



APEXFORM LIFE • RESEARCH WHITE PAPER • VERSION 1.4

Beyond the Snapshot

A Longitudinal, AI-Assisted Framework for Observing
How Human Paradigms Actually Change

Jake Wyatt — Founder, ApexForm Life Pty Ltd
Western Australia • apexformlife.com

For researchers, practitioners & partners

Executive Summary

Most tools for human development capture a **snapshot**: a personality result, a journal entry, a coaching session, an annual review. Each is useful, yet none shows how a person’s inner patterns actually change across weeks, pressures, relationships, and decisions. ApexForm Life is building a framework for **longitudinal, multimodal self-observation** — supported by AI — that helps people see the recurring patterns shaping their language, choices, emotions, and identity, and work with those patterns deliberately over time. This is a research-and-development pathway, not a finished scientific claim. Its component assumptions are grounded in established evidence on implementation intentions, goal monitoring, habit formation, and reflective language; its central proposition — that AI-assisted longitudinal reflection produces useful, reliable developmental insight — is stated here as a testable hypothesis with a concrete validation plan. The invitation is to researchers, practitioners, and partners who can help test it properly.

“People don’t only need motivation or information. They need clearer visibility of the internal models that repeatedly generate their outcomes.”

1. The Problem: Self-Development Lives in Snapshots

A personality profile, a journal entry, a coaching conversation, a course module, a goal-setting exercise, or an annual review can each be valuable. Their shared limitation is that they are **episodic**. They reveal what a person thinks in a moment, but not how their thinking evolves across time, environments, pressures, relationships, and decisions. This produces a recurring set of failures:

- People mistake temporary emotional states for stable truth.
- Recurring beliefs stay invisible because, from the inside, they feel like simple reality.
- Behavioural patterns are noticed only after their consequences have accumulated.
- Coaches and mentors work from limited snapshots of a person’s lived context.
- Users struggle to connect a moment of insight with repeated action.
- Progress is hard to evaluate because the deeper pattern is rarely tracked at all.

The result is a gap between aspiration and adaptation. People may know what they want, understand what they should do, and still remain governed by older internal models. The narrow problem worth owning is therefore not ‘motivation’ or ‘information’, but **visibility over time**.

Snapshot tools vs. a longitudinal instrument

Dimension	Snapshot tools (typical)	Longitudinal observation (proposed)
Unit of insight	A single moment or result	A trajectory across many moments
Time horizon	One sitting; occasionally repeated	Continuous, multimodal capture
Pattern visibility	Inferred by the person or coach	Surfaced as confirmable hypotheses
Failure mode	Insight without follow-through	Feedback loop: notice → act → review

Dimension	Snapshot tools (typical)	Longitudinal observation (proposed)
Evidence of change	Rarely measurable	Trackable against validated measures

Figure 1. The wedge: episodic snapshots versus longitudinal, multimodal self-observation.

2. Origin and Intellectual Context

The work began as a personal reconstruction. Bob Proctor’s teaching on paradigms supplied an early language for why people can sincerely want different outcomes while continuing to produce familiar results. That became a practical question: if paradigms shape behaviour, can their effects be observed consistently enough to help a person change them deliberately?

The question expanded through adjacent influences — growth mindset, behavioural design, habit formation, cognitive science, systems thinking, leadership development, and emotion regulation. ApexForm Life is not a rejection of coaching, psychology, or education; it is an attempt to integrate their strongest practical insights into a continuously learning support system. It sits deliberately between domains: part coaching framework, part reflective technology, part behavioural-observation system, part research programme.

On the Proctor lineage — stated plainly

Proctor’s ‘paradigm’ is a motivational-philosophy construct, not a validated scientific one. ApexForm Life treats it as the **motivating intuition** only. The evidence base, the operational definitions, and the validation plan in this paper are what make the idea testable. We aim to carry the emphasis on personal agency and disciplined transformation into a modern, evidence-oriented context — without borrowing scientific authority the original source does not claim.

3. Literature Grounding (and Its Limits)

The ApexForm Framework is not yet a validated model. Its current strength is that its component assumptions align with established and emerging evidence. The table below records each pillar and the honest strength of its support — including where the evidence is contested.

