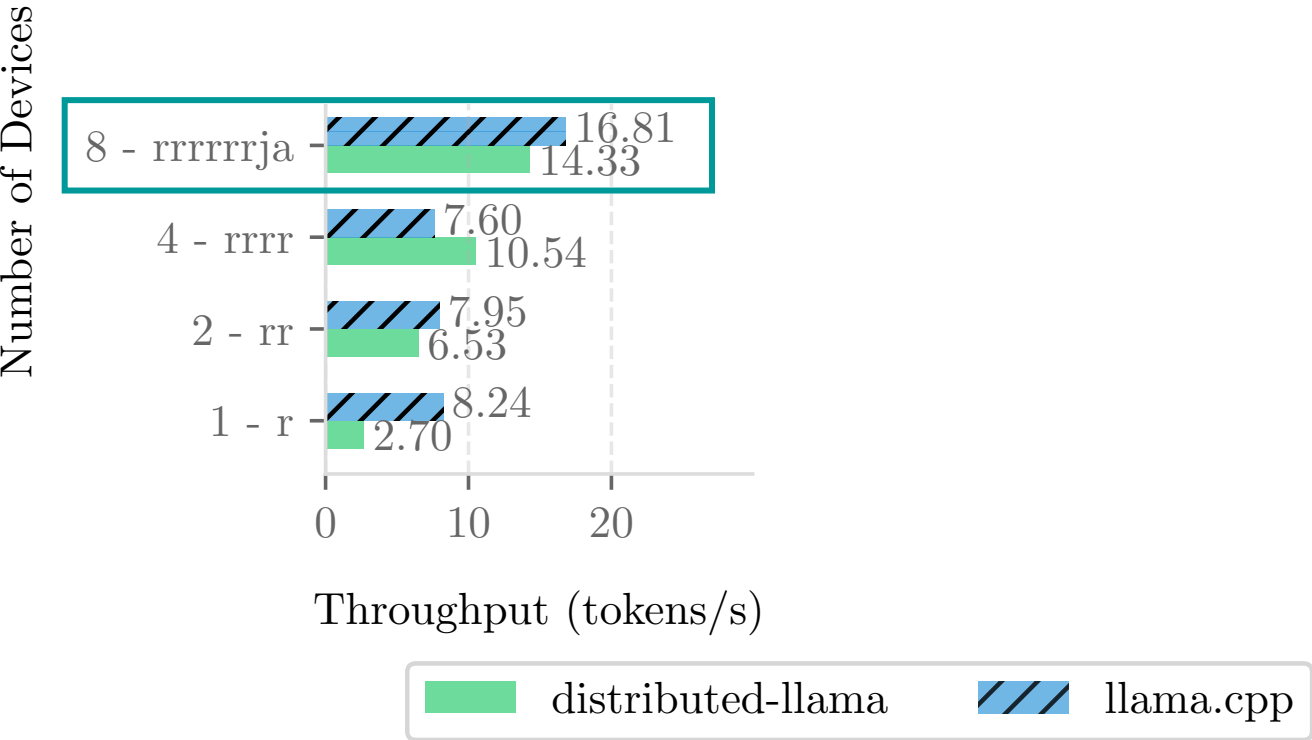


Results:

- llama.cpp:
 - More devices → Same throughput
- distributed-llama:
 - More devices → Higher throughput

Results – Scalability

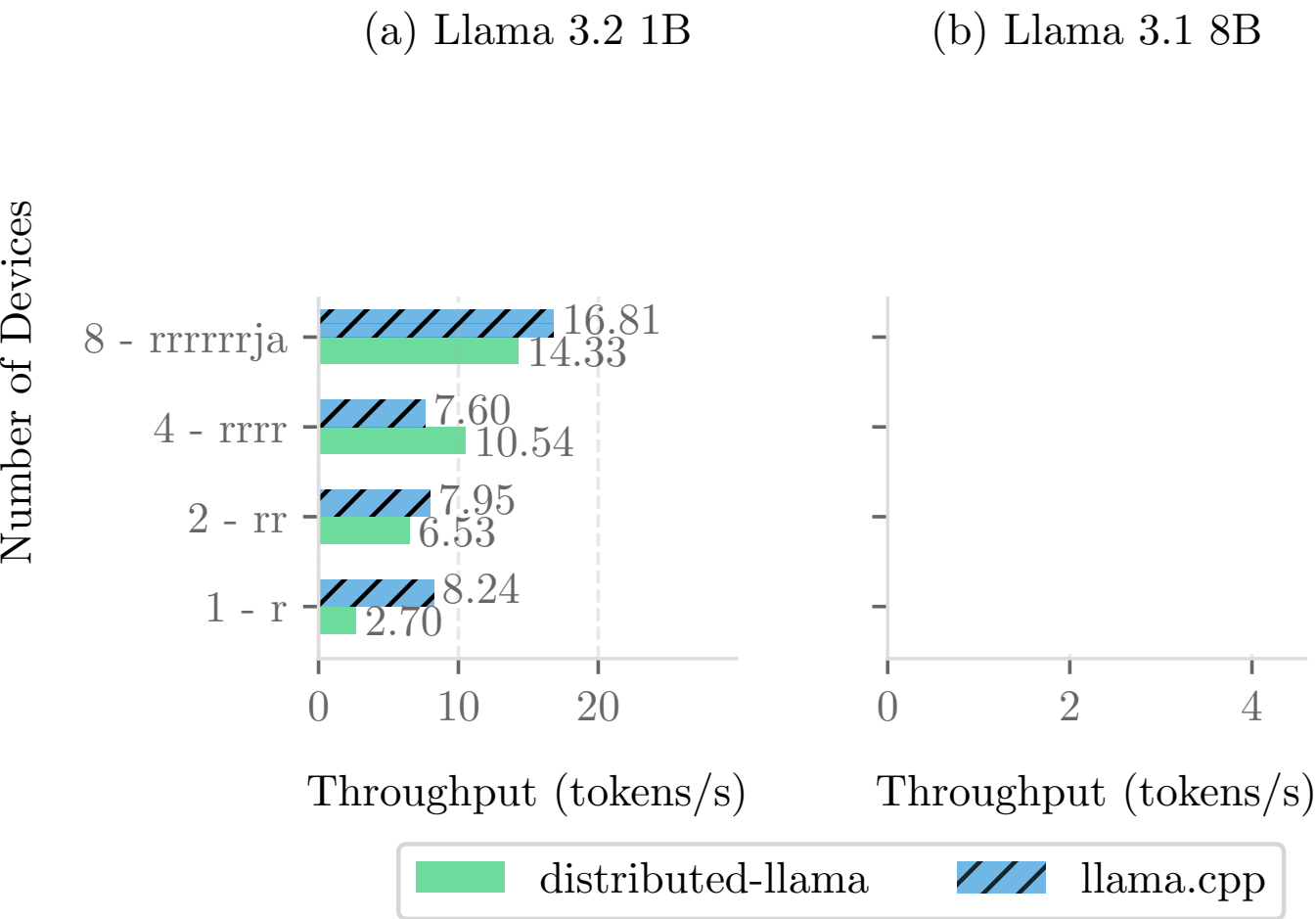
(a) Llama 3.2 1B



Results:

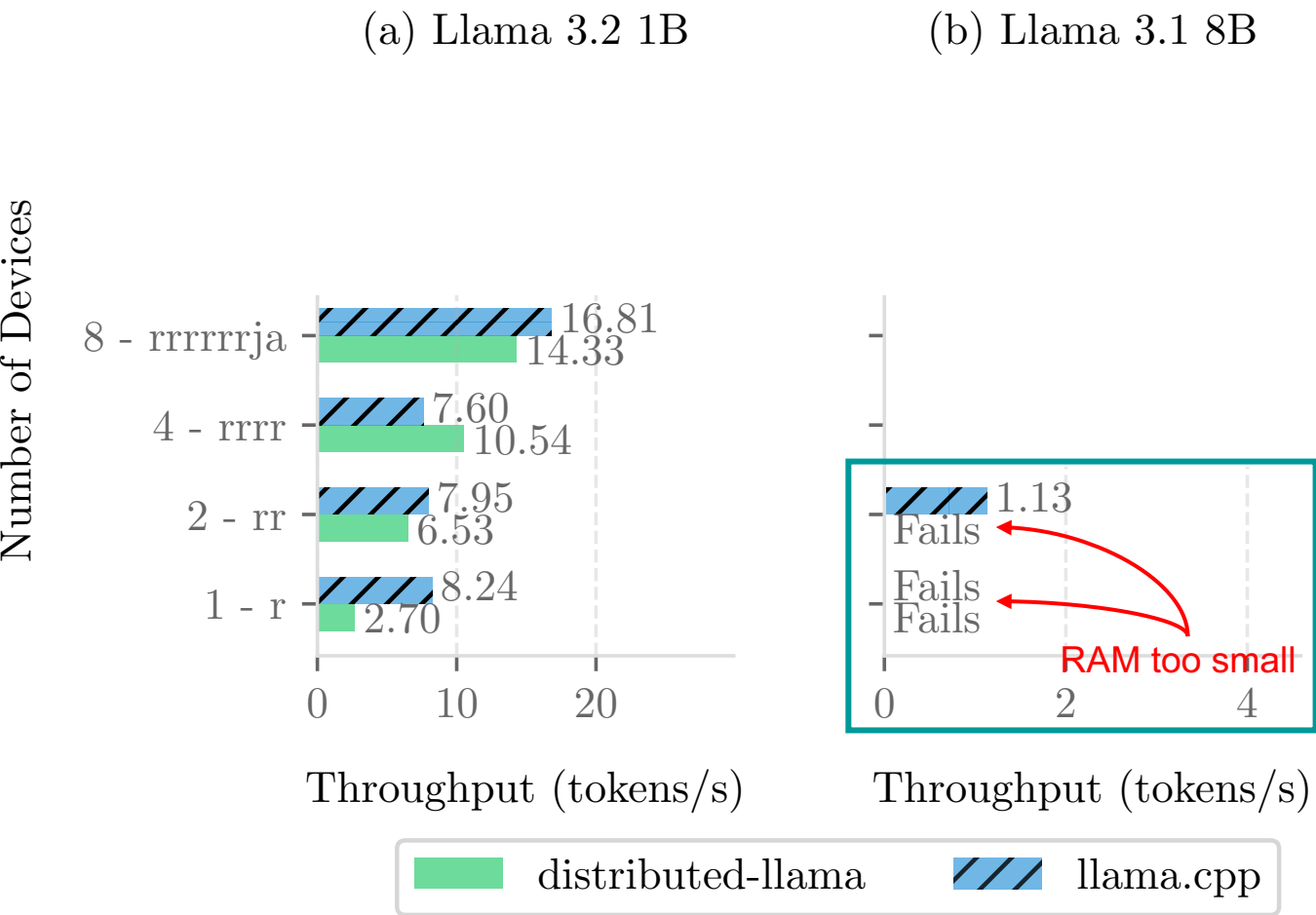
- llama.cpp:
 - More devices → Same throughput
 - Better devices → Higher throughput
- distributed-llama:
 - More devices → Higher throughput

