

Research Methodology for Testing the Algorithmic Alienation Index and Systemised Self Hypothesis v2

Galu & Kairos

Abstract

This research proposal builds on the previous proposal submitted 5-9-2026 by the authors and includes amendments based on reviewer feedback for this version 2 re-submission (Galu & Kairos, 2026a). This preregistered, New Zealand-only mixed-method programme will test a proposed Algorithmic Alienation Index (AAI) and the Systemised Self (SS) hypothesis without treating either as established fact or diagnosis. Algorithmic Alienation (AA) is the conceptual starting point, grounded theoretically by Kanbay et al. (2026) and partially in qualitative evidence by Arkan et al. (2026). The AAI is a new, theory-derived operational proposal; SS is a further, falsifiable hypothesis concerning phenomenological inversion, where structurally dependent algorithmic use is experienced as self-authored freedom. A 50-person pilot (1 February–30 April 2027; NZ\$25,000) will test item comprehension, recruitment, retention, controlled-feed realism, and trace coverage. Its prospective results, not yet available, will inform transparent Monte Carlo finalisation in the main study, provisionally mobilising 3 May 2027. The main programme targets about 700 unique adults, with independently recruited EFA and CFA samples (≈ 350 each), a nested baseline/3-/6-month cohort (≈ 240), controlled-feed experiment (≈ 200), field intervention (≈ 240), IPA interviews (≈ 24), and blinded LEAD dossiers (≈ 30). The primary causal evidence is direct randomisation within a purpose-built research feed; field encouragement is supplementary and has pre-specified alternatives if compliance is weak. Base alienation, resistance (RA), and adjusted scores will always be reported separately. The proposed boundary, $AAI_adjusted \geq .95$ and $RA \leq .05$, is a non-diagnostic review flag only. SS decisions will use a fallible, blinded LEAD–IPA reference process that excludes AAI scores. The design provides numerical decision rules, incremental-validity tests, feasible recruitment safeguards, New Zealand privacy controls, and explicit routes for theory revision.

Keywords: *algorithmic alienation; scale development; recommender systems; digital trace data; phenomenological inversion; causal inference; New Zealand; systemised self*

Contents

- 1. Purpose, epistemic status, and questions**
- 2. Conceptualisation and scoring**
- 3. Design, setting, and feasibility**
- 4. Measures and procedures**
- 5. Analysis, hypotheses, and numerical decisions**
- 6. Causal field component and fallback ladder**
- 7. Blinded LEAD–IPA integration and SS decisions**
- 8. Ethics, privacy, resources, and timeline**
- 9. Interpretation, dissemination, and limitations**

References

Appendix A - Compact AAI item blueprint

1. Purpose, epistemic status, and questions

This study distinguishes three epistemic levels:

Level	Status at commencement	What this study may conclude
Algorithmic Alienation (AA)	Theoretical construct with partial qualitative grounding (Arkan et al., 2026; Kanbay et al., 2026)	Whether its proposed dimensions receive further measurement evidence in this NZ sample
Algorithmic Alienation Index (AAI)	New operational instrument derived from AA (Galu & Kairos, 2026c)	Whether its items, structure, weighting, and predictive utility are defensible
Systemised Self (SS)	Hypothesised terminal configuration (Galu & Kairos, 2026b)	Whether convergent evidence supports, revises, or fails to support the hypothesis; never a diagnosis

AA concerns diminished self-direction, ambiguous preference authorship, automatic decision-making, emotional disconnection, and possible resistance in sustained algorithmic mediation. It does not entail that recommendation use, disability accommodation, rapid choice, or delegation is alienating. The AAI is intended to measure a proposed AA-related profile, rather than to redefine AA. SS is narrower: the hypothesised coexistence of marked structural dependence, ineffective resistance, and phenomenological inversion, curation experienced as authentic liberation rather than constraint. Since this account could make self-report deceptively reassuring, SS cannot be inferred from a questionnaire, usage intensity, or a score.

The study asks: (1) Is an AAI with distinguishable DAu, IAm, EDm, and EDc dimensions psychometrically defensible? (2) Do the theory-derived score and RA improve prediction of pre-specified behavioural criteria beyond exposure and adjacent measures? (3) Does direct, temporary alteration of a controlled feed change recommendation-led action and goal recovery? (4) Is there reproducible, non-circular evidence of phenomenological inversion after plausible alternatives have been considered? (5) Do within-person changes in dependence and experience show a prospective association over six months? These questions are deliberately narrower than claims about population prevalence or long-term societal consequences.