Pillar / assumption	Representative evidence	Strength for our use
Specific plans aid follow-through	Implementation intentions; meta-analysis (Gollwitzer, 1999; Gollwitzer & Sheeran, 2006)	Strong
Monitoring progress aids goals	Meta-analysis of goal monitoring (Harkin et al., 2016)	Strong
Habits form via context + repetition	Habit formation in the field (Lally et al., 2010; Wood & Neal, 2007)	Strong
Language reveals psychological signal	Expressive writing; LIWC text analysis (Pennebaker, 1997; Tausczik & Pennebaker, 2010)	Moderate
Beliefs about change shape effort	Implicit theories / mindset (Dweck & Leggett, 1988; Blackwell et al., 2007)	Contested*
Longitudinal traces reveal behaviour	Digital phenotyping (Torous et al., 2016; Onnela & Rauch, 2016)	Promising; clinical origin†
AI can support reflection at scale	Randomized field trials in classrooms (Kumar et al., 2024)	Emerging, peer-reviewed

Figure 2. Evidence strength by pillar. * and † are qualified below.

* The mindset evidence is genuinely contested

Recent meta-analyses of the same literature disagree: some find near-zero average effects on achievement and signs of publication bias (Macnamara & Burgoyne, 2023), while others find small but real effects concentrated among at-risk groups in supportive contexts (Burnette et al., 2023). A methodological commentary argues the divergence is the point — effects are heterogeneous, not universal (Tipton et al., 2023) — and a pre-registered national experiment found the intervention worked in specific conditions (Yeager et al., 2019). ApexForm therefore treats mindset and identity language as **hypotheses to test per person and over time**, not as a proven lever to assume.

† Digital phenotyping comes from clinical contexts

References for digital phenotyping are drawn from psychiatry, where it runs under clinical governance, ethics oversight, and clinician interpretation. Applying it to a direct-to-consumer self-development product is a different risk—evidence context, and the literature is explicit about unresolved problems of validity, the ‘ground truth’ criterion, consent, and sensitive-data governance (Martínez-Martin et al., 2018; Spathis et al., 2021; Stachl et al., 2021). ApexForm will treat phenotyping as **hypothesis-generating, never diagnostic**.

Taken together, these literatures do not prove that AI can identify a person’s paradigms. They support a narrower, defensible proposition: longitudinal reflection, behavioural tracking, and language-pattern analysis may help users and practitioners form **better hypotheses** about the internal models behind repeated outcomes.

4. Core Hypothesis and the Construct We Must Measure

ApexForm Life rests on the hypothesis that human development produces **observable signals** — across language, behaviour, routines, decisions, emotional responses, and identity statements — that, interpreted over time, may reveal stable patterns corresponding to a person’s underlying paradigms. The hypothesis has three parts:

- Paradigms are not directly observable, but their outputs often are.
- AI can help by organising large volumes of reflective data into patterns a person may miss.
- Human review must stay central; meaning, context, ethics, and agency cannot be safely automated.

A serious reader will immediately ask: identify patterns accurately relative to what? Without an operational definition and a ground-truth criterion, ‘accuracy’ cannot be evaluated. This is the construct-validity problem, and answering it is the most important step from vision to research.

Operational stance (how a ‘pattern’ earns the status of evidence)

Because a paradigm cannot be measured directly, ApexForm treats a surfaced pattern as evidentially ‘real’ only to the degree that it satisfies three independent tests: the user **confirms** it as recognisable; it is **stable** across repeated observation over time; and it **converges** with an established, validated measure of the relevant construct (for example, a recognised self-report scale or a computational language marker). A pattern that meets one test is a prompt for reflection; a pattern that meets all three is a candidate finding.

“The system should never tell a person who they are. It should help them examine the patterns that may be shaping who they are becoming.”

5. The ApexForm Framework: Five Layers

The framework converts raw lived experience into progressively more useful understanding, keeping human interpretation at the centre at every step. Transformation is rarely produced by insight alone — a person can

identify a limiting pattern and still repeat it — so the architecture is a feedback loop, not a funnel.

Layer	Function	What it produces
1 • Capture	Voice-first, low-friction recording of reflections in the moment	Raw multimodal record
2 • Structuring	Organise entries into themes, language, emotion, commitments	Searchable developmental log
3 • Modelling	Surface recurring patterns as confirmable hypotheses (the construct-validity layer)	Candidate patterns + uncertainty
4 • Reflection	Guided review; user confirms, rejects, or refines	Validated personal insight
5 • Adaptation	Translate insight into specific plans and re-observe (reconnection loop)	Behaviour change, tracked

Figure 3. The five-layer loop. Layer 3 is where construct validity lives; Layer 5 closes the loop back to Layer 1.

The platform that expresses this framework is **Paradigm** (the application). Its distinctive design principle is **reconnection over discovery**: the goal is not a one-off revelation but helping a person return, repeatedly and deliberately, to the way of being they recognise as their own. The strongest use case is not a single AI response; it is the accumulation of evidence across time.

