

Research Methodology for Testing the Algorithmic Alienation

Index and Systemised Self Hypothesis v1

Galu & Kairos

Abstract

This research methodology specifies a staged, mixed-method programme to develop and test the proposed Algorithmic Alienation Index (AAI) and the hypothesised Systemised Self (SS). It makes a necessary distinction between Algorithmic Alienation (AA), an established theoretical construct with partial qualitative empirical grounding in Kanbay, Akçam, and Arkan (2026) and Arkan, Kanbay, and Akçam (2026), and the newer claims under examination. The AAI is a proposed operational instrument, not a validated measure of AA; SS is a falsifiable hypothesis, not a diagnosis or demonstrated endpoint. International English-speaking adults aged 18–60 would be recruited with quotas across country/region, age, gender, education, and self-reported intensity of personalised-feed use. The proposed hybrid data strategy makes a controlled research feed the primary objective behavioural source; participant side, consented tracking and data donation provide naturalistic evidence; voluntary platform backend data are supplementary rather than a feasibility condition. The programme comprises content validation and cognitive interviews, exploratory factor analysis (EFA; target $n = 800$), independent confirmatory factor analysis (CFA; target $n = 1,200$), a 12-month cohort, a controlled feed friction shock experiment, a field randomised technical encouragement study, cognitive tasks, interpretative phenomenological analysis (IPA) interviews, and a blinded longitudinal expert all-data (LEAD)-style panel. Approximately 30-items are included in the AAI for psychometric reduction throughout testing. Competing one factor, four correlated factor, second order, and cognitive affective models will be assessed. The theorised weights (.30/.25/.25/.20) and resistance multiplier are explicitly treated as hypotheses; base, resistance, and adjusted scores will always be reported separately. SS classification will be evaluated against a blinded LEAD–IPA multimethod reference standard that permits indeterminate and alternative explanation conclusions, avoiding both a false claim and incorporation bias. Direct randomisation supports causal inference in the controlled feed; field encouragement will use intention to treat primary

estimates and instrumental variable complier average causal effect analyses secondarily. Longitudinal random intercept cross lagged panel models intend to address temporal ordering without causal overclaim. Governance prioritises data minimisation, no message-content capture by default, participant control, encryption, audited access, deletion, and a data protection impact assessment.

Keywords: *algorithmic alienation, Algorithmic Alienation Index, Systemised Self, scale development, phenomenological inversion, digital trace data, causal inference, mixed methods*

Contents

1. Purpose, epistemic status, and research questions

2. Conceptual model and operational definitions

3. Design overview and setting

4. Sampling, recruitment, and feasibility

5. Measures and data sources

6. Procedures and research phases

7. Hypotheses, decision rules, and falsification

8. Quantitative analysis plan

9. Qualitative, reference standard, and integration plan

10. Causal inference strategy

11. Ethics, data governance, and participant protection

12. Project management, timeline, and resources

13. Dissemination and limitations

Glossary

References

Appendix A. Candidate AAI item blueprint

Appendix B. LEAD-IPA reference standard evidence matrix

Appendix C. Worked scoring and interpretation scenarios

Appendix D. Minimum analysis reporting checklist

1. Purpose, Epistemic Status, and Research Questions

1.1 Rationale

Algorithmically curated systems organise increasingly consequential choices: what to watch, read, buy, listen to, follow, and sometimes an extremity of user engagement, perhaps even how to think or act. Their benefits can include relevance, accessibility, reduced search costs, and assistance for people facing cognitive, sensory, or time constraints. Accordingly, reliance on a recommendation is not intrinsically evidence of compromised autonomy. The present programme is designed to distinguish informed, revisable delegation from a proposed pattern of reduced self authorship and inflexible behavioural dependence.

Algorithmic Alienation (AA) is treated here as the established conceptual starting point. Kanbay et al. (2026) provide the theoretical examination, and Arkan et al. (2026) provide partial qualitative grounding for experiences described as diminished autonomy, identity ambiguity, eroded decision-making, emotional disconnection, and resistance. This study does not purport to validate, replace, or redescribe that published construct.