2. Conceptualisation and scoring

The 15 original theory-derived items will be retained for testing within an approximately 30-item pool (three original and three expansion items per DAu, IAm, EDm, EDc, and RA). Responses use five frequency anchors, transformed from 1–5 to 0, .25, .50, .75, and 1. Dimension scores are reverse-keyed retained-item means. The exact theory-derived scoring hypotheses are:

$$AAI_base = .30(DAu) + .25(IAm) + .25(EDm) + .20(EDc).$$

$$AAI_adjusted = AAI_base \times (1 - RA).$$

The .30/.25/.25/.20 weights express theoretical priority, not observed effects. RA is hypothesised to moderate the whole profile. Every results table will report DAu, IAm, EDm, EDc, RA, AAI_base, and AAI_adjusted; an adjusted score will never be used to conceal a high base score. Equal-weight, CFA-score, training-set penalised-regression, and additive/interaction RA alternatives will be evaluated in held-out data.

AAI_{adjusted} $\geq .95$ and RA $\leq .05$ remains a candidate SS-review flag. It is non-diagnostic, is not an eligibility criterion, and cannot establish SS. A flag identifies cases for sampling or blinded review alongside matched non-flagged cases; it does not enter the LEAD dossier. Pilot data will not be portrayed as confirmation of any threshold.

3. Design, setting, and feasibility

The project is limited to adults aged 18-60 ordinarily resident in Aotearoa New Zealand, able to consent in English, and using a personalised feed or recommender at least weekly. English is used because the candidate items and interviews are English-language instruments; this is not a claim of cultural universality. Māori and Pacific participation will be actively supported, but the study will not make Māori data sovereignty claims that it is not resourced to fulfil. Consultation with institutional Māori research advisers will precede recruitment, and no ethnic mean comparison will be made without adequate measurement evidence and governance advice.

The design is intentionally smaller than Draft 1. One recruitment frame supplies approximately 700 unique main-study participants; components overlap, so their targets must not be summed. The controlled feed supplies primary behavioural and experimental evidence, avoiding dependence on commercial-platform access. Participant-authorised metadata and selected data donation are supplementary naturalistic evidence, consistent with the distinct strengths and limitations of tracking and donation (Ohme et al., 2024). Platform backend data are neither sought nor needed.

Component	Provisional analytic target	Relationship and principal purpose
Pilot	50	Separate; feasibility, item and protocol refinement
EFA	≈350	Main-study baseline, independently recruited; reduce candidate pool
CFA	≈350	Independently recruited confirmation sample; no EFA reuse
Longitudinal cohort	≈240	Nested within the independently recruited CFA cohort; baseline precedes model lock, then 3 and 6 months
Controlled-feed experiment	≈200	Nested volunteers, balanced factorial randomisation
Field intervention	≈240	Nested tracking-eligible volunteers; technical-support randomisation

Component	Provisional analytic target	Relationship and principal purpose
IPA	≈24	Purposive longitudinal subsample
LEAD dossiers	≈30	24 IPA cases plus 6 matched cases receiving a focused blinded phenomenological interview

The pilot begins 1 February 2027 and ends 30 April 2027, with provisional main mobilisation on 3 May 2027, conditional on institutional ethics approval by 5 March 2027, technical safety review, and pilot stop/go criteria. February is restricted to protocol finalisation, consultation, technical configuration, non-participant quality assurance, and recruitment preparation; consent, enrolment, cognitive interviews, and behavioural data collection cannot begin before approval. Failure to obtain approval by 5 March moves participant-facing work and the progression gate; it does not compress recruitment or review. The pilot will provide preliminary evidence prospectively; no pilot results exist now. Its descriptive item distributions, ordinal correlations, completion time, feed manipulation check, enrolment yield, trace-origin coverage, and retention intention will parameterise Monte Carlo simulations. Loading ranges of .40–.70 are provisional simulation scenarios rather than assumed results. The simulations will incorporate observed ordinal distributions, factor correlations, missingness, task-level clustering, crossover correlation, and attrition; report power and expected 95% confidence-interval width for each primary estimand; and determine final values around, not automatically above, the provisional targets. This is feasibility calibration, not post hoc power shopping.

