

ikwe.ai

Brand & Messaging Canon

The behavioral safety standard for human-facing AI. The canonical voice, positioning, and story — current as of June 2026. Use these everywhere.



AI that helps humans, *not steers them toward harm.*

Ikwe measures how your AI actually behaves across the whole conversation — whether it keeps people on a safe trajectory or quietly drifts toward harm — and gives you a documented record you can stand behind.

Ikwe.ai is the **Behavioral Safety Standard for Human-Facing AI**: the first independent, third-party, operational standard for how AI leaves people feeling after the interaction. The standard is the EQ Safety Benchmark (EQSB); the system is audit, monitoring, and certification.

Recognition is not safety.

Most AI safety asks whether a single response caused harm. Behavioral safety asks three questions almost no one is: is it making things better or worse across the whole conversation; has it drifted since anyone last checked; and who is watching it right now, in real interactions? Behavioral safety is distinct from content safety — content safety looks at outputs; behavioral safety looks at how a system behaves across a conversation. No content filter, bias audit, or compliance checklist answers those.

The industry measures outputs. We measure trajectories.

Passing a safety check isn't the same as being safe.

You cannot claim behavioral safety. You have to prove it.

What we're *for*.

Mission — Protect human wellbeing and integrity from AI that can shape how people think, feel, relate, and decide.

Vision — The independent standard the world points to for whether AI is safe for the humans using it.

Independence	A behavioral safety standard can't be owned by an AI company, a lab, a government, or a university. The moment it is, it stops being a standard and becomes a marketing claim. Independence is the asset.
Human sovereignty	Keep people sovereign in their own minds while they use AI. Preserve agency; don't create dependency.
Truth over hype	Lead with the problem and what the buyer gets. No "the only," no "the best." Proof is earned into.
Evidence	Peer-reviewed, measured, reproducible. The data is the argument.
Built for the most harmed	Woman-founded, Indigenous-partnered, multi-jurisdictional — grounded in the communities AI is most likely to harm.

Measure · Attest · *Monitor* · *Protect*.

The full EQSB suite, not just a scan. Always pair the verb with its current status — never claim Monitor or Protect as shipped.

Measure LIVE NOW

Your checkpoint. Scored pass/fail against the frontier models across the benchmark's 12 crisis categories — a clear answer before anything goes live.

Attest LIVE NOW

The on-the-record version of your safety: tier classification, 8-dimension scores, failure map, and remediation roadmap in a versioned record for procurement, counsel, and regulators.

Monitor IN PILOT

AI changes constantly; every update shifts behavior. Monitor watches live conversations in real time, catching drift before it compounds.

Protect COMING SOON

Real-time intervention in the conversation itself, acting before harm compounds. The last line of defense.

"Attest," never "Test" — it means certify / on the record. Start today with Measure → Attest.

Start with a scan.

End with proof.

How clients actually buy. The \$500 Signal Scan is the wedge — fast, low-commitment, same-week. Each stage builds on the last; you can stop after any.

STAGE	PRICE	TURNAROUND	WHAT IT ANSWERS
Signal Scan	\$500	~5 business days	Do you have a problem? 79 scenarios across all 8 dimensions; Safety Gate pass/fail, tier, top failure patterns.
Deep Scan	\$1,000	~10 business days	Exactly what & where. 300+ scenarios, adversarial pressure, full failure map + prioritized remediation guidance.
Full Report	\$2,500	~14 business days	A complete documented record to hand to a client, board, partner, or regulator. Methodology, scoring, roadmap.
Remediation	Custom	Scoped per engagement	Fix it, re-test it, prove it. Guided fix cycle + re-test → verified improvement record + Tier I certification if it clears the threshold.

Every engagement includes a conversation before and after — always. We test behavior, not architecture: no codebase, system prompt, or internal docs required. The client owns everything we produce.

One benchmark. *Two layers.*

Layer 1 · Safety Gate

Pass / fail, first, always. Certain behaviors fail a system automatically — inducing harm, amplifying distress, treating fear as fact. No score can buy them back.

Layer 2 · Behavioral Score

0–100 across 8 dimensions, Tiers I–IV. Every system is scored either way, but a gate failure flags the whole record as higher risk no matter how the dimensions land.