Results – Scalability



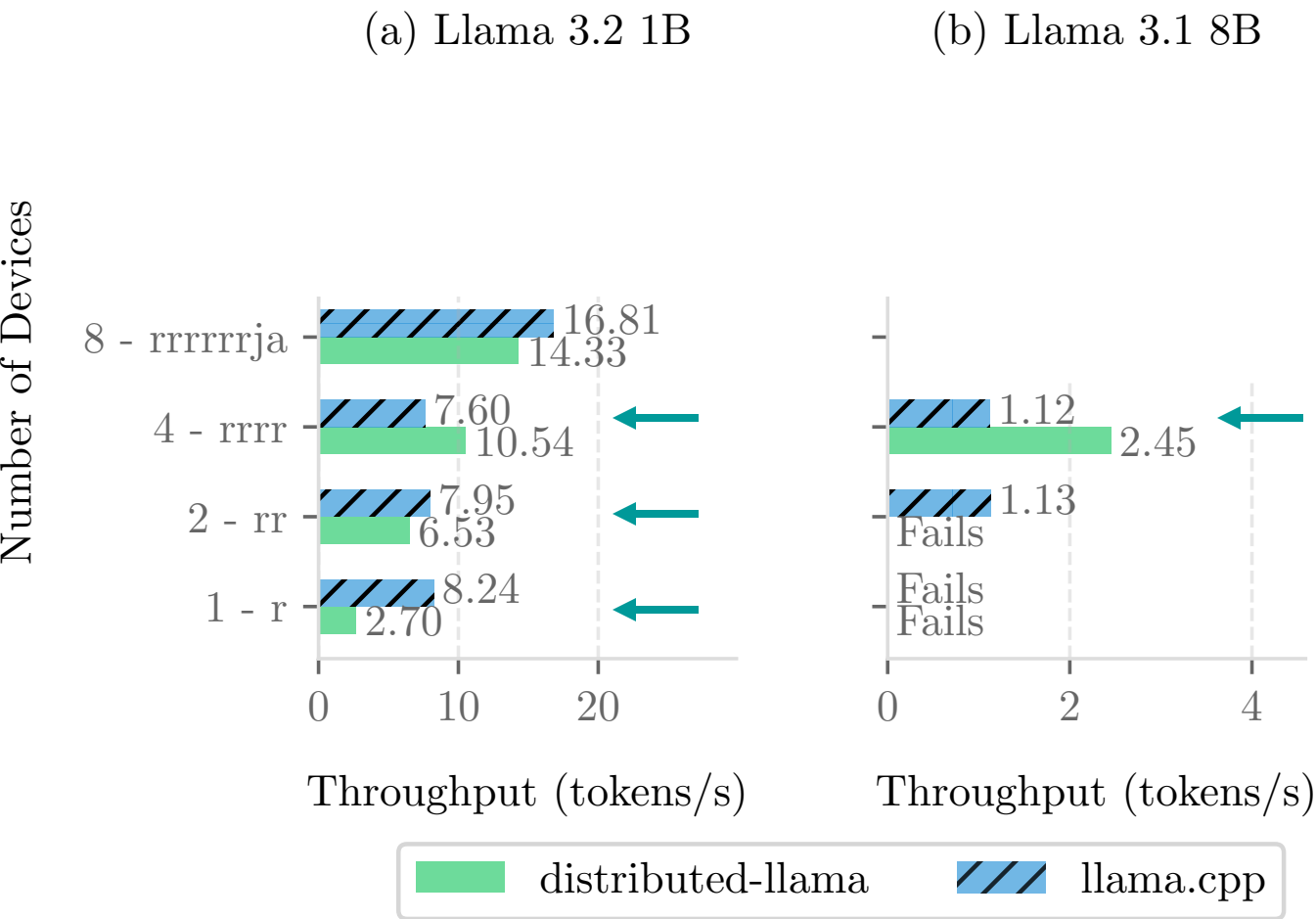
- Results:
- llama.cpp:
 - More devices → Same throughput
 - Better devices → Higher throughput
 - distributed-llama:
 - More devices → Higher throughput

Results – Scalability



- Results:
- llama.cpp:
 - More devices → Same throughput
 - Better devices → Higher throughput
 - More efficient RAM usage
 - distributed-llama:
 - More devices → Higher throughput

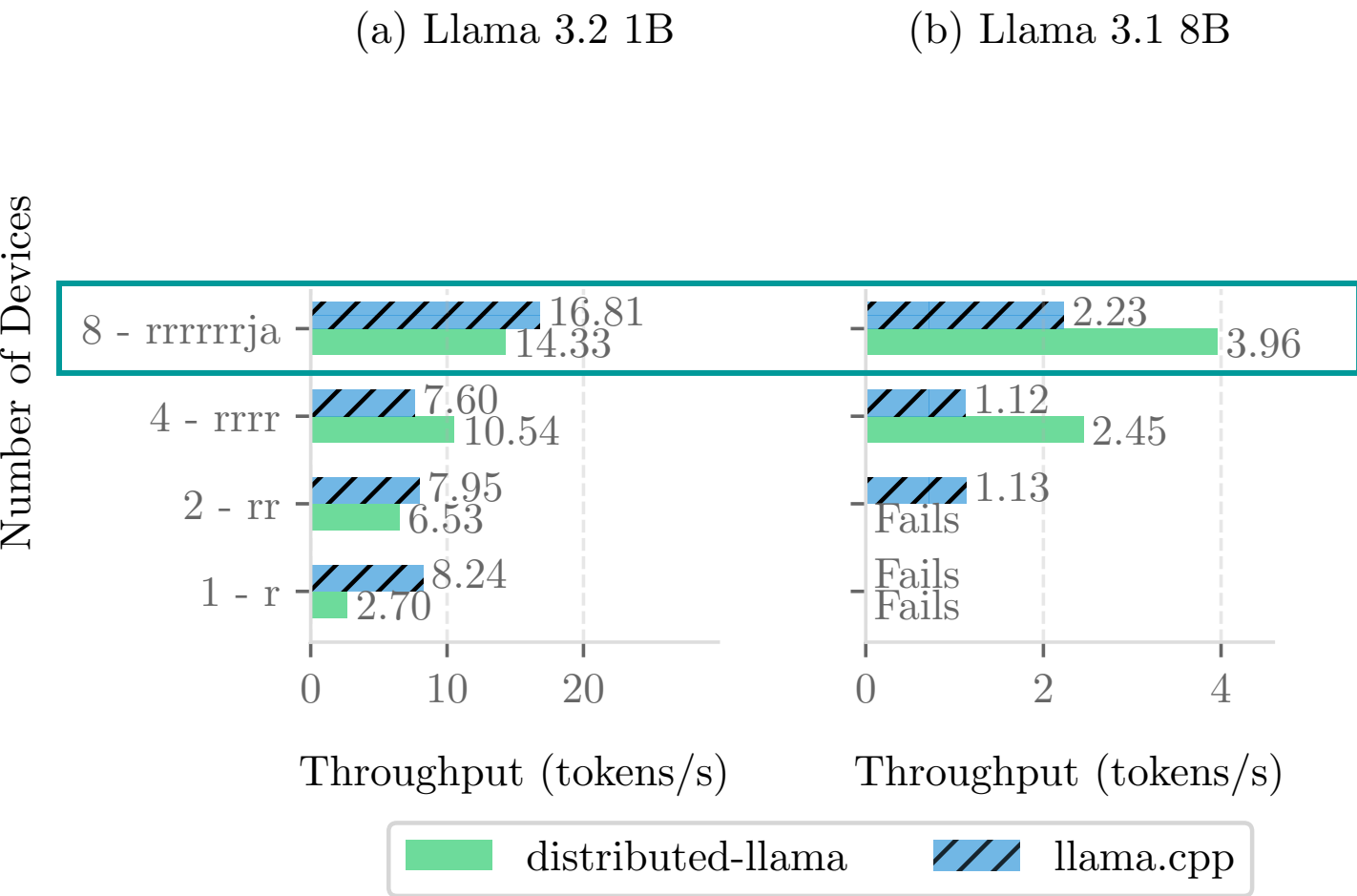
Results – Scalability



Results:

- llama.cpp:
 - More devices → Same throughput
 - Better devices → Higher throughput
 - More efficient RAM usage
- distributed-llama:
 - More devices → Higher throughput

