

Stage 1 Prompt Variation Generation

Base prompt: "Explain the working principle of transformer architecture."

Variation generation template

Definitions

- Isolation rules
- Sub-property

Constraints

- Length control
- Output format

LLM as generator



Prompt candidates

Var 1 "Could you explain how the Transformer architecture works and describe its underlying working principle?"

Var 2 "Explain the working principle of the Transformer architecture, including how self-attention ..."

...

Stage 2 Prompt Variation Validation

Prompt variation scoring

Property scoring

	P ₁	P ₂	P ₃
S ₁	■	■	■
S ₂	■	■	■

Semantic similarity

Var ₁	■	■	0.60
Var ₂	■	■	0.80

Candidates filtering

(higher is better)

Property score $\geq T_A$
Semantic similarity $\geq T_B$

Final prompt dataset

Base prompts
Prompt variations

Stage 3 On-Device Inference Profiling

Lightweight LLMs

Deployment

Profiling metrics

Energy-quality pareto



Runtime: MLC-LLM
Mobile devices



- Latency / power
- Prefill / decode
- Rel., Corr., Coh.
- Cmp., Inst., Cons.



System Implementation

Host server setup

- Prompt variation generation and validation
- Profiling log collection and analysis



Stage 1

Stage 2

Inference results collection

Response quality 0-1, 6 metrics

Rel.	■	.91
Corr.	■	.89
Coh.	■	.90
Cmp.	■	.87
Inst.	■	.83
Cons.	■	.96

Validated prompts

Prompt dataset

"Use a numbered list..."

On-device inference

On-device LLM setup



Mobile devices

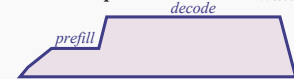
- Pixel 7
- Pixel 8 Pro

LLMs

- Gemma-2-2B
- Llama-3.2-1B
- Qwen-2.5-0.5B
- Qwen-2.5-1.5B
- SmolLM2-360M

Stage 3

Inference power



Inference latency