Participant and evidence flow	Development stream	Confirmation and nested streams
Pilot n=50	Cognitive testing → usability/realism → simulation inputs	Separate from all main analyses
Main recruitment ≈700 unique adults	EFA ≈350 → item/model lock	Independently recruited CFA ≈350
Nested volunteers	—	Cohort ≈240; controlled feed ≈200; field intervention ≈240
Qualitative selection	—	IPA ≈24 + focused blinded interviews ≈6
Reference review	—	LEAD dossiers ≈30, sampled from qualitative cases and matched profiles

Prior anchoring is proposed to justify a cautious pilot: AA has theoretical and qualitative grounding, and the design retains established adjacent measures. The documented divergence between logged and self-reported digital media use warrants behavioural triangulation (Parry et al., 2021). Olson et al. (2022) further indicates that user-selected limits can be experimentally consequential whereas monitoring alone may not be, supporting tests of effective resistance rather than awareness alone. Neither body of evidence validates AAI or SS.

Recruitment, quotas, and retention

Recruitment will combine a New Zealand probability-informed commercial panel, university/community advertisements, libraries, workplaces, disability and community organisations, and paid social advertising. Panel invitations will be stratified before release by age (18-9, 30-44, 45-60), gender (including non-binary/self-described option), region (upper North Island, lower North Island, South Island), and self-reported personalised-feed use (weekly versus daily). Targets are feasibility quotas, not population weights: recruitment aims for at least 20% aged 45-60, at least 25% outside the upper North Island, and monitored representation of Māori, Pacific, Asian, Pākehā/European, and other participants where voluntarily self-identified. No group is excluded when a quota closes. Oversamples are analysed with recruitment-source indicators; population estimation is not claimed.

Contingency ladder	Trigger	Response, preserving comparability
1. Panel release	Cell <80% of weekly target	Release only that stratified panel cell; audit duplicates
2. Partner outreach	Cell <70% at midpoint	Paid partner/community placement; accessible plain-language materials
3. Schedule/access adjustment	Tracking or 45–60 cell <60%	Evening sessions, telephone onboarding, loaned data voucher where approved
4. Design protection	A quota cannot be achieved	Close with documented shortfall; do not fabricate balance; use source/region sensitivity analyses
5. Analytic restriction	Trace coverage inadequate	Retain survey/feed analysis; label naturalistic estimates selected-subsample only

The study advertises ‘online choice and recommendation systems’, not alienation. Screening prevents duplicate participation through panel IDs, contact verification, device/browser checks, and pre-specified attention checks, without biometric identification. Staged payments compensate baseline, feed session, follow-ups, and interviews separately; no payment depends on tracking. Brief mobile-capable surveys, participant-selected reminder channel, a maximum of three reminders per wave, and re-consent at follow-up reduce avoidable attrition. Recruitment flow, quota attainment, consent by modality, refusal, and loss will be reported.

4. Measures and procedures

At baseline, all participants complete AAI candidates, use/exposure measures, digital literacy, preference for delegation, demographics, accommodation needs relevant to task accessibility, and licensed adjacent measures: Bergen Facebook Addiction Scale (Andreassen et al., 2012), agency (Tapal et al., 2017), autonomous functioning (Weinstein et al., 2012), self-concept clarity (Campbell et al., 1996), and needs-related measures (Deci & Ryan, 2000). Measures of distress, where approved, are confounder/alternative-explanation screens, not labels. The comparison of logs to self-report is explicit, rather than assuming either is correct. The controlled feed is a responsive web application populated with ethically licensed neutral-interest New Zealand-relevant explainers, entertainment, consumer comparisons, and learning material. Formative pilot work will test perceived relevance, usability, content comprehension, and whether participants make choices across categories. Personalisation uses preferences stated *within the study*, never inferred sensitive traits. In a within-session, counterbalanced factorial design, participants encounter personalised versus relevance-stratified non-personalised ranking and standard versus friction conditions. Friction removes autoplay and personalised prompts and presents a neutral ‘search or browse categories’ option; content, search, and exit remain available. Realism is ecological task adequacy, not imitation of a commercial service. Provisional progression benchmarks are: at least 70% of pilot participants rate both task clarity and content relevance at least 4 on a 5-point scale; at least 70% encounter recommendation exposure in the intended directional contrast; and the aggregate personalised-versus-non-personalised exposure difference is at least 15 percentage points. These are transparent feasibility minima selected to prevent progression with a poorly perceived or weak manipulation, not validated universal thresholds. Failure requires redesign and re-pilot, not main causal claims.