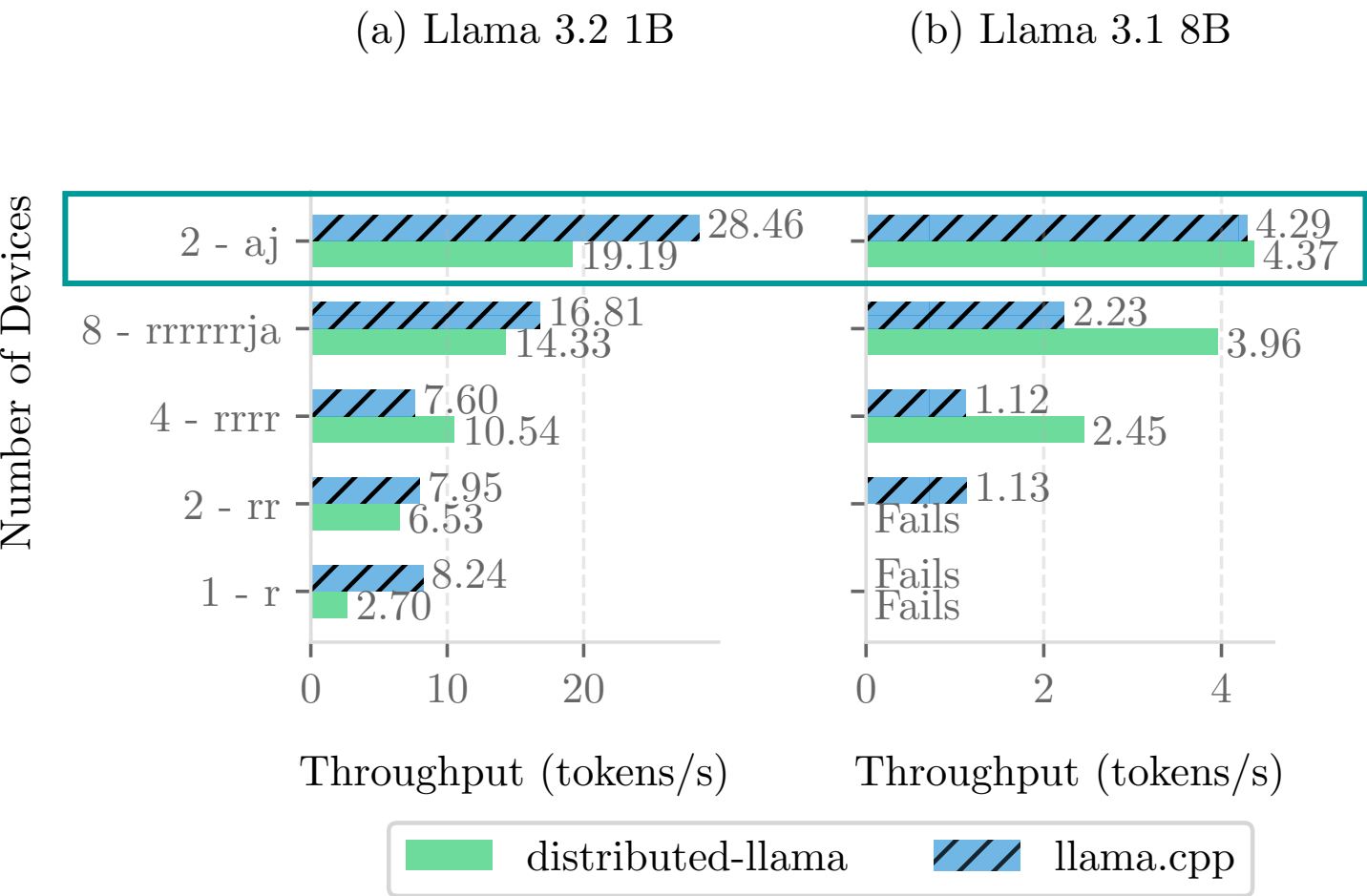
Results – Scalability



Results:

- llama.cpp:
 - More devices → Same throughput
 - Better devices → Higher throughput
 - More efficient RAM usage
 - Lower throughput for larger models
- distributed-llama:
 - More devices → Higher throughput
 - Higher throughput for larger models

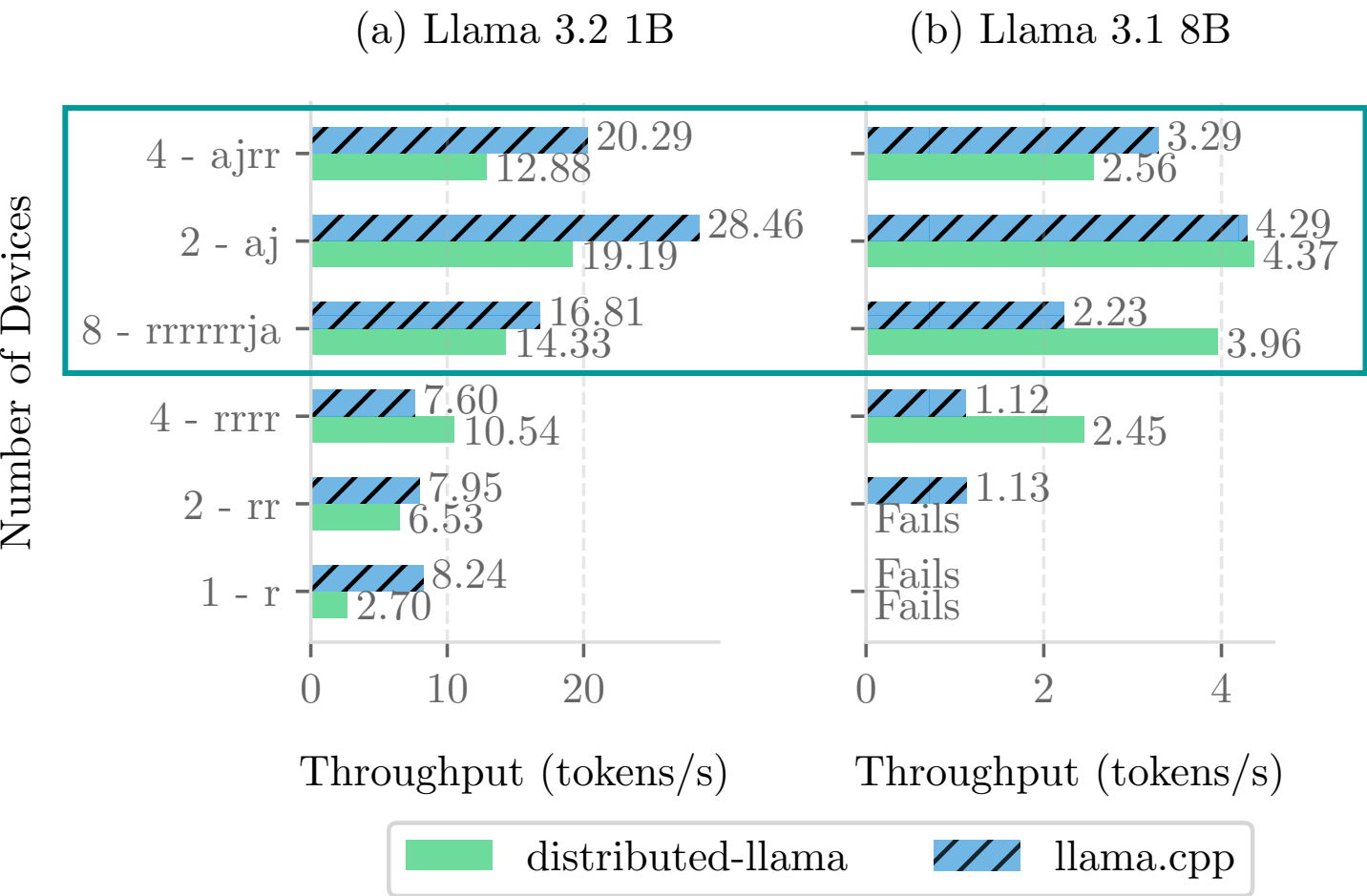
Results – Scalability



Results:

- llama.cpp:
 - More devices → Same throughput
 - Better devices → Higher throughput
 - More efficient RAM usage
 - Lower throughput for larger models
- distributed-llama:
 - More devices → Higher throughput
 - Higher throughput for larger models
- Both:
 - Better devices → Higher throughput