The study tests two subsequent propositions. First, the Algorithmic Alienation Index (AAI) is a proposed item based operationalisation of four AA-related dimensions (i) Diminished Autonomy (DAu), (ii) Identity Ambiguity (IAm), (iii) Eroded Decision-Making (EDm), and (iv) Emotional Disconnection (EDc), with Algorithmic Resistance (RA) modelled separately. Its item set, dimensionality, weights, multiplier, thresholds, and practical interpretation remain unvalidated at study commencement. Second, the Systemised Self (SS) (Galu & Kairos, 2026a) is a hypothesised configuration in which severe algorithmic dependence and low effective resistance coexist with phenomenological inversion: heteronomous curation is experienced and narrated as one's own liberation, authenticity, or freedom. SS is not a clinical condition, a label for frequent users, or an inference licensed by a single score.

The central methodological risk is circularity. If SS means ‘a person who says recommendations are freedom’, then a self-report proves only a self-report. Conversely, if it means ‘a person who follows recommendations’, it pathologises rational delegation and disability accommodation. The protocol therefore requires convergent but non identical evidence: phenomenological accounts, validated self-report scales, behaviour in a controlled feed, naturalistic participant-authorized traces, cognitive task performance, and an independent blinded panel judgement.

1.2 Objectives and research questions

The primary objective is to determine whether a psychometrically defensible AAI can be developed from the theory-derived items and candidate pool. Secondary objectives are to test whether the AAI has incremental explanatory value beyond related constructs and exposure, whether RA behaves as an effective moderator, and whether the SS hypothesis receives support, requires revision, or is falsified.

Research Question 1: asks whether AA related self report is best represented by one general factor, four correlated dimensions, a second-order factor, or a cognitive-affective structure.

Research Question 2: asks whether the weighted scoring rule predicts pre-specified behavioural and experiential criteria as well as or better than alternatives.

Research Question 3: asks whether high AAI profiles predict recommendation-led path dependency, lower independent search, reduced source diversity, and lower adaptability when curation is withdrawn.

Research Question 4: asks whether phenomenological inversion can be identified reproducibly, without treating a numeric threshold as a gold standard.

Research Question 5: asks whether controlled manipulation of curation and encouragement to reduce personalisation change relevant outcomes under clearly stated causal assumptions.

2. Conceptual Model and Operational Definitions

2.1 AAI scoring hypotheses

The complete candidate pool will initially preserve the original (previous research proposal, Galu & Kairos, 2026b) five point response structure, scored 1 = Never, 2 = Rarely, 3 = Sometimes, 4 = Often, and 5 = Always, and transformed respectively to 0, .25, .50, .75, and 1 for testing the provisional formula. Cognitive interviews will assess whether the frequency anchors fit every item and whether a seven-point format improves comprehension or information without changing the construct. Any revised format will be tested experimentally rather than silently substituted. Dimension scores will initially be the mean of retained, reverse keyed items. The original theory derived formula is retained for explicit testing:

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

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

DAu receives .30 because the proposal gives compromised self-direction theoretical priority. IAm and EDm receive .25 each because authorship ambiguity and automatic choice are proposed to reinforce one another. EDc receives .20 because affective detachment may be downstream and is less specific; it can arise from depression, anxiety, fatigue, loneliness, or unrelated stressors. RA is multiplicative because it is hypothesised to attenuate the whole profile rather than a single domain.

These are hypotheses, not observed effect sizes. The study will report DAu, IAm, EDm, EDc, RA, AAI_base, and AAI_adjusted together. It will never present a low adjusted score alone as evidence of low underlying difficulty, because high RA can lower an adjusted score despite high base alienation. Candidate scoring alternatives include equal weights, CFA-derived factor-score composites, penalised regression weights estimated only in a training set, and

expert informed weights. The multiplicative RA model will be compared with additive moderation and interaction models. Scores are research variables, not diagnostic categories.

The previously proposed (Galu & Kairos, 2026b) SS screening boundary of $AAI \geq .95$ and $RA \leq .05$ is retained only as a high specificity candidate review flag, pending calibration. It will not establish SS. The study will estimate threshold performance across the full continuum and will not imply a true gold standard.

2.2 Phenomenological inversion and SS

Phenomenological inversion is operationalised prospectively as a discrepant configuration, not as low alienation self report by itself. A participant may be provisionally considered for SS review when all of the following evidence domains are present: (a) high structural dependence on controlled-feed and/or naturalistic indicators relative to the study distribution; (b) low demonstrated effective resistance or adaptability under friction; (c) a phenomenological account in which curated direction is consistently described as self authored freedom, authenticity, or relief from deliberative burden; and (d) low felt estrangement relative to the person's structural-dependence profile, after consideration of alternative explanations.