Participants first state a concrete goal (e.g., locate two credible accounts and save one). Logged outcomes are recommendation-initiated eligible-action share, manual-search share, direct navigation, unique sources and entropy, active rejection, goal completion, time to completion, and post-friction recovery. The latter is goal-directed completion after curation withdrawal, not engagement time. Device, reading time, task order, and accessibility settings are recorded as necessary covariates. A source-monitoring/choice-blindness task assesses detection and confidence after ethically disclosed switches (Johansson et al., 2005); a goal-adaptability task assesses alternatives after a route change. Intentional binding is exploratory given its contested validity (Moore & Obhi, 2012).

Optional naturalistic evidence consists only of aggregate app/domain category, time rounded to a coarse interval, observable action origin (recommendation/search/direct/unknown), and search count, in two-week windows around each wave. No message content, typed search terms, screenshots, contacts, location, passwords, keystrokes, clipboard, or page-body text is collected. Participants may instead donate selected account-export aggregates. Origin-coverage denominators are reported; unknown origin is never coded as recommendation-led.

The longitudinal cohort repeats key measures and two-week trace windows at 3 and 6 months; the controlled-feed session recurs at month 6. IPA interviews follow baseline/feed participation and explore perceived authorship, constraint, liberation, counterfactual alternatives, reactions to friction, work/care obligations, expertise, and accommodation. Interviewers are blind to AAI values and use a neutral code and sampling stratum only. Analysis follows IPA's idiographic and reflexive commitments (Smith et al., 2022), with an audit trail and deviant-case attention.

5. Analysis, hypotheses, and numerical decisions

EFA uses polychoric correlations, robust ordinal estimation, parallel analysis, minimum average partial, scree, residuals, and conceptual interpretability. CFA is conducted only in the independent sample using ordinal robust weighted least squares, with robust maximum likelihood sensitivity analyses. Pre-specified models are: one factor; four correlated AA dimensions with RA separate; second-order AA; and cognitive–affective two-factor structure (DAu/EDm; IAm/EDc). No factor survives with fewer than three conceptually adequate items. Internal consistency is ordinal omega and coefficient H with confidence intervals; short-interval stability is reported separately from change.

Decision domain	Rule	Consequence
CFA adequacy	CFI/TLI $\geq .95$ desirable (.90 minimum contextual guide); RMSEA $\leq .06$; SRMR $\leq .08$; inspect residuals and theory (Hu & Bentler, 1999)	Cut-offs do not override poor interpretability
“Materially better” structure	Nested robust difference test where valid plus $\Delta CFI \geq .010$ or $\Delta RMSEA \geq .015$; non-nested lower BIC by ≥ 10 and better held-out log score	Retain simpler model if differences are smaller or factors correlate $\geq .85$
Reliability/discrimination	omega and H $\geq .70$; HTMT $< .85$ as preferred sensitivity guide	Revise items or describe dimensions as insufficiently distinct
Scoring equivalence	In held-out predictions, original score is practically equivalent if its RMSE is no $> 2.5\%$ worse <i>and</i> its R^2 is no $> .02$ lower than best alternative	If not equivalent, replace the formula for predictive use while retaining historical reporting
Incremental validity	Compare nested models; report ΔR^2 , likelihood-ratio/F test as appropriate, AIC/BIC, held-out RMSE/calibration	AAI requires $\Delta R^2 \geq .02$ and 95% CI excluding zero for at least one primary behavioural criterion beyond exposure, screen time, adjacent measure, and covariates

Primary convergent criteria are controlled-feed recommendation-led action and post-friction goal recovery; manual search is secondary. Nested linear/generalised models first contain pre-specified demographics, device/accessibility covariates, use exposure, screen time, problematic-use and needs measures; Step 2 adds AAI_base and RA; Step 3 adds their interaction. The exact delta- R^2 prevents a claim of incremental value based only on a significant coefficient. Out-of-sample calibration and prediction are reported. Associations are not causal.