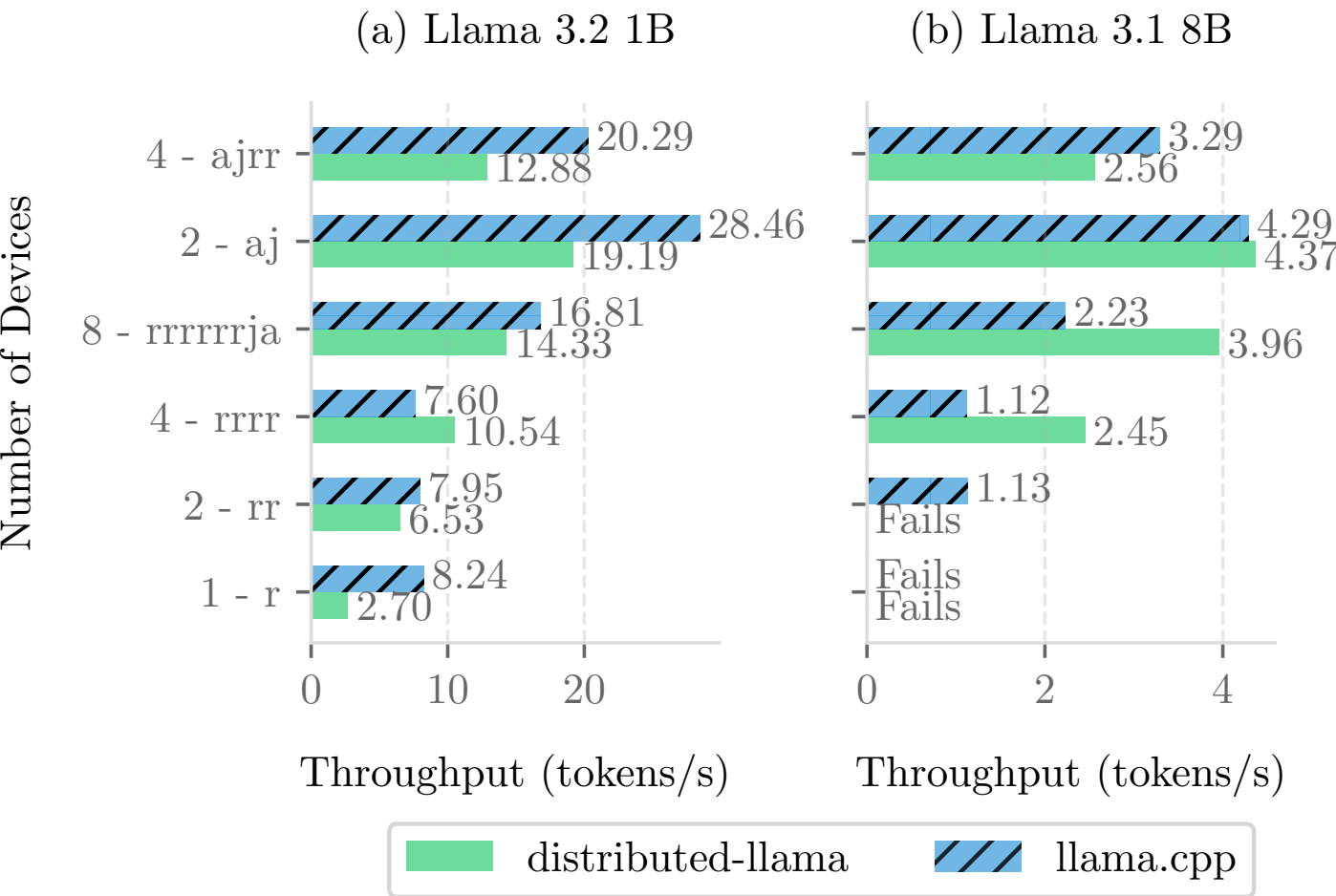
Results – Scalability



Results:

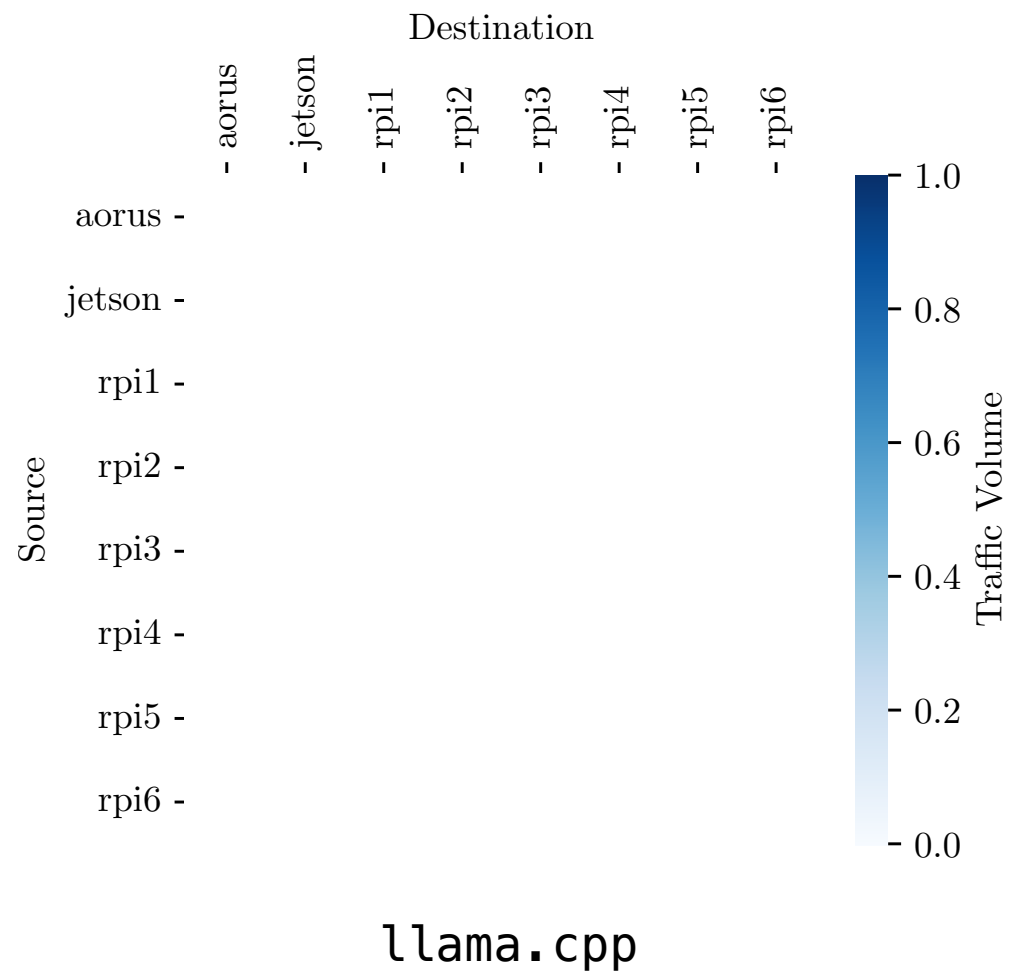
- llama.cpp:
 - More devices → Same throughput
 - Better devices → Higher throughput
 - More efficient RAM usage
 - Lower throughput for larger models
- distributed-llama:
 - More devices → Higher throughput
 - Higher throughput for larger models
- Both:
 - Better devices → Higher throughput
 - Less throughput including worse devices

Results – Scalability

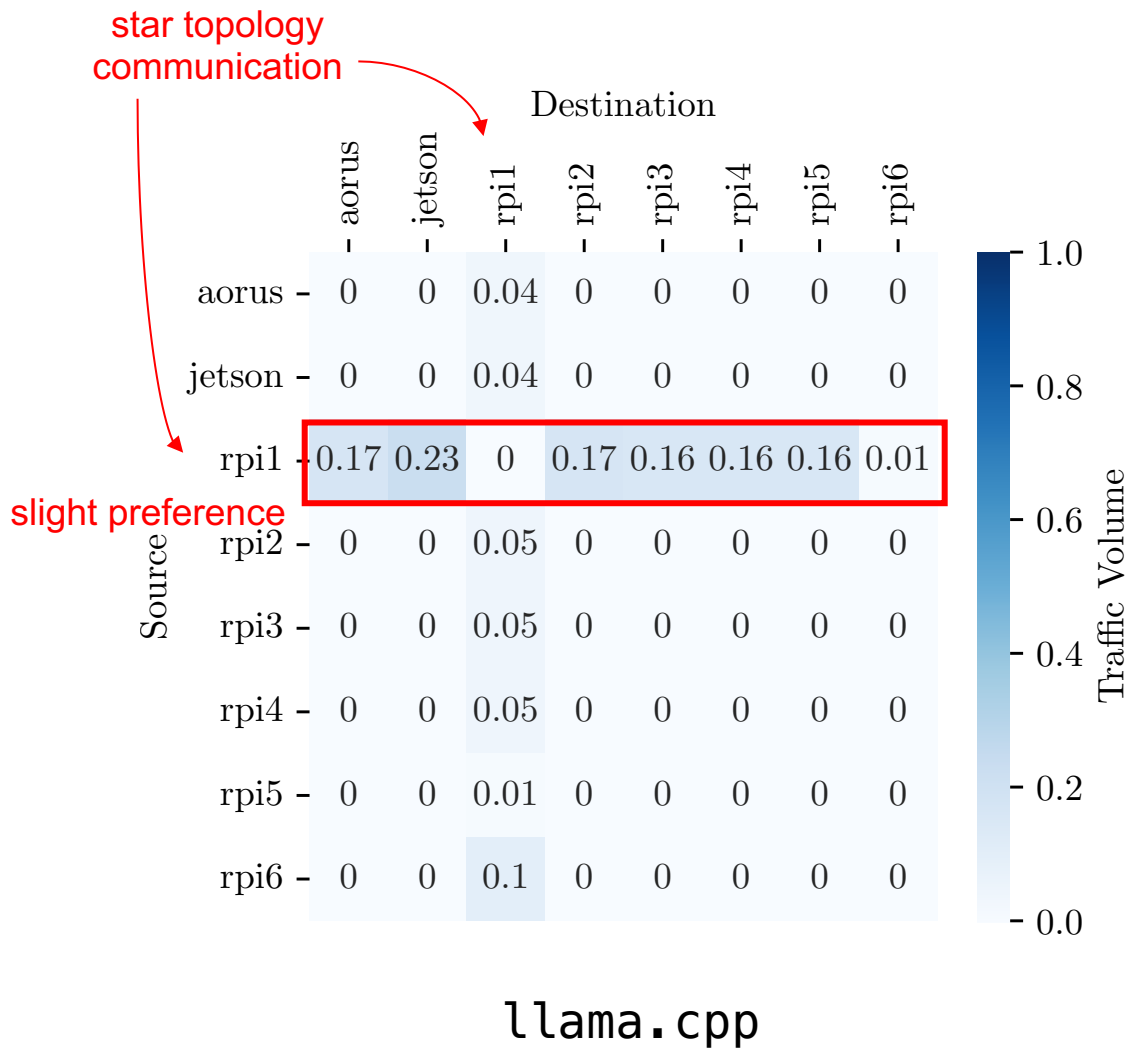


- Results:
- llama.cpp:
 - More devices → Same throughput
 - Better devices → Higher throughput
 - More efficient RAM usage
 - Lower throughput for larger models
 - distributed-llama:
 - More devices → Higher throughput
 - Higher throughput for larger models
 - Both:
 - Better devices → Higher throughput
 - Less throughput including worse devices