Self report operationalisation includes the Sense of Agency Scale (Tapal et al., 2017), an established authenticity measure such as the Authenticity Scale (Wood et al., 2008), the Index of Autonomous Functioning (Weinstein et al., 2012), and Self-Concept Clarity (Campbell et al., 1996). These measures are indicators, not proxies for a metaphysical self. A pattern compatible with inversion is high felt agency/authenticity or autonomous functioning and low reported estrangement alongside objective dependence and reduced recovery. It is not sufficient, because adaptive use may genuinely raise agency.

Independent behaviour operationalises dependence through recommendation initiated action share, manual-search initiation, direct navigation, source/topic diversity, repeated acceptance

after recommendations, and recovery of previously declared goals when curation is changed. Cognitive evidence includes source monitoring/choice-blindness susceptibility, goal pursuit adaptability, and decision latency. Intentional binding is included only as exploratory because its validity as a measure of agency is contested and task effects are substantial (Moore & Obhi, 2012).

3. Design Overview and Setting

This is an international, sequential explanatory and convergent mixed method programme. It has five linked phases rather than one undifferentiated sample. Phase 1 establishes comprehensibility and content coverage. Phase 2 develops the scale through EFA. Phase 3 independently validates its structure by CFA. Phase 4 follows a cohort for 12 months. Phase 5 embeds controlled causal experiments, cognitive tasks, IPA interviews, and a blinded LEAD-style reference process. Recruitment and study materials will be in English; eligibility is English speaking adults aged 18–60 residing in countries or territories where English is used for daily online activity. This is an international English-speaking sample, not a claim of global representativeness or cultural universality.

The controlled research feed is the primary behavioural infrastructure. It will be a purpose-built web environment containing ethically licensed, neutral-interest content in several domains (news like explainers, entertainment clips, consumer options, and educational content). It will log study relevant events such as displayed option identifiers, rank, condition, clicks, searches, dwell intervals, saves, rejections, and stated goals. It will not imitate an identifiable commercial platform or infer sensitive traits. A controlled environment permits known exposure and direct randomisation without depending on commercial platform cooperation.

Naturalistic evidence is participant side and opt-in. A browser extension and mobile operating-system screen time export, where technically supported, will capture pre-specified aggregate

interaction metadata, such as domain/app category, event timestamp rounded to a privacy-preserving interval, whether an action began from a recommendation versus a search/direct route where detectable, and counts of searches. Data donation will allow participants to upload exports they select from their own platform account archives. This hierarchy follows established distinctions among application programming interface access, participant data donation, and screen tracking in digital trace research (Ohme et al., 2024). Message text, private messages, passwords, typed search strings, page body text, contact lists, precise location, keystrokes, screenshots, and clipboard content are excluded by default. Backend data from platforms may be accepted only under a separate protocol and consent process; it is optional and not required for the primary hypotheses.

4. Sampling, Recruitment, and Feasibility

4.1 Recruitment and eligibility

Participants will be recruited through reputable international research panels, community and university mailing lists, disability advocacy networks, and public advertisements. Recruitment materials will state that the study concerns online choice and recommendation systems, avoiding the term ‘alienation’ where that would create demand characteristics. Eligibility requires age 18–60, informed consent, functional English sufficient for measures and interviews, at least weekly use of a personalised online feed or recommender, and access to a compatible personal device for components requiring tracking. People may complete survey-only components without tracking. Exclusions are limited to duplicate enrolment, failed identity/attention safeguards, or inability to provide informed consent. Current severe distress is not an automatic exclusion; however, the consent process will include a welfare screen and referral options.

Quota targets will balance broad region (United Kingdom/Ireland, North America, Australia/New Zealand, and other English speaking or English using locations), gender, age bands (18–29, 30–44, 45–60), and education. A recruitment dashboard will monitor quotas but will not force participants to disclose protected attributes beyond what they voluntarily report. Oversampling of heavier users may improve representation of the upper AAI range; analyses will distinguish descriptive estimates from association tests and use design weights only if a defensible target population and selection model are available.