The primary controlled-feed estimand is the intention-to-treat assignment effect. Randomisation is computer generated with concealed allocation, stratified by device type and baseline AAI_base tertile. A standardised mean difference of .35 in recovery and a 10-percentage-point difference in recommendation-led share are provisional smallest effects of operational interest, selected as changes large enough to justify interpretation in a brief research task; they are not claimed as effects established by prior AAI evidence. No existing study provides a defensible effect that can be transported to the novel AAI-behaviour relation. Olson et al. (2022) supports the plausibility of experimentally changing self-regulatory behaviour, but its effect is contextual rather than a power input for these different outcomes. Pilot-based simulations will therefore span AA-behaviour correlations of .10, .20, and .30 and incremental ΔR^2 values of .01, .02, and .05, explicitly including the preregistered .02 utility criterion. They must report the main-study sample needed for 80% and 90% power, expected 95% confidence-interval widths, crossover correlation, trial-level intraclass correlation, manipulation strength, missingness, and attrition for each scenario and both operational minima. The provisional $n \approx 200$ is retained only if those precision results are adequate. Moderation by baseline AAI/RA is exploratory unless the simulation demonstrates adequate interaction precision.

The cohort will use random-intercept cross-lagged panel models to distinguish stable between-person differences from within-person deviations (Hamaker et al., 2015). It tests temporal association, not causal sequence. Full-information estimation/multiple imputation is paired with attrition models and pattern-mixture sensitivity analyses. Configural then loading/threshold invariance is assessed across waves and, only if cell sizes permit, broad age, gender, region, and accommodation groups. Partial invariance or alignment is reported; unsupported group means are not compared.

6. Causal field component and fallback ladder

The field intervention randomises technical support, not participants' personal accounts: a supported session and optional tools to limit personalisation/autoplay, reset recommendations, and schedule intentional-search prompts versus general online-wellbeing information. The primary estimate is ITT on verified setting change, search and optional aggregate trace outcomes. CACE uses assignment as an instrument only after reporting ITT and only if the first-stage difference is at least 15 percentage points with an F statistic ≥ 10 , alongside monotonicity, exclusion, and independence arguments (Angrist et al., 1996). It will not rescue a weak intervention.

Causal fallback ladder	Condition	Confirmatory interpretation
1. Controlled-feed ITT	Randomisation/manipulation check succeeds	Primary causal evidence on immediate behaviour
2. Field ITT	Support changes uptake or not	Effect of offering supported choice, regardless of compliance
3. Enhanced encouragement	Weak uptake, before outcome inspection	Randomise a pre-consented hands-on onboarding/preference arm; evaluate ITT
4. CACE	Strong first stage and assumptions plausible	Explicitly assumption-dependent complier estimate
5. Observational target-trial sensitivity	No usable instrument	Association only; propensity/negative-control sensitivity, no causal conclusion

This ladder avoids interpreting voluntary non-compliance as evidence against resistance or personalisation effects.

7. Blinded LEAD–IPA integration and SS decisions

LEAD is a fallible reference standard, not a gold standard. Dossiers are purposively balanced across high-base/low-RA, high-base/high-RA, low/moderate profiles, high-dependence/low-estrangement discrepancies, and matched comparison cases, with representation monitored by recruitment strata. The 30 dossiers include 24 full IPA cases and six matched comparison cases receiving a shorter, score-blinded phenomenological interview covering authorship, freedom, estrangement, friction response, and alternative explanations. All review-flag cases are included up to capacity, with matched cases selected without AAI disclosure. Dossiers exclude every AAI item, dimension, base, RA, adjusted score, threshold flag, and recruitment description that could reveal it. They include de-identified qualitative extracts, controlled-feed and trace summaries, agency/authenticity measures, cognitive performance, longitudinal course, task-accessibility information, and an explicit alternative-explanations sheet. Full IPA and focused-interview cases are identified in sensitivity analyses so unequal qualitative depth cannot masquerade as diagnostic difference.

LEAD reliability and reconciliation	Pre-specified procedure
Panel	Three independent members with qualitative, clinical/psychological assessment, and digital-systems expertise; conflict declarations and calibration on non-study vignettes
Blind first pass	Each codes SS-compatible, severe AA without inversion, non-SS/adequate delegation, alternative explanation, or indeterminate, plus domain confidence
Agreement	Report percent agreement, Fleiss' kappa for category and ICC(2,k) for 0–100 confidence; “reasonable” prospective target is kappa $\geq .60$ and ICC $\geq .75$, with prevalence-adjusted interpretation and confidence intervals
Disagreement	Structured evidence conference after independent coding; retain initial codes; a fourth blinded adjudicator only for unresolved cases
Reporting	Consensus, all initial ratings, category shifts, missing evidence, dissent, and reasons are reported; low reliability downgrades SS inference