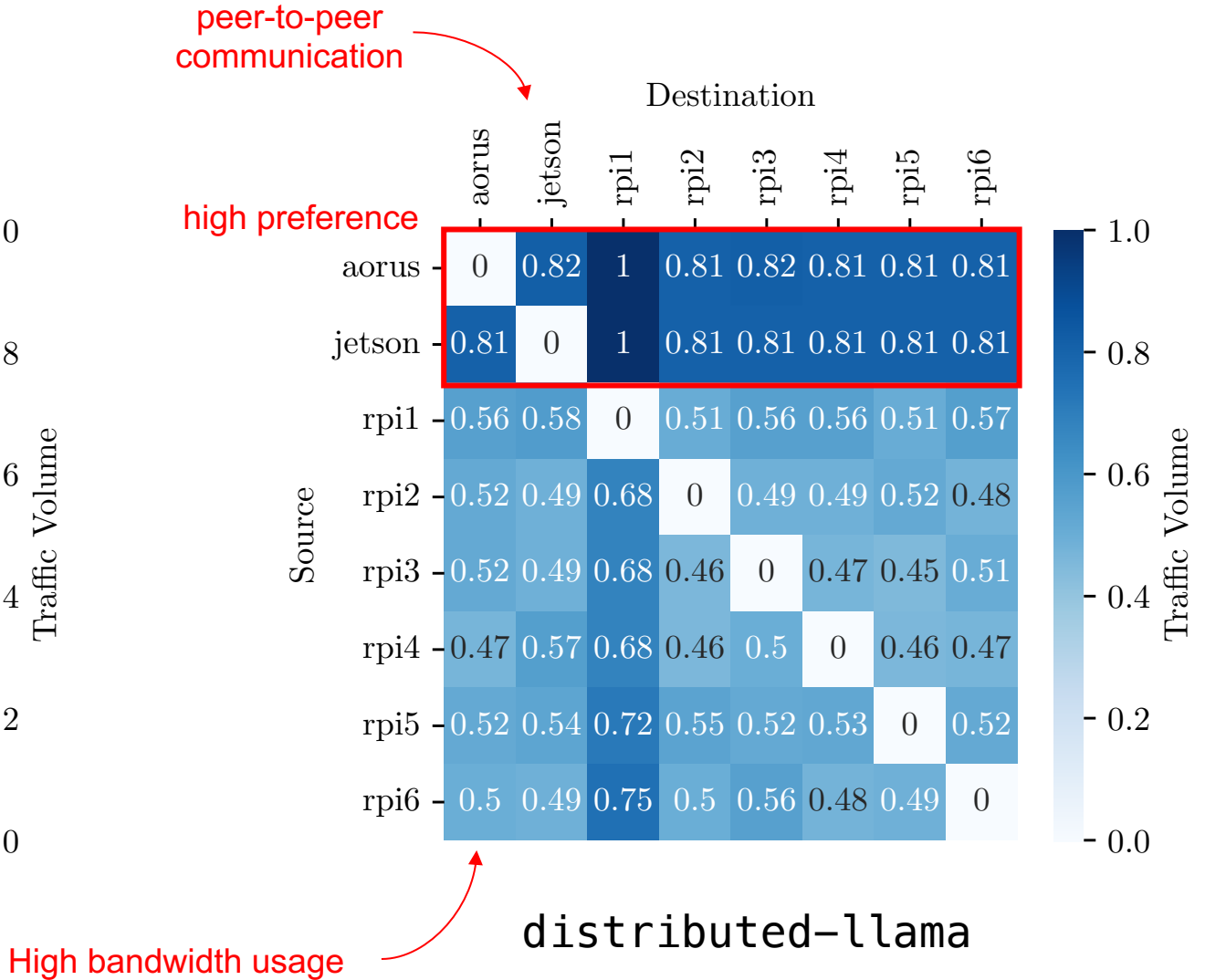
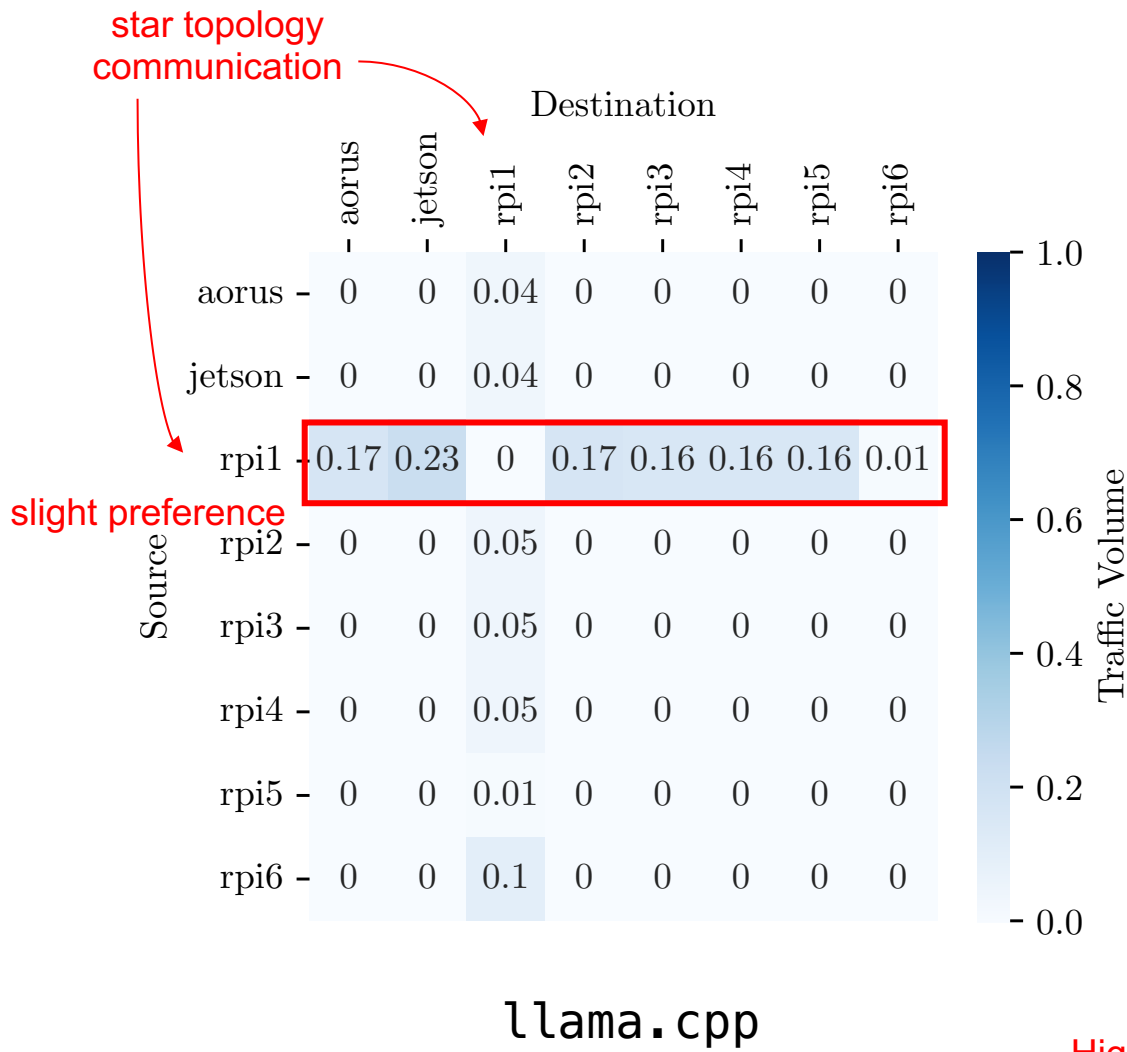
Results – Traffic Volume



Results – Traffic Volume



Results – Traffic Volume



Takeaways – It depends...

Choose your distributed LLM inference framework:

1. Heterogeneous or homogeneous cluster
2. High or low number of devices
3. Available Bandwidth
4. Inference throughput

- **llama.cpp:**
 - Few powerful, homogeneous devices
 - Lower bandwidth usage
- **distributed-llama:**
 - Maximum horizontal scalability
 - Many devices (even heterogeneous)
 - Higher throughput for large models

Machine operator
is happy



John is happy



Contact

Philippe Buschmann
PhD Candidate

Siemens AG
Foundational Technologies
Research & Predevelopment
Industrial Networks & Wireless Germany
FT RPD CED INW-DE
Friedrich-Ludwig-Bauer-Str. 3
85748 Garching, Germany