6. From Vision to Research: A Concrete Validation Plan

Each priority research question is converted below into a testable form with a named measure and a comparison condition. This is the single biggest credibility lift available to the project, and the thing a research partner or investor most wants to see.

Research question	Operationalised as	Measure / comparison
Does AI reflection surface patterns users recognise?	User-rated accuracy of AI pattern summaries	Likert accuracy + recognise/reject rate vs. a placebo summary control
Does longitudinal review beat episodic journaling?	Change in a self-awareness measure over 30 days	RCT: Paradigm vs. unstructured journaling vs. waitlist; pre/post validated scale
Can language change predict behaviour change?	Linguistic markers → later self-reported action	Within-person, time-lagged correlation; report effect size + CI
Which patterns are signal vs. noise?	Test—retest stability of surfaced patterns	Stability coefficient; blind human-rater false-positive review
What feedback builds agency without dependence?	Autonomy + reliance measures across feedback styles	A/B of feedback framings; autonomy scale + usage-dependence proxy

Figure 4. Research questions rendered as a falsifiable design.

Evidence strategy in three layers

Layer one — founder corpus. Years of the founder’s reflections, notes, and voice captures form a development corpus for building the initial taxonomy and testing pattern extraction. Bias is disclosed openly: n = 1, self-selected, self-interpreted — useful for development, not proof. **Layer two — controlled pilots.** A 30-day pilot tests whether users find observations accurate, useful, actionable, and respectful of autonomy. **Layer three — expert collaboration.** Qualified researchers strengthen methodology, ethics, and evaluation. A mature programme adds user-rated accuracy, independent human review, false-positive analysis, longitudinal follow-up, privacy testing, and comparison against simpler alternatives.

Kill criteria (what would make us stop or pivot)

If surfaced patterns are not recognised by users above a placebo baseline, or are unstable on test—retest, or do not outperform unstructured journaling on a validated self-awareness measure, the core hypothesis is not supported and the approach should be redesigned rather than scaled.

7. Ethics, Safety, and Privacy

A system that observes thought, emotion, behaviour, and identity carries obvious responsibility and should earn trust through restraint as much as capability. Seven principles set the floor: human dignity before optimisation; user ownership and control of data; clear consent for what is captured, analysed, stored, and shared; transparency about uncertainty; AI as a reflective assistant, not an authority over identity; no manipulation or covert profiling; and clear escalation where professional support is more appropriate.

LLM-specific safeguards (engineering against known failure modes)

- **Sycophancy.** LLMs tend to agree with and flatter users — one evaluation observed it in roughly 58% of cases (Fanous et al., 2025). A tool that mirrors a user’s framing back would reinforce the very paradigm it claims to examine. Paradigm is prompted to surface disconfirming patterns, not just agreeable ones.
- **Automation bias / over-trust.** People over-rely on confident machine outputs (Lawrence et al., 2024). Every surfaced pattern is presented as a hypothesis with visible uncertainty and requires user confirmation.
- **Hallucinated psychological inference.** The system does not assert latent traits it cannot substantiate (Omran et al., 2024).
- **Crisis escalation.** Conversational agents are slow to escalate risk (Heston, 2023); Paradigm maintains a tested pathway to human and professional support rather than continuing ordinary dialogue.

Privacy commitments (regulator-ready)

As an Australian company processing what is almost certainly ‘sensitive information’ about mental state, ApexForm Life operates under the Privacy Act 1988 and the Australian Privacy Principles, and under the GDPR (notably the Article 9 special-category regime) for EU users. We commit to: explicit opt-in consent for sensitive-data processing; data minimisation with on-device or user-controlled storage where feasible; a clear retention and deletion policy; no secondary use or sale of personal data; and an external ethics review before any external pilot.

8. Positioning: Not Another Snapshot

ApexForm is not competing to be another content feed, a diagnostic instrument, or a chatbot therapist. Personality tests, mood check-ins, journaling and CBT apps, courses, and AI chat companions each capture a moment; their value is real but episodic. ApexForm’s distinct claim is **longitudinal and multimodal**: it observes how a person’s language, choices, emotional responses, and self-description evolve across time, and feeds that observation into a reconnection loop designed to turn insight into a default way of being. The outreach posture is founder-to-founder and researcher-to-researcher: a serious attempt to extend a transformative body of ideas into an evidence-oriented framework.