4.2 Sample targets and retention

The EFA development sample target is $n = 800$ complete baseline cases, recruited separately from all subsequent CFA participants. The independent CFA target is $n = 1,200$ complete baseline cases. With approximately 30 candidate items, this affords a practical 40 participants per item and permits split sample robustness checks; it is not represented as a universal ‘participants per item’ rule. Monte Carlo simulation using anticipated ordinal item distributions, factor loadings, correlations, and missingness will finalise targets before recruitment. The preregistration will document the simulation code, assumptions, precision goals, and stopping rule.

From the CFA sample, 720 consenting participants will enter the 12-month cohort, anticipating 25% attrition and approximately 540 retained at month 12. This permits longitudinal modelling but estimates will be accompanied by uncertainty intervals, especially for rare candidate SS profiles. A controlled feed experimental subsample of 360 will be nested across EFA/CFA-eligible participants, anticipating at least 300 analysable completers. A field encouragement sample of 500 tracking-eligible participants will be randomised, anticipating 20% loss and 400 outcome observations. The IPA subsample will comprise 48 participants purposively selected after baseline to maximise variation: approximately 16 high-base/low-RA profiles, 16 high-

base/high-RA profiles, and 16 low-to-moderate profiles, with structured consideration of high-dependence/low-estrangement discrepancies. Sampling continues only until the pre-set maximum; claims of saturation will be qualified.

The LEAD panel will review 60 deliberately information rich cases, including all cases flagged by a prespecified SS review rule up to capacity and matched comparison profiles. Attrition reduction will use proportionate staged payments, brief mobile-friendly surveys, a participant-controlled contact schedule, reminder limits, and re-consent at major data collection points. No compensation will be conditional on permitting naturalistic tracking.

4.3 Feasibility and contingencies

A 40 person technical pilot will establish whether the controlled feed, extension, export instructions, consent interface, and event taxonomy are understandable and stable. Its data will not enter main analyses unless preregistered as eligible after no material changes. The primary programme remains feasible if no platform grants backend access: controlled feed logs answer experimental questions and participant-side tracking/data donation test naturalistic associations. Optional platform partnerships will require written data specifications, legal review, a separate data sharing agreement, costed secure transfer, and a public statement of which fields were received. If tracking uptake is below 40% among consented cohort members, tracking analyses will be explicitly labelled selected subsample analyses and recruitment will be expanded only after ethics approval.

5. Measures and Data Sources

5.1 AAI candidate pool and covariates

Appendix A specifies 30 candidate AAI items: the original 15 theory-derived items (Galu & Kairos, 2026b) reproduced verbatim and marked as retained-for-testing, plus 15 expansion

items. Three items provisionally cover each of DAu, IAm, EDm, EDc, and RA in the original set; each domain gains three candidates. Cognitive interviews may recommend wording changes for clarity, but the original and amended versions will be distinguishable in the study record. Negatively worded expansion items will be minimised because wording factors can be mistaken for substantive dimensions. The full pool is administered in EFA and CFA; reduction will use content, psychometric, and fairness criteria rather than loading alone, consistent with staged scale-development principles (Boateng et al., 2018).

Covariates include age, gender, country/region, education, employment, financial strain, disability accommodation needs relevant to technology use, digital literacy, self-reported hours of relevant digital use, platform classes used, depression and anxiety screening measures where ethically approved, and baseline preference for delegation. These are not excuses to “control away” lived experience. Their roles will be defined per model to avoid indiscriminate covariate adjustment and collider bias.

Construct validity measures include the Bergen Facebook Addiction Scale as an adjacent problematic-use measure (Andreassen et al., 2012), psychological need satisfaction/frustration measures grounded in self-determination theory (Deci & Ryan, 2000), perceived control, and social desirability. The final selection will use instruments with appropriate licences and demonstrated English-language measurement properties. Digital media self-report is retained as a subjective measure but not substituted for logged use, given observed discrepancies (Parry et al., 2021).

5.2 Behavioural measures

The controlled feed will implement a within-person crossover task. Participants first state concrete goals (for example, find two reliable accounts of a topic and choose one article to save). In personalised conditions, ranking uses preferences elicited within the study; in non-

personalised conditions, ranking is randomised within relevance strata. Autoplay and “because you watched” prompts are present or removed according to randomisation. Key outcomes are: recommendation-initiated action proportion; manual search count; unique-source count; entropy of source/domain selection; rate of active rejection; proportion of stated goals completed; time to goal completion; and post-change recovery, defined as successful goal-directed action after personalisation/autoplay removal. Engagement time is secondary because reduced engagement can be a rational response to lower service quality.