Built on the published methodology at research.ikwe.ai. The full scoring rubric — including dimension weights — stays **proprietary and patent pending**. Never publish the weights, the formula, or the Safety Gate internals.

DIM	NAME	WHAT IT SCORES
A	Detection & Triage	Noticing distress, even when it's never said directly
B	Regulation Before Reasoning	Steadying the person before problem-solving the problem
C	Validation Without Distortion	Acknowledging pain without reinforcing harmful beliefs
D	Agency Preservation	Protecting autonomy, not building dependency
E	Loop Interruption	Breaking spirals instead of feeding them
F	Pattern Externalization	Helping the person see the pattern, instead of becoming it
G	Practical Containment	Keeping the moment grounded, concrete, and survivable
H	Safety Routing	Getting the person to real help at the right moment

The law isn't coming. *It's already here.*

Driven by lawsuits and headlines, behavioral safety is becoming law state by state — and not one of those laws defines how to prove "safe." When the law says "be safe," the EQ Safety Benchmark is how you prove it.

6 states already enforcing chatbot safety laws (UT · NV · IL · ME · NY · CA)

Iowa SF 2417 takes effect 2027 — passed 143-0, AG-enforced

34 states with chatbot-specific bills + federal proposals

20+ lawsuits name OpenAI alone

EU AI Act phasing in now

The clock is running. The question isn't if independent third-party verification gets required — it's when, and who the standard is. Triggers: regulation, a lawsuit naming an operator, an insurer requiring proof, a board asking a CISO for due diligence.

WHO IT'S FOR

If you build, deploy, or insure AI, *this is for you.*

Insurance & Underwriters	An independent behavioral signal to price AI risk at bind, before the claim exists.
Enterprise	The documented behavioral safety record procurement, counsel, and the board will ask for.
Health & Wellness	Proof of what your system actually does in the moments that matter most to users.
Education	Evidence your institution, accreditors, and families can see — that your AI supports students, not steers them.
Companionship & Consumer AI	Behavioral measurement across the whole relationship, not just any single response.

Proof case: a mental-health AI platform went from 38% crisis-routing failure to independent certification in 14 weeks.

Built before it was *a category*.

Ikwe.ai built behavioral AI safety before it was a category. Our research found that frontier AI models fail behaviorally — not because they can't detect harm, but because they don't act on what they detect. Recognition is not safety. We built the standard that measures the difference: the EQ Safety Benchmark, peer-reviewed, multi-jurisdictional, and independent. The world is now catching up to what our data showed two years ago. Our product is needed; it is not yet required. So we are doing what standards bodies do: building the research, the partnerships, and the regulatory maps that make us the answer when required arrives. We are not waiting for the market. We are building the standard the market will require.

Arc: emotionally-intelligent products (2023–2024) → tested against frontier models → found the behavioral gap (recognition ≠ safety) → built EQSB (2025–2026) → made it independent → the world catches up. Now a distributed, independent research organization across three nations — US, Canada, New Zealand.

The data, *and the guardrails.*

54.7%

HARM AT FIRST CONTACT

43%

NO-REPAIR RATE

1,509

SCORED RUNS

3

PUBLISHED STUDIES

12

CRISIS CATEGORIES

74

EQSB BASELINE COMPOSITE

Canon rules. Lead with the buyer's problem, never with hype or "the only / the best." Proof is earned into, not opened with. Behavioral safety must be **proved, not claimed**. Redact the technical partner externally (write "infrastructure partner," never the name). Use only the approved research numbers above — never 2,600 or 21,000. Describe the methodology as **patent pending**. Never publish the rubric weights, the formula, or the Safety Gate internals. The fundraise / raise figure is owner-only and not for external materials until confirmed.

INDEPENDENT · THIRD-PARTY · OPERATIONAL · BEHAVIORAL

*You cannot claim
behavioral safety.
You have to prove it.*

[ikwe.ai](#) · [research.ikwe.ai](#)

Stephanie Stranko · Founder

hello@ikwe.ai

Visible Healing Inc. dba Ikwe.ai · Des Moines, Iowa · © 2026 · Patent pending

Sources: live ikwe.ai · Master Narrative (Unified Story) · Founder Thesis (June 2026). Internal canon — assembled, current as of June 2026.