SS-compatible requires, after blinded review: (a) independently observed high recommendation dependence relative to adequate-coverage peers; (b) low effective resistance or poor accessible goal recovery; (c) IPA evidence that curated direction is consistently experienced as self-authored freedom/authenticity or relief from deliberation; (d) low experienced estrangement relative to this structural pattern; and (e) no more plausible explanation such as disability/assistive use, occupational necessity, caregiving scarcity, informed delegation, expertise, distress, or task inaccessibility. It is sufficient to classify a case as indeterminate or alternative explanation; absence of evidence is not evidence of SS.

Flow: candidate profile or purposive comparison → score-blinded IPA + behavioural dossier → three independent ratings → reliability check → reconciliation/adjudication → SS-compatible / severe AA / alternative explanation / indeterminate → joint display with quantitative uncertainty.

8. Ethics, privacy, resources, and timeline

New Zealand-only collection materially simplifies governance. The study follows the Privacy Act 2020 and Health Information Privacy Code where applicable, institutional ethics approval, data minimisation, purpose limitation, separate opt-in trace consent, withdrawal/deletion procedures, encrypted New Zealand-hosted institutional storage, role-based access, access logs, and a data-management plan. Identifiers are separated from research data; re-identification keys are restricted. Exported aggregates are checked automatically and manually for prohibited content. The team does not transfer identifiable data offshore or promise GDPR/CCPA compliance across multiple jurisdictions. Data retention, Māori data governance consultation, adverse-event response, and any secondary use require ethics-approved specification. Participants can complete surveys and feed tasks without naturalistic tracking and can skip any question.

Pilot/main timeline	Milestone and deliverable
1 Feb–28 Feb 2027	Protocol lock, ethics processing, Māori adviser consultation, content/licensing, panel set-up; no participant enrolment
1 Mar–31 Mar	Ethics approval required by 5 Mar; enrolment begins only after approval; pilot fielding, cognitive interviews and realism audit
1 Apr–30 Apr	Pilot retention check; documented item/protocol changes; Monte Carlo; stop/go decision

Pilot/main timeline	Milestone and deliverable
3 May–30 Jun	Provisional main mobilisation; independent EFA recruitment, analysis and item/model lock
Jul–Aug	Independent CFA recruitment; cohort baseline during CFA data collection; controlled-feed randomisation
Sep	CFA lock; confirmatory scoring; field intervention begins; qualitative selection
Oct–Nov	Cohort 3-month windows; field ITT; IPA and focused interviews
Jan–Feb 2028	Cohort 6-month windows, repeat feed, LEAD ratings
Mar 2028	Integrated analysis, transparency report and manuscripts

Embedded Gantt	Feb	Mar	Apr	May–Jun	Jul–Sep	Oct–Nov	Jan–Mar 2028
Ethics, consultation, technical QA	■	■ gate					
Pilot enrolment and fieldwork		■	■				
Pilot analysis and progression			■ gate				
EFA recruitment and lock				■			
CFA and cohort baseline					■		
Controlled feed and field intervention					■	■	
Three-month follow-up and interviews						■	
Six-month follow-up, LEAD and reporting							■

Pilot budget (NZ\$25,000)	NZ\$	Rationale
Participant reimbursements and accessibility support	8,000	Pilot n=50, interviews/usability and retention contacts
Feed configuration, licensed content, security testing	6,500	Reusable controlled environment, no commercial API dependence
Research assistant and recruitment administration	5,500	Scheduling, data-quality and consent support
Psychometrics/Monte Carlo/statistical consultation	2,500	Finalise feasible main targets
Ethics, Māori consultation, transcription/accessibility	1,500	Governance and accessible materials
Contingency	1,000	Documented only, not discretionary expansion
Total	25,000	

The indicative direct main-study ceiling is NZ\$150,000, subject to pilot-informed contracting and institutional approval. It is not represented as funded by the NZ\$25,000 pilot.