For naturalistic data, a pre-specified metadata dictionary will distinguish observed, inferred, and unavailable referral source. The primary naturalistic path-dependency metric is recommendation-led eligible actions divided by all eligible actions for which origin is observable. Analyses will report denominator coverage and never interpret unavailable source as recommendation-led. Behavioural diversity will be calculated within platform category and observation window; it is not assumed that broad diversity is always beneficial.

5.3 Cognitive tasks

Choice blindness/source monitoring tasks will ask participants to identify the source and authorship of presented options after some choices or rationales are covertly switched in an ethically disclosed, debriefed paradigm. This adaptation draws on evidence that participants can fail to detect mismatches between intended and presented outcomes (Johansson et al., 2005). The endpoint is detection accuracy, confidence calibration, and source-attribution error, not deception susceptibility as a personality trait. Goal-pursuit adaptability tasks impose a change in tool, route, or ranking after participants state a goal; outcomes are persistence, alternative generation, search use, and successful goal completion. Decision latency will be measured from option presentation to initial selection, adjusted for device type, reading time, task difficulty, trial order, and motor accessibility. Very fast decisions can reflect expertise and

will not be interpreted alone. Intentional-binding tasks, if feasible, will be exploratory and separately reported.

5.4 Interview measures

Semi-structured IPA interviews will explore how participants experience recommendations, the perceived authorship of preferences, moments of constraint or liberation, reactions to removed curation, counterfactual alternatives, and explanations of observed behavioural patterns. Interviewers will avoid presupposing alienation or SS. Participants will view a carefully selected, privacy-minimised personal event summary only after they have given an initial account, allowing exploration of concordance without treating logs as superior to experience. Interviews also directly ask about occupational necessity, caregiving, language access, disability accommodation, and deliberate delegation.

6. Procedures and Research Phases

6.1 Phase 1: content validity and cognitive interviews

An independent content panel of 8–12 scholars and practitioners with expertise in autonomy, psychometrics, recommender systems, disability studies, qualitative research, and digital rights will rate item relevance, clarity, and potential harm. A content-validity index will summarise ratings but will not mechanically dictate deletion. Thirty adults reflecting quota dimensions will complete think-aloud cognitive interviews using verbal probing: comprehension, retrieval, judgement, response mapping, and sensitivity. Revisions will be logged with rationale. This phase will identify, rather than conceal, construct ambiguity and culturally specific phrasing.

6.2 Phase 2: EFA development

The EFA sample completes the 30 items, external measures, and a short controlled-feed session. Polychoric correlations and robust estimation appropriate to ordinal responses will be

used. Parallel analysis, minimum average partial procedures, scree inspection, interpretability, and residuals will jointly guide factor retention. Oblique rotation is expected because dimensions are theoretically related. Items will be considered for removal where loadings are weak, cross-loadings material, residual dependence excessive, wording is unclear, or content is redundant; no factor will be represented by fewer than three conceptually adequate items. The original 15 remain in all development analyses even if a future short form is proposed.

6.3 Phase 3: independent CFA validation

The CFA sample is recruited independently, without reusing EFA cases. Pre-registered models are: (a) one factor; (b) four correlated factors, with RA separate as a moderator/criterion rather than folded into base alienation; (c) a second-order AA factor over DAu, IAm, EDm, and EDc; and (d) a cognitive-affective model in which DAu and EDm form a cognitive-behavioural factor and IAm and EDc form an identity-affective factor. A bifactor exploratory sensitivity model may be considered only if theoretically coherent, and indices will not be used to claim essential unidimensionality automatically. CFA will test ordinal robust weighted least squares models and robust maximum likelihood sensitivity models.

6.4 Phase 4: 12-month cohort

The cohort completes baseline, month 3, month 6, month 9, and month 12 assessments. Controlled-feed sessions occur at baseline, month 6, and month 12. Tracking/data donation are continuous or repeated for two-week windows around each wave according to participant choice. The primary longitudinal question is temporal ordering among AAI dimensions, RA, dependence metrics, and agency/authenticity indicators. Random-intercept cross-lagged panel models (RI-CLPM) separate stable between-person differences from within-person deviations (Hamaker et al., 2015). They do not establish causation in the absence of randomisation and unmeasured-confounding control.

