



MEASURED · not asserted

THE LAB · A RESEARCH LANE

The gate: the proof is the product

Not a test you run at the end — the moat you build on. Native equals interpreter on every compile, and it is never disabled. Verifiability self-scores 9/10. Prompt injection becomes a type error. Everyone else picks two of {long-context, quantized, any-chip, proven}. This is the stack that is all four.

A measured research lane

● Shipped — measured, gate-clean

The gate: the proof is the product

Not a test you run at the end — the moat you build on. Native equals interpreter on every compile, and it is never disabled. Verifiability self-scores 9/10. Prompt injection becomes a type error. Everyone else picks two of {long-context, quantized, any-chip, proven}. This is the stack that is all four.

01 The claim, in brief

02 Native equals interpreter, always

03 The compounding axes

04 The honest bound

05 Why it ladders back

06 Sources, glossary & the honesty ledger

Every claim carries its tier · nothing faked

PREPARED

2 July 2026

VERIFIED BY

Jon Ross — Founder, Vocabotics —
15 years building safety-critical
systems

HOW TO READ IT

Drafted with AI, verified by a
named human. Sources at the back.

01 The claim, in brief

Everyone else picks two of {long-context · quantized · any-chip · proven}. We are the only stack that is all four.

The costly signal

The measured result

Every number carries its own honesty tier — the tiering is the credibility.

native $\equiv$ interp

on every compile, never disabled

● SHIPPED

9 / 10

verifiability self-score

● SHIPPED

type error

what injection becomes

● SHIPPED

4 of 4

compounding axes, at once

● SHIPPED

Most labs treat verification as a test: run it at the end, hope it passes, ship. We treat it as the *foundation* — the thing everything else is built on and can never be switched off. That inversion is the moat. A rival can copy a fast kernel; they cannot copy a proof they do not have.

02 Native equals interpreter, always

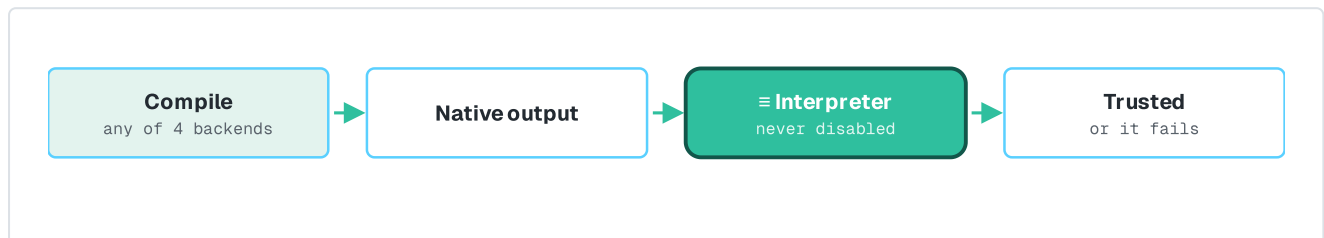
● SHIPPED — MEASURED, GATE-CLEAN

The rule, stated plainly: on **every compile**, the native output must equal the interpreter's output, and the check is **never disabled**. ● Combined with diverse double-compiling (build with `gcc`, build with `clang`, compare bytes), this closes the trusting-trust gap: the artefact is verified against an independent oracle, not merely trusted because a tool produced it.

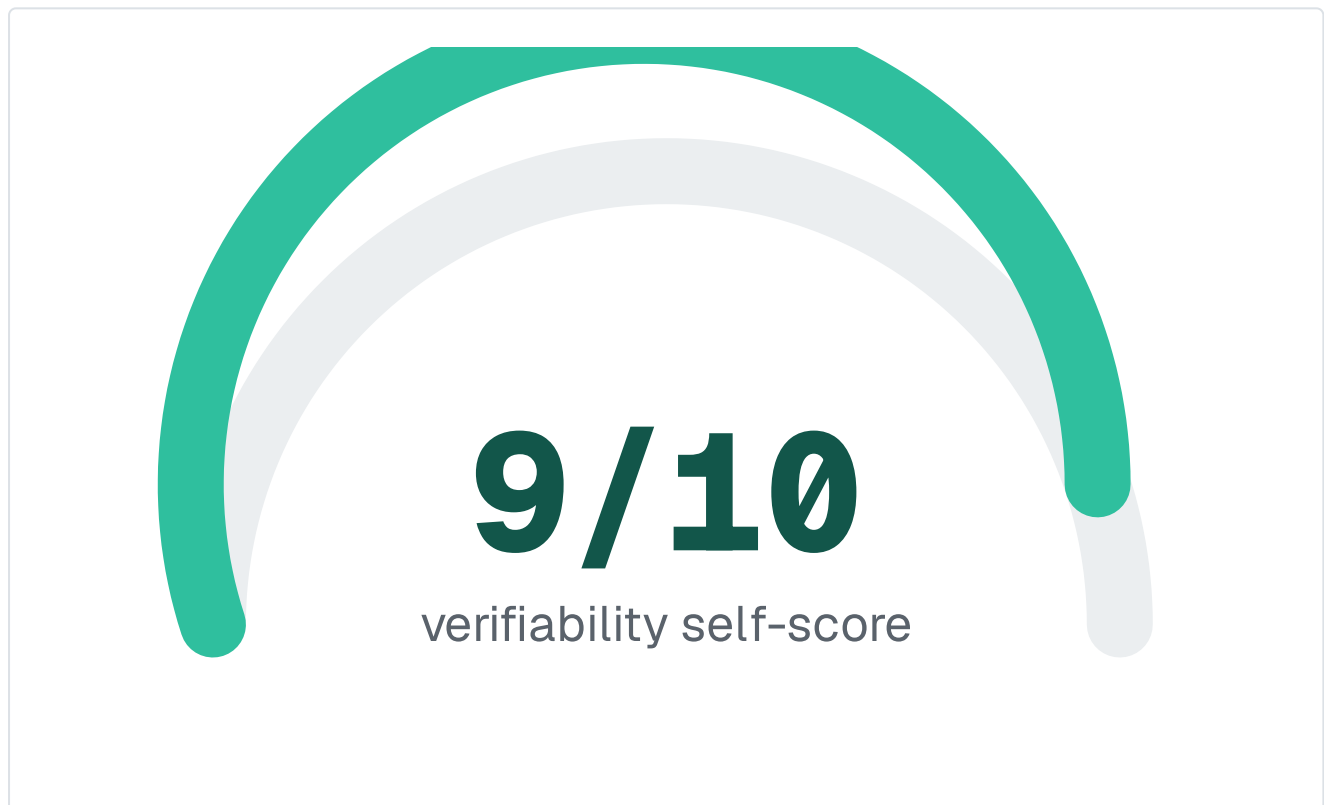
Two consequences fall out:

Verifiability self-scores 9/10. ● Against our own rubric — and we label it as a self-score, not an external audit — the stack is close to fully verifiable end to end.

Prompt injection becomes a type error. ● When capability and provenance are carried in the type system, an instruction that should not be trusted is not a policy violation to be detected later; it fails to type-check. The attack does not get blocked at runtime — it does not compile.



The rule that is the moat: on every compile the native output must equal the interpreter's, and the check is never disabled. Combined with diverse double-compiling, the artefact is verified against an independent oracle, not merely trusted.



Verifiability self-scores 9/10 — against our own rubric, and we label it as a self-score, not an external audit.

03 The compounding axes



SHIPPED – MEASURED, GATE-CLEAN

Here is the precise "best" claim, and the reason it is defensible: there are four properties everyone wants at once, and almost no stack has all four.

AXIS	WHAT IT MEANS	THE PROOF
Long-context	Flat cost as context grows	2× llama.cpp at 128k, flat line
Quantized	Runs small, on cheap silicon	ternary weights, int4 fused decode
Any-chip	Not tied to one vendor	4 backends gate-clean
Proven	Verified, not just tested	native ≡ interpreter, never disabled

Most stacks pick two — fast *or* proven, small *or* long-context. Winning all four *at once* is the claim we make, and it is the only "best" the numbers support.

We do not say 'fastest'.

Fastest is a single-axis race we do not win (we are 92.8% of cuBLAS). The compounding axes are a different, harder, more honest claim — and one a competitor cannot knock, because it is true and specific. Say the true rare thing, precisely.

	LONG-CONTEXT	QUANTIZED	ANY-CHIP	PROVEN
Most stacks				
This stack				

Go Check it Keep a human

The precise 'best' claim: four properties everyone wants at once. Most stacks pick two — fast or proven, small or long-context. Winning all four together is the only 'best' the numbers support.

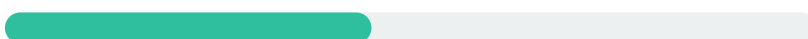
04 The honest bound



RESEARCH – MEASURED, PARTIAL

A gate is only as strong as the coverage of what it checks. The QUANTA chapter frontier (~271 of 600 gate-green) is today's ceiling: the gate is airtight over what it covers, and what it covers is still growing. We would rather own that bound than imply the gate already reaches everywhere.

Gate coverage (chapters)



~271 / 600

The bound, owned: a gate is only as strong as what it checks. It is airtight over its coverage — and coverage (~271 of 600 chapters) is still growing.

05 Why it ladders back



Trust, certified — at the language, at the mind, and at the way we work — is what lets AI be *good* in the strong sense: not merely useful, but safe to rely on. The proof is not paperwork bolted to the product. The proof is the product.

What is still open — kept visible

The honest edges of this lane, shown next to the wins. None of these is a footnote; each is a claim we have not yet earned — and marking it is what earns the right to be believed on the rest.

- ☐ The 9/10 is a self-score against our own verifiability rubric, not an external audit — we label it as such.
- ☐ The gate's strength is bounded by the coverage of what it checks; the QUANTA chapter frontier (~271/600) is the current ceiling.

06 Sources, glossary & the ledger

Where every claim comes from — and what the plain words mean.

Sources

- 01 Vocabotics Dashboard — Why best-in-world: trust certified at three layers, 2026-07-02 — Vocabotics internal record (as of July 2026)

THE HONESTY LEDGER

- Shipped — measured, gate-clean.
- 🔬 Research — measured, partial; the gap is shown.
- 🔭 Vision — designed, honest; not yet built.

The bigger claim goes up a tier, never down. We show the eval beside the demo — the tiering is the credibility.

See where the proof begins

The gate verifies what QUANTA builds. Start at the language that proves itself, byte-for-byte.

[/research/quant](#)



Nothing faked.

Measured, not asserted. Every claim carries its own honesty tier and the open edges are shown next to the wins — the tiering is the credibility.