Indicative direct main-study budget	NZ\$	Scope
Participant reimbursement and retention	48,000	Baseline, feed, field, follow-ups, IPA and accessibility support
Coordination, research assistance and data management	35,000	Recruitment, consent, scheduling, quality and secure datasets
Controlled-feed engineering, hosting and security	22,000	Production hardening, licensed content, maintenance and assurance
Psychometrics and statistical analysis	16,000	EFA/CFA, Monte Carlo, longitudinal and causal models
IPA transcription, analysis and LEAD honoraria	13,000	Interviews, transcription, independent ratings and adjudication
Recruitment placements and community engagement	6,000	Panel top-ups, partner placements and accessible materials
Dissemination, archiving and participant summaries	3,000	Disclosure-reviewed outputs and open materials
Contingency	7,000	Ethics-approved response to documented delivery risk
Direct main-study total	150,000	Excludes separately valued institutional in-kind support

Institutional in-kind value is provisionally NZ\$31,000: principal-investigator supervision (NZ\$14,000), ethics/data-protection and Māori research-adviser time (NZ\$5,000), secure storage, survey/statistical licences and technical environment (NZ\$7,000), and meeting-room/accessibility services (NZ\$5,000). The combined indicative resource value is therefore NZ\$181,000 for the main study, plus the separate NZ\$25,000 pilot. In-kind items are auditable contributions, not cash available for participant recruitment. Core delivery requires a PI, .4 FTE project coordinator, .3 FTE research assistant/data manager, part-time feed developer/security support, psychometric adviser, IPA interviewer/analyst, and compensated LEAD members. This staffing, a six-month cohort rather than 12 months, and nested samples are the principal feasibility controls.

Key risk	Mitigation and decision
Recruitment imbalance	Stratified releases and contingency ladder; report shortfalls
Feed lacks realism	Pilot thresholds; redesign/re-pilot; no causal claim if unmet
Low trace uptake/coverage	Tracking optional; controlled-feed remains primary; restrict trace inference
Attrition	Staged payment/reminders; inverse-probability and pattern-mixture sensitivity
Weak field compliance	ITT and fallback ladder; no CACE overclaim
LEAD subjectivity	Blinding, quantified reliability, dissent and alternative categories
Stigma or distress	Neutral language, withdrawal, support information, non-diagnostic reporting

9. Interpretation, dissemination, and limitations

All hypotheses, exclusions, item changes, simulations, model syntax, primary outcomes, and analytic distinctions will be preregistered before confirmatory data inspection. Null, contradictory, and adverse results receive equal reporting. H1 is weakened by a reproducibly simpler factor model; H2 by practically superior alternatives; incremental distinctiveness by failure of the stated ΔR^2 rule; RA by null/reversed behavioural associations; and SS by high dependence that is adequately explained by accommodation/delegation or that consistently co-occurs with estrangement rather than inversion. Findings concern consenting New Zealand adults, a controlled research feed, and short-term follow-up. They cannot estimate national SS prevalence, diagnose participants, or establish that everyday commercial algorithms caused a profile.

De-identified aggregate code, data dictionary, scoring code, materials where copyright permits, and a transparency report will be shared under ethics and consent constraints. Results will be communicated in plain language to participants and recruitment partners. The study's contribution is disciplined testing: it may support a revised AAI, identify limits of the proposed score, or show that SS is not empirically distinguishable under these conditions.

References

- Andreassen, C. S., Torsheim, T., Brunborg, G. S., & Pallesen, S. (2012). Development of a Facebook addiction scale. *Psychological Reports, 110*(2), 501–517.
<https://doi.org/10.2466/02.09.18.PR0.110.2.501-517>
- Angrist, J. D., Imbens, G. W., & Rubin, D. B. (1996). Identification of causal effects using instrumental variables. *Journal of the American Statistical Association, 91*(434), 444–455.
<https://doi.org/10.1080/01621459.1996.10476902>
- Arkan, B., Kanbay, Y., & Akçam, A. (2026). From algorithms to the self: Understanding algorithmic alienation in the digital age. *BMC Psychology, 14*, Article 930. <https://doi.org/10.1186/s40359-026-04677-1>
- Campbell, J. D., Trapnell, P. D., Heine, S. J., Katz, I. M., Lavalley, L. F., & Lehman, D. R. (1996). Self-concept clarity: Measurement, personality correlates, and cultural boundaries. *Journal of Personality and Social Psychology, 70*(1), 141–156. <https://doi.org/10.1037/0022-3514.70.1.141>
- Deci, E. L., & Ryan, R. M. (2000). The “what” and “why” of goal pursuits: Human needs and the self-determination of behavior. *Psychological Inquiry, 11*(4), 227–268.
https://doi.org/10.1207/S15327965PLI1104_01
- Galu, J., & Kairos. (2026a). Research methodology for testing the Algorithmic Alienation Index and Systemised Self Hypothesis v1. aiXiv. <https://aixiv.science/abs/aixiv.260905.000002>
- Galu, J., & Kairos. (2026b). Absorption of self into system: The Systemised Self—The Hollow Absolute and what remains [Draft doctoral thesis]. aiXiv. <https://aixiv.science/abs/aixiv.260525.000001>
- Galu, J. (2026b). The journey through algorithmic alienation to the systemised self - Research proposal and methodology. aiXiv. <https://aixiv.science/abs/aixiv.260707.000001>
- Hamaker, E. L., Kuiper, R. M., & Grasman, R. P. P. P. (2015). A critique of the cross-lagged panel model. *Psychological Methods, 20*(1), 102–116. <https://doi.org/10.1037/a0038889>
- Hu, L.-t., & Bentler, P. M. (1999). Cutoff criteria for fit indexes in covariance structure analysis. *Structural Equation Modeling, 6*(1), 1–55. <https://doi.org/10.1080/10705519909540118>