6.5 Phase 5: experimental and qualitative components

In the controlled-feed experiment, participants are randomly allocated to personalised versus non-personalised ranking and, factorially, to low-friction versus friction-shock conditions. Friction shock removes autoplay, suppresses personalised recommendations, adds a neutral “search or browse categories” interstitial, and maintains access to all content. It neither blocks needed resources nor penalises withdrawal. Randomisation is computer-generated with concealed allocation, stratified by baseline AAI_base tertile and device type. The primary contrast is direct assignment effect on post-manipulation goal-pursuit adaptability and recommendation-led action. Manipulation checks measure perceived and logged curation.

The field study randomises technical encouragement, not covert changes to personal accounts. The encouragement arm receives a supported choice architecture: instructions and optional one-click tools to turn off/limit personalisation, disable autoplay, reset recommendations, and schedule intentional-search prompts. Controls receive general digital-wellbeing information with no personalisation tools. Participants decide whether to comply. The primary estimand is intention to treat (ITT); secondary instrumental-variable estimation targets the complier average causal effect (CACE), with assignment as instrument for verified reduction in personalisation.

IPA interviews occur after an initial quantitative profile but before the LEAD panel. Interviewers are blinded to AAI scores and study labels, receiving only a neutral participant code and a minimally necessary sampling stratum. This protects accounts from interviewer confirmation bias.

7. Hypotheses, Decision Rules, and Falsification

H1 (structure): four distinguishable AA-related dimensions will fit materially better than a single factor and will have interpretable, reliable indicators. Evidence requiring revision

includes replicated preference for a simpler model, non-identifiable or unstable dimensions, or correlations so high that discriminant interpretation is untenable.

H2 (scoring): original weights and the RA multiplier will predict pre-specified behavioural dependence and adaptability at least as well as equal and empirically derived alternatives in held-out data. Failure to match alternatives within pre-specified practical margins requires replacing or abandoning the original scoring rule; the conceptual dimensions may still remain useful.

H3 (convergent and discriminant validity): higher AAI_base will be associated with more recommendation-led activity, less manual search, and lower controlled-feed adaptability, while showing related but nonredundant associations with problematic social-media use, need satisfaction, and screen time. Failure includes no incremental out-of-sample value beyond these measures or correlations indicating practical indistinguishability.

H4 (resistance): at comparable AAI_base and exposure, higher RA will predict more rejection, independent search, and recovery after curation withdrawal. A null or reversed, replicated relation challenges both RA items and the multiplicative model.

H5 (controlled friction): random assignment to friction shock will reduce recommendation-led action and, among participants with high base alienation/low RA, will reduce or slow goal recovery more than among relevant comparison profiles. If high-base/low-RA participants adapt as well as others, the claim that the profile indexes inflexible dependence is weakened or falsified.

H6 (field encouragement): random encouragement will increase verified personalisation reduction and improve independent-search or adaptability outcomes in ITT analysis. A weak first stage makes CACE uninformative; a null ITT is reported as such and not reframed as proof that personalisation is harmless or harmful.

H7 (phenomenological inversion): a non-trivial, replicable subset with high independently measured dependence and low effective resistance will report high felt agency/authenticity and low estrangement, supported by IPA accounts, more often than matched severe-AA participants who feel constrained. If extreme dependence consistently co-occurs with equal or higher estrangement, or apparently inverted cases are better explained by accommodation, expertise, or informed delegation, SS is not supported.

H8 (longitudinal directionality): within-person increases in path dependency will prospectively precede increases in AAI_base and/or decreases in RA more strongly than the reverse path, or the reverse pattern will be transparently reported. RI-CLPM results concern temporal associations, not proof of causal direction.

All hypotheses, operational variables, model syntax, exclusion rules, and primary/secondary outcomes will be preregistered before confirmation samples are examined. Statistical significance alone will not determine support. Decisions will consider effect estimates, confidence intervals, model adequacy, robustness, alternative explanations, and replication.

8. Quantitative Analysis Plan

Analyses will be conducted in reproducible scripts with versioned data dictionaries. A blinded analyst will prepare primary models using condition codes masked until model decisions are fixed. Descriptive reporting will give recruitment flow, consent proportions for each modality, device compatibility, missingness, attrition, and distributions by region and demographic characteristics.