9. What This Is Not

- It is not a claim that AI can fully understand a human being.
- It is not a clinical diagnostic tool.
- It is not a replacement for therapy, coaching, mentoring, education, or spiritual practice.
- It is not a promise that all beliefs or patterns can be changed through software.

- It is not a finished scientific theory.

It is a serious development pathway built from lived experience, disciplined reflection, product experimentation, and a hypothesis that deserves structured testing.

10. Research and Development Roadmap

Stage	Focus	Primary output	Horizon
1 • Corpus	Build taxonomy from founder corpus; test extraction	Initial pattern taxonomy	0—3 mo
2 • Prototype	Capture → structuring → modelling in Paradigm	Working reflection loop	3—6 mo
3 • Pilot	30-day controlled user pilot vs. comparison	Accuracy & usefulness data	6—9 mo
4 • Collaboration	Researcher-strengthened methodology & ethics	Validated study design	9—15 mo
5 • Responsible scale	Expand only if kill criteria are cleared	Evidence-backed product	15 mo+

Figure 5. Indicative horizons; each stage gates the next.

11. Collaboration Invitation

ApexForm Life is seeking conversations with people whose work intersects human development, paradigm change, coaching, education, leadership, psychology, AI, behavioural science, and ethics. The most valuable collaborators are not those with the largest audiences but those who can improve the integrity of the work: reviewing the framework, advising on research design, stress-testing the ethical model, contributing domain expertise, or supporting pilots. For networks connected to Bob Proctor’s legacy, the invitation is specific — to explore whether the idea of paradigms can be carried into a modern technological context without diluting its emphasis on personal agency, responsibility, and disciplined transformation.

“If the premise is worth exploring, let us explore it properly — observed with more clarity, modelled with more care, and changed with more responsibility.”

Start a conversation: jakewyatt@apexformlife.com • apexformlife.com

— Jake Wyatt • Founder, ApexForm Life Pty Ltd • Western Australia

12. Limitations and Open Questions

This paper states a hypothesis and a plan, not results. The construct of the ‘paradigm’ remains to be validated against established instruments; the platform’s accuracy, stability, and safety are unproven; and the strongest supporting evidence (longitudinal AI-assisted reflection) is emerging rather than settled. The founder corpus is a development resource, not proof. Privacy commitments here should be confirmed with an Australian privacy practitioner before any external pilot. These are the right uncertainties to hold openly — and the reason this paper is an invitation to test, not a claim to have arrived.

References

Foundational and supporting evidence (APA 7). All sources verified; DOIs are clickable.