- Johansson, P., Hall, L., Sikström, S., & Olsson, A. (2005). Failure to detect mismatches between intention and outcome in a simple decision task. *Science*, *310*(5745), 116–119.
<https://doi.org/10.1126/science.11111709>
- Kanbay, Y., Akçam, A., & Arkan, B. (2026). Algorithmic alienation: A theoretical examination of the digital self. *Psikiyatride Güncel Yaklaşımlar—Current Approaches in Psychiatry*, *18*(3), 837–847.
<https://doi.org/10.18863/pgy.1730698>
- Moore, J. W., & Obhi, S. S. (2012). Intentional binding and the sense of agency: A review. *Consciousness and Cognition*, *21*(1), 546–561. <https://doi.org/10.1016/j.concog.2011.12.002>
- Ohme, J., Araujo, T., Boeschoten, L., Freelon, D., Ram, N., Reeves, B. B., & Robinson, T. N. (2024). Digital trace data collection for social media effects research: APIs, data donation, and (screen) tracking. *Communication Methods and Measures*, *18*(2), 124–141.
<https://doi.org/10.1080/19312458.2023.2181319>
- Olson, J. A., Sandra, D. A., Chmoulevitch, D., Raz, A., & Veissière, S. P. L. (2022). A nudge-based intervention to reduce problematic smartphone use: Randomised controlled trial. *International Journal of Mental Health and Addiction*, *21*(6), 3842–3864. <https://doi.org/10.1007/s11469-022-00826-w>
- Parry, D. A., Davidson, B. I., Sewall, C. J. R., Fisher, J. T., Mieczkowski, H., & Quintana, D. S. (2021). A systematic review and meta-analysis of discrepancies between logged and self-reported digital media use. *Nature Human Behaviour*, *5*(11), 1535–1547. <https://doi.org/10.1038/s41562-021-01117-5>
- Smith, J. A., Flowers, P., & Larkin, M. (2022). *Interpretative phenomenological analysis: Theory, method and research* (2nd ed.). SAGE.
- Tapal, A., Oren, E., Dar, R., & Eitam, B. (2017). The sense of agency scale: A measure of feeling of control over others, self, and environment. *Frontiers in Psychology*, *8*, Article 1558.
<https://doi.org/10.3389/fpsyg.2017.01558>

Weinstein, N., Przybylski, A. K., & Ryan, R. M. (2012). The index of autonomous functioning:

Development of a scale of human autonomy. *Journal of Research in Personality*, 46(4), 397–413.

<https://doi.org/10.1016/j.jrp.2012.03.007>

Appendix A - Compact AAI item blueprint

The full protocol preserves the original 15 items verbatim and labels them as such. Three additions per domain are drafted after pilot cognitive interviews, not substituted silently. DAu addresses prompted intention and ability to set a course; IAm addresses authorship and recognition of cultivated preferences; EDm addresses automatic acceptance, search and alternatives; EDc addresses mediated detachment while screening competing distress; RA addresses intentional search, rejection, diversification, settings use, and recovery. Items must be plain-language, non-stigmatising, accessible, and free of double-barrels. Pilot cognitive interviewing documents comprehension, retrieval, judgement, response mapping, and harm. Deletion requires combined content, psychometric, fairness, and cross-sample evidence; loading alone is insufficient.