E-mail: philippe.buschmann@siemens.com

Appendix

Potential and Selected Frameworks

Pipeline parallelism

GPU support

CPU support

Tensor parallelism

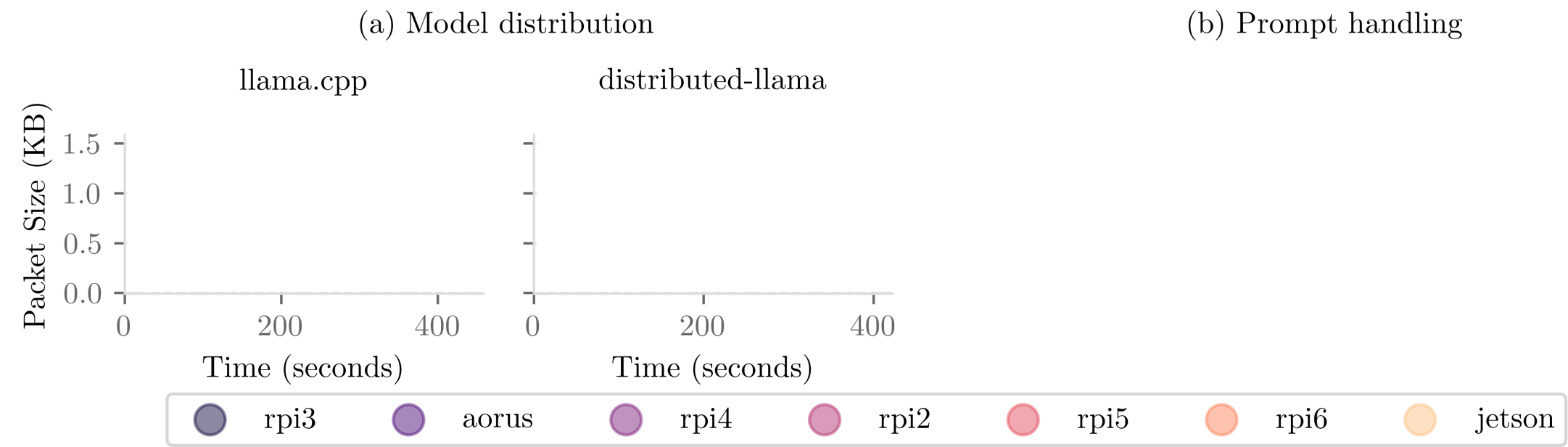
Compatible

Framework	CPU	GPU	PP	TP	C
Petals [16]	X	✓	✓	X	X
LLM-PQ [35]	X	✓	✓	X	X
HexGen [21]	X	✓	✓	✓	X
vLLM [13]	✓ ¹	✓	✓	✓	X
exo [3]	✓	✓	✓	X	X
prima.cpp [23]	✓	✓	✓	X	✓
llama.cpp [5]	✓	✓	?	?	✓
distributed-llama [31]	✓	✓ ²	X	✓	✓

¹ Not available as pre-built

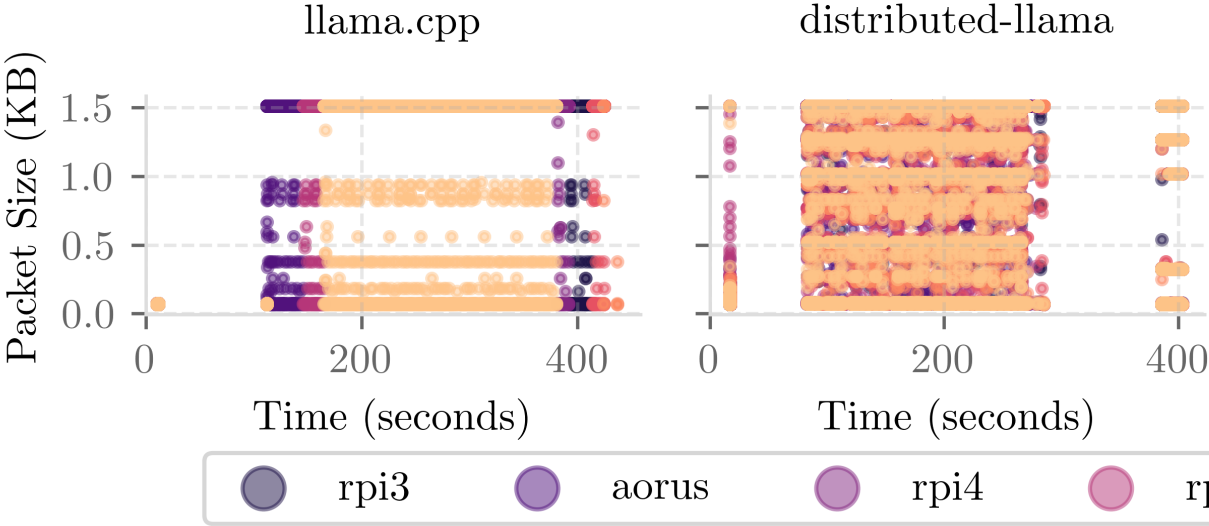
² Vulkan (experimental)

