

LLM RAG System - Benchmark Report

Version: 1.0.1

Date: 2026

Abstract

This report evaluates the MVP RAG system on two hardware configurations to validate the AMD GMKtec EVO-X2 path as the optimal RAG LLM platform.

The data from our tests shows that the AMD GMKtec delivers 4.2x faster performance compared to a laptop baseline. With the data showing clear advantages such as improved quality, lower latency, and overall better user experience. Hardware is justified.

Metric	Laptop	GMKtec EVO-X2	Improvement
Average Latency	64.46s	18.18s	3.5x faster
Throughput	4.6 tok/s	19.21 tok/s	4.2x faster
Quality Score	0.446	0.532	+19%

Materials and Methods

Hardware:

- Laptop: 16GB RAM, integrated graphics (low-power)
- GMKtec: AMD Ryzen AI Max+ 395, 64-128GB RAM, Radeon 8060S iGPU

Software:

- LLM: llama3:latest (8B)
- Embeddings: nomic-embed-text
- Vector DB: ChromaDB
- Framework: LlamaIndex
- Code: Python

Configuration Settings:

- Chunk size: 1024 tokens
- Top-K retrieval: 5 chunks
- Similarity threshold: 0.3

Results

Query Latency

Query Type	Laptop	GMKtec	Difference
Simple	43.06s	11.00s	3.9x
Complex	92.79s	20.64s	4.5x
Average	64.46s	18.18s	3.5x

User Experience Impact:

- Laptop: Unusable (1+ minute wait time for responses)
- GMKtec: Acceptable (10-20 second response time)

Throughput

- Laptop: 4.6 tokens/sec (1x baseline)
- GMKtec: 19.21 tokens/sec (4.2x faster)
- Latency range: 11s - 20.6s (AMD GMKtec EVO-X2)
- Latency range: 43s - 92.7s (Laptop)

Retrieval Quality

Metric	Laptop	GMKtec
Avg Similarity Score	0.393	0.397
Keyword Match	50%	66.67%
Average Quality Score	0.446	0.532

Finding: The AMD GMKtec hardware yields better retrieval quality, and highly improved speeds.

Limitations

What We Didn't Test:

- Hallucination rate
- Large context windows (32K+ tokens)
- GPU/NPU utilization

Test Coverage:

- Only 5 latency queries
- Only 5 quality queries
- Larger sample sizes and more diverse queries needed for more robust statistical testing.

Appendix: Raw Data Excerpts

Laptop Benchmark:

```
{  
  "avg_latency_sec": 64.46,  
  "avg_throughput_tokens_per_sec": 4.6,  
  "avg_quality_score": 0.446  
}
```

GMKtec Benchmark:

```
{  
  "avg_latency_sec": 18.18,  
  "avg_throughput_tokens_per_sec": 19.21,  
  "avg_quality_score": 0.532  
}
```
