

Report summary

Safetensors unknown arch | F32 | 4 tensors | ~0M params | 0 MB | 1 blocking issue / 2 warnings / 3 informational

Model

Filename	overlapping-tensors.safetensors
Format	safetensors
Architecture	—
Size	2.1 KB
SHA-256	8ed50ec8cdb4027f2f63f46afeeb81233e7919e2463a336452e42515092e20ff
Scan mode	static-deep
Scanned at	2026-08-10T20:38:48+00:00Z
Report ID	670955388bd6489ba2ad1de006f9a2e5

Findings (6)

SEVERITY	CATEGORY	DETAIL
CRITICAL	Integrity	Overlapping tensor data regions (1 pairs) Pairs: [('blk.0.attn_q.weight', 'blk.0.attn_k.weight')] Recommendation: Overlapping regions enable aliasing attacks where two tensors share memory; reject this file
MEDIUM	Memory Alignment	Misaligned tensor data (1 tensor(s)) Tensors not aligned to 8-byte boundary: ['blk.0.attn_k.weight']. Misalignment can cause SIGBUS on ARM and degrade SIMD performance on x86. Likely indicates manual file editing or a non-standard exporter. Recommendation: Re-export using the official safetensors library to ensure correct alignment.
LOW	Provenance	Sparse __metadata__: missing author, license, source/model_name Embedded metadata has 1 key(s) but lacks: author, license, source/model_name. Without provenance, the model origin cannot be verified from the file alone. Recommendation: Add author, license, and source fields to __metadata__ when exporting.
INFO	Precision	Full-precision weights (F32) — not quantized All 4 tensors use 32-bit float. Higher numerical fidelity but 2× the memory of F16 and 4-8× of quantized variants. Recommendation: Consider converting to BF16 or INT8 for production inference if accuracy allows.
INFO	Structure	Dtype breakdown: {'F32': 4} 4 tensors, ~448 parameters Recommendation: Verify dtype matches expected model precision
INFO	Format	Safetensors format provides strong sandboxing Unlike pickle-based formats, safetensors cannot execute arbitrary code during loading Recommendation: Safetensors is the recommended format for sharing model weights safely

Runtime compatibility

runtime	supported	notes
HuggingFace Transformers	No	—
vLLM	No	—
llama.cpp	No	—
ONNX Runtime	No	—

Generated 2026-08-21 06:51 UTC by LLMScan.online (engine v2.0.0, report schema v1.0). Full interactive report: <https://llmscan.online/report/670955388bd6489ba2ad1de006f9a2e5> — Static analysis only; not a guarantee of safety. See llmscan.online/docs/how-it-works for methodology.