Results – Traffic Patterns



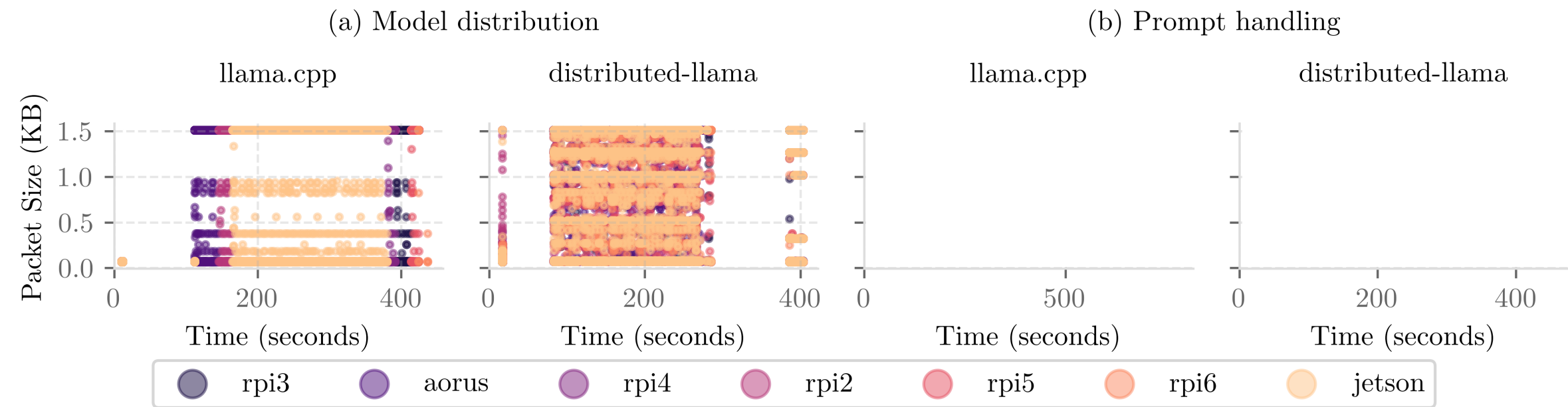
Results – Traffic Patterns

(a) Model distribution

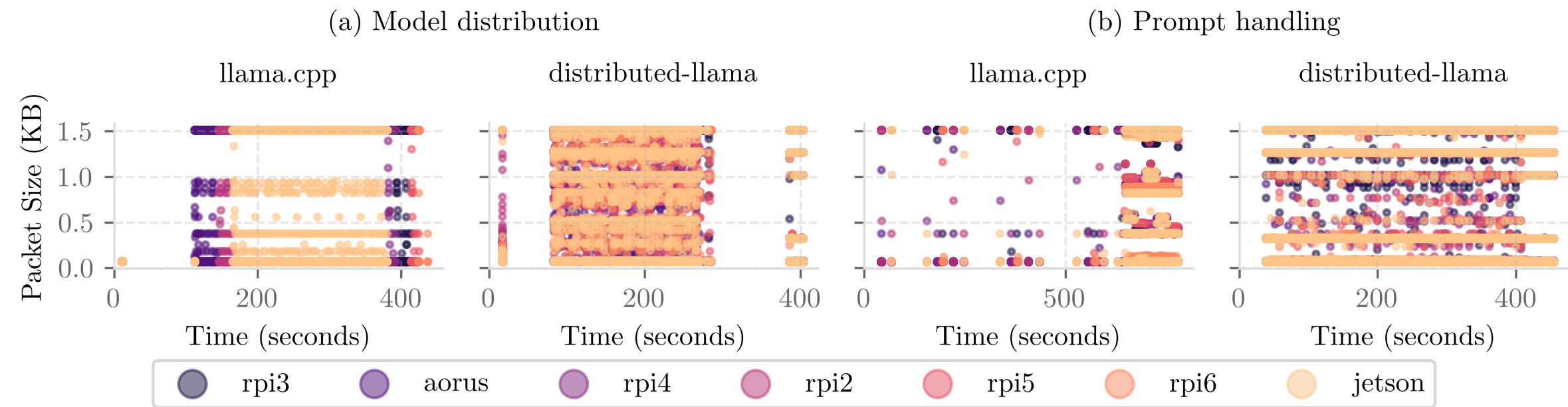


(b) Prompt handling

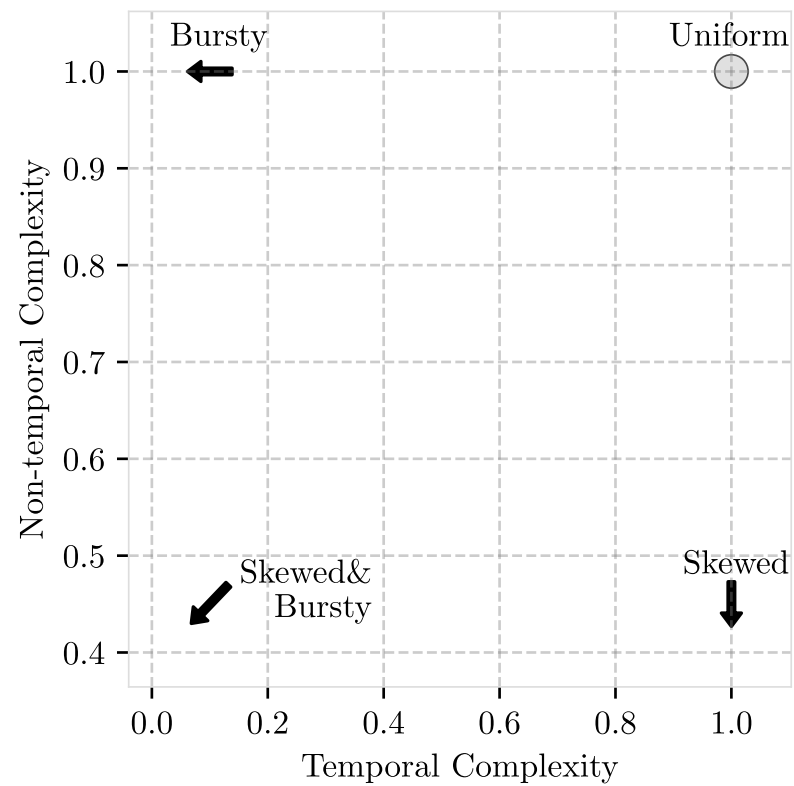
Results – Traffic Patterns



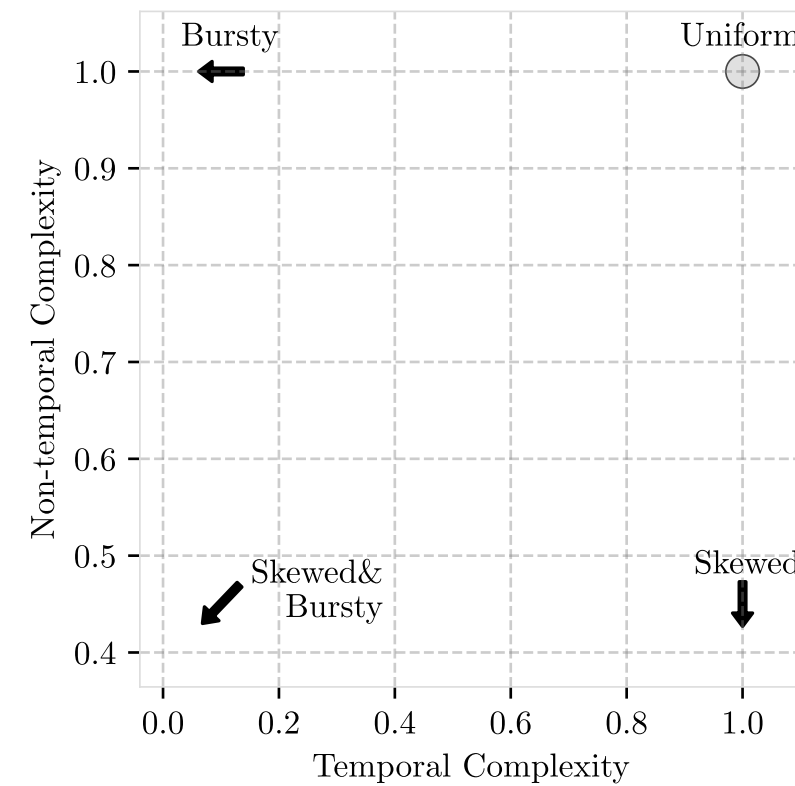
Results – Traffic Patterns



Results – Traffic Structure

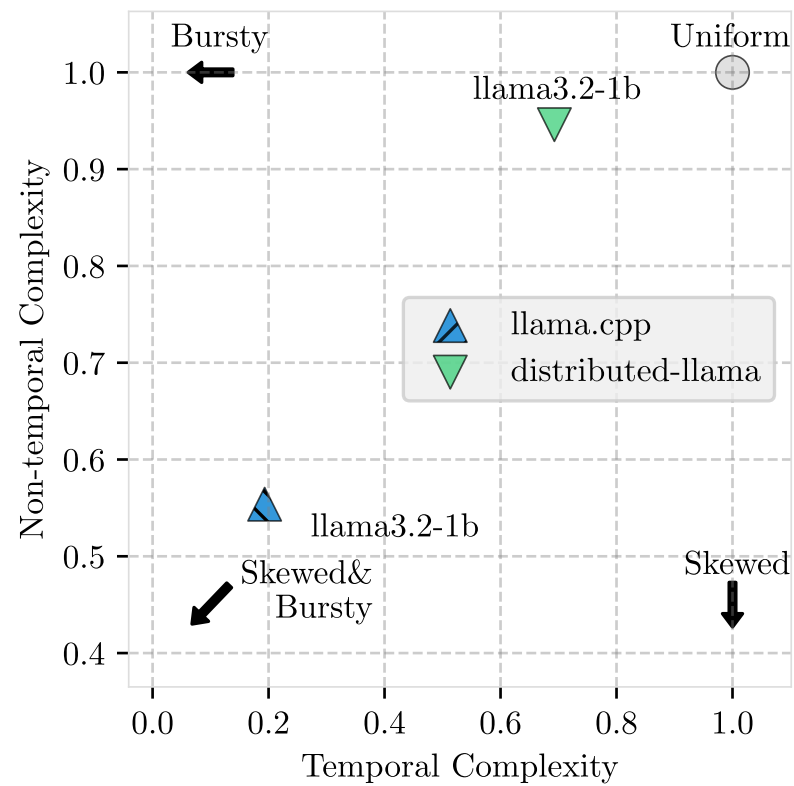


Model Distribution

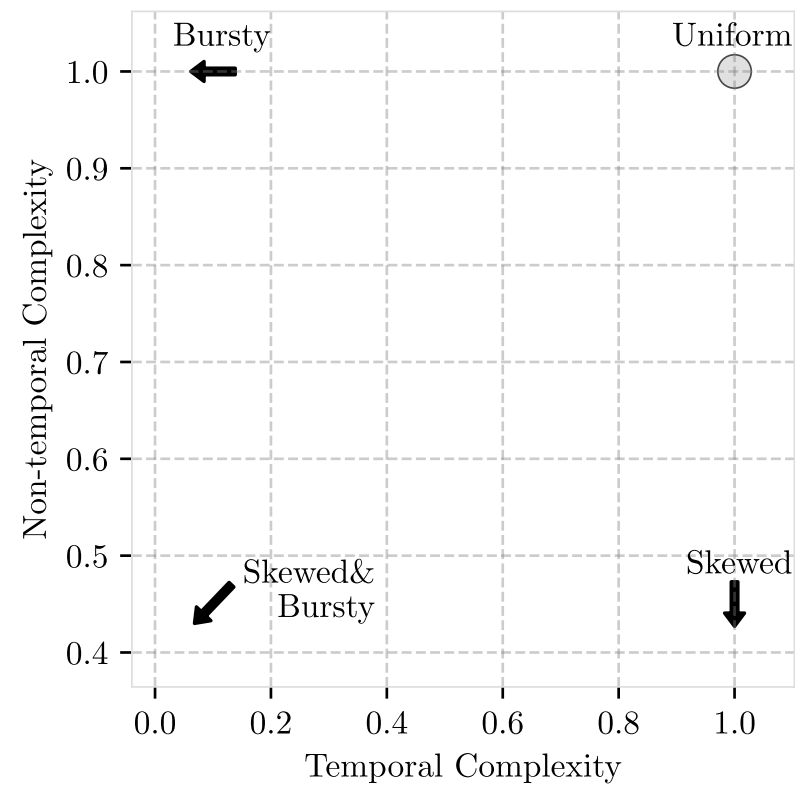


Prompt Handling

Results – Traffic Structure