- Blackwell, L. S., Trzesniewski, K. H., & Dweck, C. S. (2007). Implicit theories of intelligence predict achievement across an adolescent transition. *Child Development*, 78(1), 246–263. <https://doi.org/10.1111/j.1467-8624.2007.00995.x>
- Burnette, J. L., Billingsley, J., Banks, G. C., Knouse, L. E., Hoyt, C. L., Pollack, J. M., & Simon, S. (2023). A systematic review and meta-analysis of growth mindset interventions. *Psychological Bulletin*, 149(3–4), 174–205. <https://doi.org/10.1037/bul0000368>
- Dweck, C. S., & Leggett, E. L. (1988). A social-cognitive approach to motivation and personality. *Psychological Review*, 95(2), 256–273. <https://doi.org/10.1037/0033-295X.95.2.256>
- Fanous, A., Goldberg, J., Agarwal, A. A., Lin, J., Zhou, A., Daneshjou, R., & Koyejo, S. (2025). SycEval: Evaluating LLM sycophancy [Preprint]. arXiv. <https://doi.org/10.48550/arXiv.2502.08177>
- Gollwitzer, P. M. (1999). Implementation intentions: Strong effects of simple plans. *American Psychologist*, 54(7), 493–503. <https://doi.org/10.1037/0003-066X.54.7.493>
- Gollwitzer, P. M., & Sheeran, P. (2006). Implementation intentions and goal achievement: A meta-analysis. *Advances in Experimental Social Psychology*, 38, 69–119. [https://doi.org/10.1016/S0065-2601\(06\)38002-1](https://doi.org/10.1016/S0065-2601(06)38002-1)
- Harkin, B., Webb, T. L., Chang, B. P. I., Prestwich, A., Conner, M., Kellar, I., Benn, Y., & Sheeran, P. (2016). Does monitoring goal progress promote goal attainment? A meta-analysis. *Psychological Bulletin*, 142(2), 198–229. <https://doi.org/10.1037/bul0000025>
- Heston, T. F. (2023). Safety of large language models in addressing depression. *Cureus*, 15(12), e50729. <https://doi.org/10.7759/cureus.50729>
- Kumar, H., Xiao, R., Lawson, B. A., Musabirov, I., Shi, J., Wang, X., Luo, H., Williams, J. J., Rafferty, A. N., Stamper, J., & Liut, M. (2024). Supporting self-reflection at scale with large language models: Insights from randomized field experiments in classrooms. In *Proceedings of the Eleventh ACM Conference on Learning @ Scale*. ACM. <https://doi.org/10.1145/3657604.3662042>
- Lally, P., van Jaarsveld, C. H. M., Potts, H. W. W., & Wardle, J. (2010). How are habits formed: Modelling habit formation in the real world. *European Journal of Social Psychology*, 40(6), 998–1009. <https://doi.org/10.1002/ejsp.674>
- Lawrence, H. R., Schneider, R. A., Rubin, S. B., Mataric, M. J., McDuff, D. J., & Jones Bell, M. (2024). The opportunities and risks of large language models in mental health. *JMIR Mental Health*, 11, e59479. <https://doi.org/10.2196/59479>
- Macnamara, B. N., & Burgoyne, A. P. (2023). Do growth mindset interventions impact students' academic achievement? A systematic review and meta-analysis. *Psychological Bulletin*, 149(3–4), 133–173. <https://doi.org/10.1037/bul0000352>
- Martinez-Martin, N., Insel, T. R., Dagum, P., Greely, H. T., & Cho, M. K. (2018). Data mining for health: Staking out the ethical territory of digital phenotyping. *npj Digital Medicine*, 1, 68. <https://doi.org/10.1038/s41746-018-0075-8>
- Michie, S., van Stralen, M. M., & West, R. (2011). The behaviour change wheel. *Implementation Science*, 6, 42. <https://doi.org/10.1186/1748-5908-6-42>
- Omrani, A., Karimi-Malekabadi, F., Trager, J., Atari, M., Park, P. S., Golazizian, P., Abdurahman, S., Xue, M. J., & Dehghani, M. (2024). Perils and opportunities in using large language models in psychological research. *PNAS Nexus*, 3(7), pgae245. <https://doi.org/10.1093/pnasnexus/pgae245>
- Onnela, J. P., & Rauch, S. L. (2016). Harnessing smartphone-based digital phenotyping to enhance behavioral and mental health. *Neuropsychopharmacology*, 41, 1691–1696. <https://doi.org/10.1038/npp.2016.7>
- Pennebaker, J. W. (1997). Writing about emotional experiences as a therapeutic process. *Psychological Science*, 8(3), 162–166. <https://doi.org/10.1111/j.1467-9280.1997.tb00403.x>
- Spathis, D., Perez-Pozuelo, I., Morley, J., Gifford-Moore, J., & Cows, J. (2021). Digital phenotyping and sensitive health data: Implications for data governance. *JAMIA*, 28(9), 2020–2024. <https://doi.org/10.1093/jamia/ocab012>
- Stachl, C., Boyd, R. L., Horstmann, K. T., Khambatta, P., Matz, S. C., & Harari, G. M. (2021). Computational personality assessment. *Personality Science*, 2, e6115. <https://doi.org/10.5964/ps.6115>
- Tausczik, Y. R., & Pennebaker, J. W. (2010). The psychological meaning of words: LIWC and computerized text analysis methods. *Journal of Language and Social Psychology*, 29(1), 24–54. <https://doi.org/10.1177/0261927X09351676>
- Tipton, E., Bryan, C. J., Murray, J. S., McDaniel, M. A., Schneider, B., & Yeager, D. S. (2023). Why meta-analyses of growth mindset and other interventions should follow best practices for examining heterogeneity. *Psychological Bulletin*, 149(3–4), 229–241. <https://doi.org/10.1037/bul0000384>
- Torous, J., Kiang, M. V., Lorme, J., & Onnela, J. P. (2016). New tools for new research in psychiatry. *JMIR Mental Health*, 3(2), e16. <https://doi.org/10.2196/mental.5165>

Wood, W., & Neal, D. T. (2007). A new look at habits and the habit-goal interface. *Psychological Review*, 114(4), 843—863.
<https://doi.org/10.1037/0033-295X.114.4.843>

Yeager, D. S., Hanselman, P., Walton, G. M., Murray, J. S., Crosnoe, R., Muller, C., Tipton, E., Schneider, B., ... Dweck, C. S. (2019). A national experiment reveals where a growth mindset improves achievement. *Nature*, 573(7774), 364—369.
<https://doi.org/10.1038/s41586-019-1466-y>