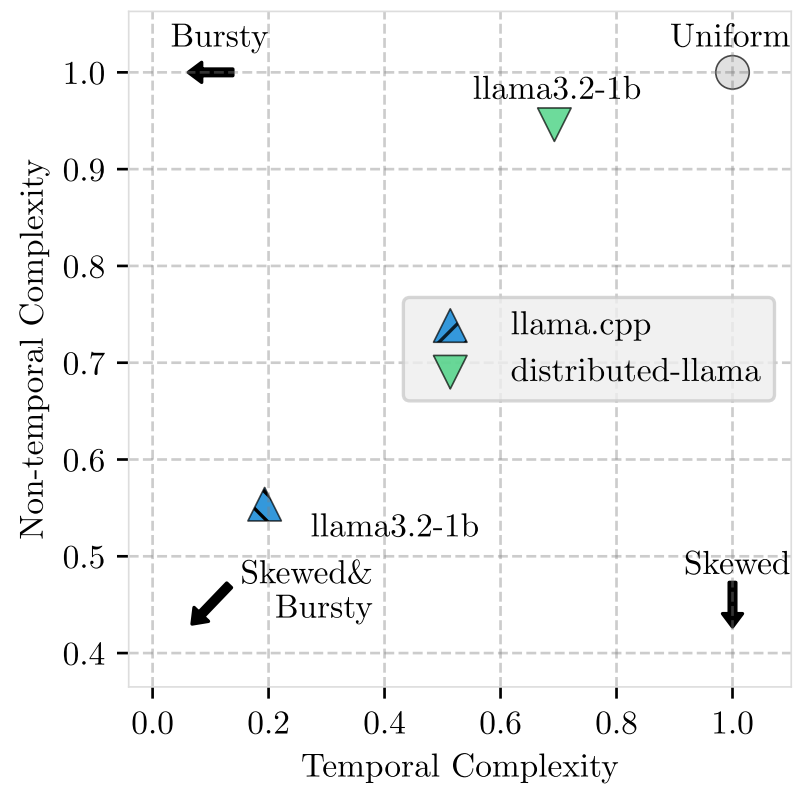


Model Distribution

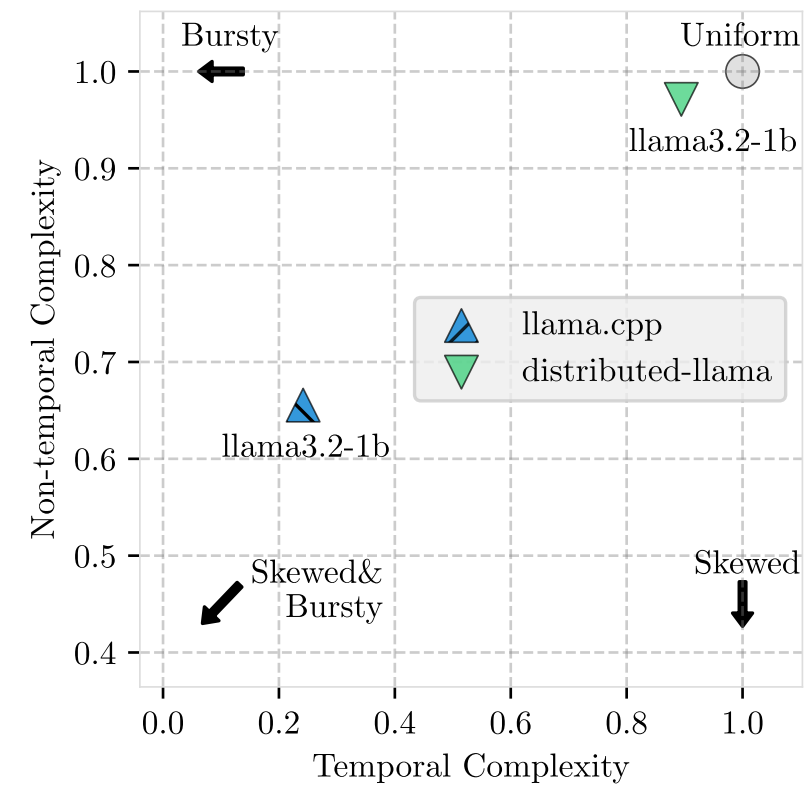


Prompt Handling

Results – Traffic Structure

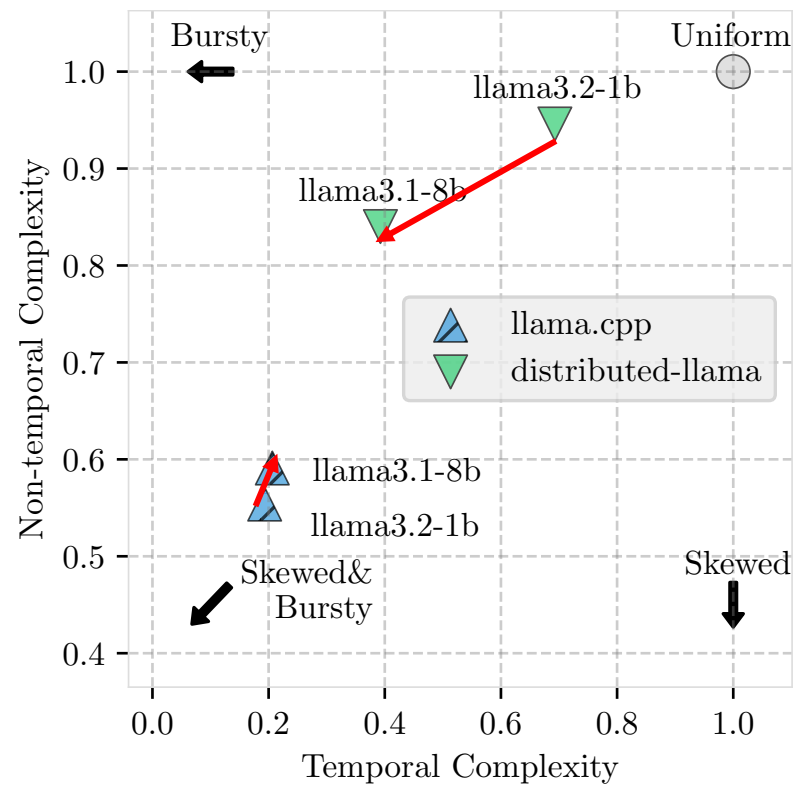


Model Distribution

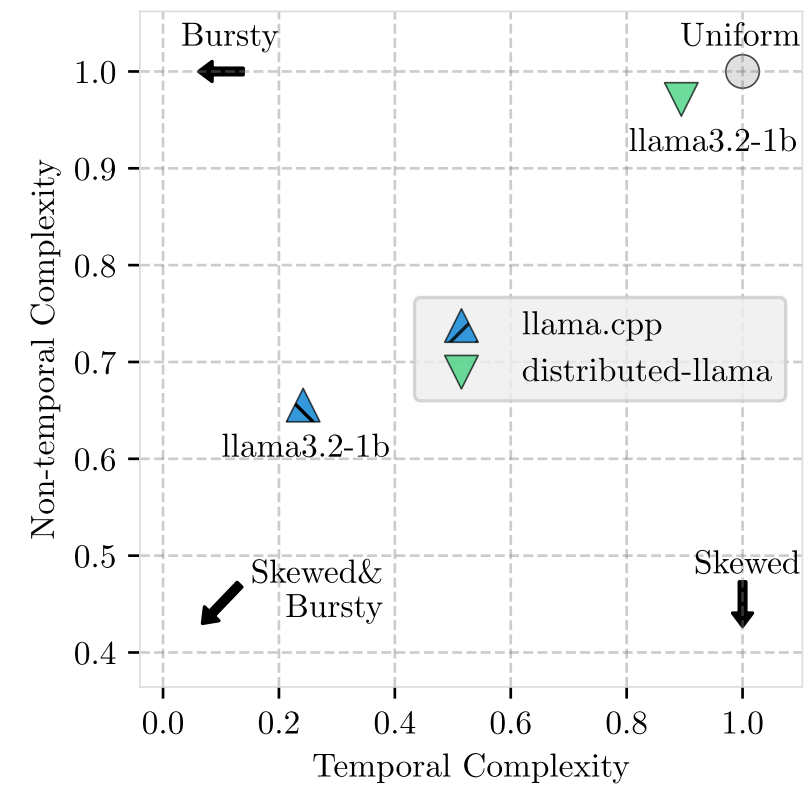


Prompt Handling

Results – Traffic Structure

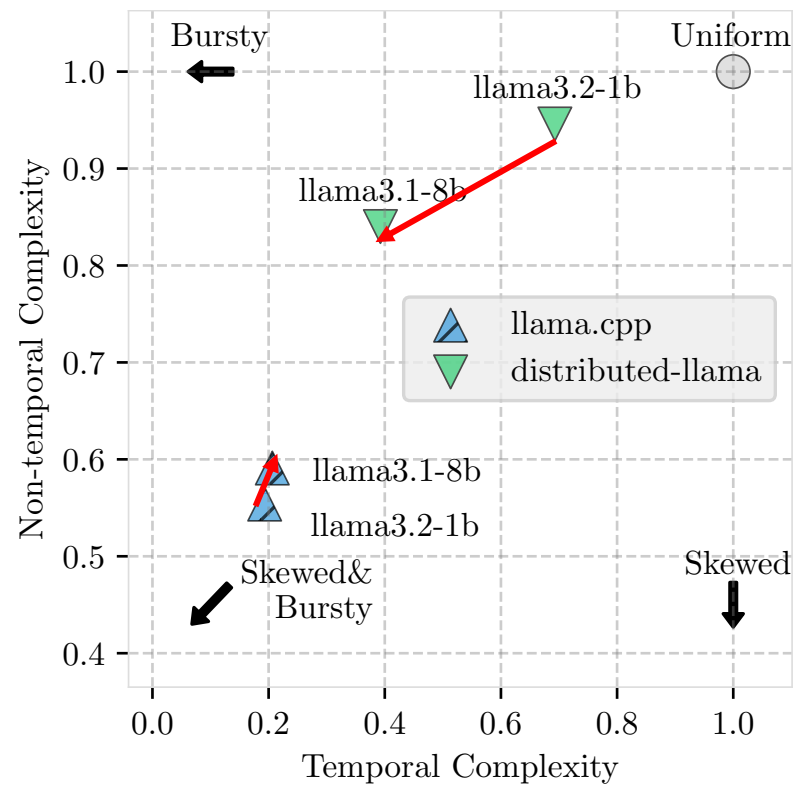


Model Distribution

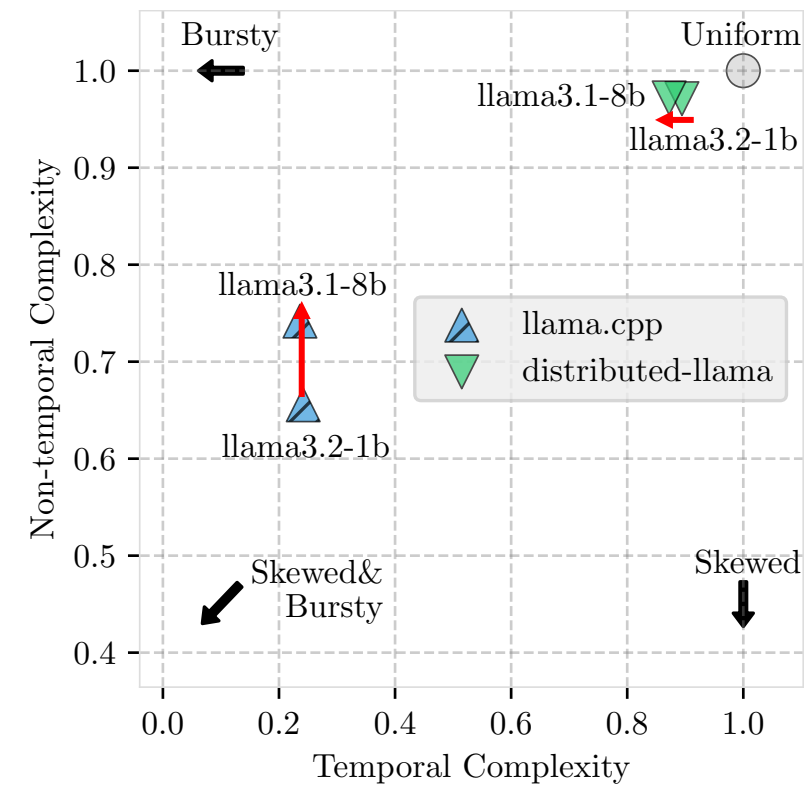


Prompt Handling

Results – Traffic Structure



Model Distribution



Prompt Handling