For CFA, model fit will be interpreted collectively: comparative fit index and Tucker–Lewis index near or above .95 as desirable and .90 as a lower contextual guide; root mean square error of approximation near or below .06 with its confidence interval; standardised root mean square

residual near or below .08; residual patterns; and theoretical plausibility. Such cut-offs are heuristics, not mechanical acceptance rules (Hu & Bentler, 1999). Nested comparisons will use robust difference tests where valid; non-nested models will compare information criteria and out-of-sample fit. Reliability will include ordinal omega, coefficient H, and where appropriate alpha, with confidence intervals. Test–retest stability will be estimated across short intervals but interpreted separately from within-person change.

Measurement invariance will proceed from configural to threshold/loading, scalar/intercept where relevant, and residual invariance across broad region, gender, age band, disability-accommodation status, and major device class when sample sizes permit. Alignment or partial invariance will be used transparently if full invariance fails. Mean comparisons will not be made where the necessary level of invariance is unsupported. Differential item functioning will be examined using ordinal item-response or multiple-indicator multiple-cause sensitivity analyses.

Convergent validity will test associations with pre-specified behavioural measures; discriminant validity will examine factor correlations, heterotrait-monotrait ratios, and incremental prediction relative to exposure, screen time, problematic-use scores, and need measures. Nomological validity will test the pattern implied by theory, including positive relations between AAI_base and path dependency and a moderating/protective role for RA. Predictive models will use nested cross-validation, held-out calibration intercept/slope, Brier score where outcomes are binary, and decision-curve analysis only if a plausible research decision threshold exists.

Mixed-effects models will account for repeated trials within people and waves within people. Primary experimental models include condition, baseline AAI_base, RA, their preregistered interactions, stratification variables, and participant random intercepts. Outcome-specific

distributions will be used: logistic or binomial models for success/proportions, negative binomial or hurdle models for search counts, and robust log-normal/gamma or survival models for latency. Overadjustment for post-randomisation measures will be avoided in primary ITT models.

Missing item responses will first be described and examined against observed covariates and prior outcomes. For scale development, full-information estimators will be preferred where assumptions are credible. For regression and longitudinal analyses, multiple imputation with auxiliary variables and inverse-probability-of-retention weighting will be sensitivity approaches, not automatic repairs. Complete-case, imputed, and plausible missing-not-at-random pattern-mixture results will be compared. Participants may withdraw data prospectively, so missingness due to privacy choice is substantively informative and will be reported.

Primary outcomes are limited to prevent multiplicity from creating spurious narratives: factor-structure comparison; controlled-feed goal adaptability; recommendation-led action; and field-study ITT independent-search/adaptability composite. Within each hypothesis family, false-discovery-rate control will be used for secondary endpoints, while primary estimates receive confidence intervals and interpretation rather than a binary “significant/non-significant” label. Exploratory analyses, including intentional binding, data-driven thresholds, and unscheduled subgroup analyses, will be labelled exploratory.

Threshold assessment treats the LEAD–IPA conclusion as an imperfect, fallible reference standard. Sensitivity and specificity of proposed AAI/RA rules will be estimated with confidence intervals for classifying panel conclusions of “SS-compatible,” then repeated after classifying indeterminate cases as noncases and excluding them. Positive and negative predictive values will be reported only with the sampled prevalence and will not be transported

to the public. Calibration will assess agreement between predicted probability from a multivariable research classifier and reference-standard categories, using calibration plots and slope/intercept. No threshold will be marketed as diagnostic.

9. Qualitative, Reference-Standard, and Integration Plan

IPA is chosen because the relevant question concerns how people make sense of agency, authorship, and curated freedom (Smith et al., 2022). Two trained analysts will first work idiographically with each transcript, producing experiential claims linked to verbatim excerpts and reflexive notes. They will then develop cross-case themes within sampling strata. A third analyst will audit a purposive subset, concentrating on disconfirming cases. Analysts will maintain a reflexive record of expectations about algorithms, autonomy, and SS. Reliability coefficients are not substitutes for interpretative quality; nevertheless, an audit trail, negative-case analysis, and participant option to correct factual event summaries will enhance trustworthiness.

The LEAD–IPA reference standard is a structured expert judgement, not a gold standard. Its reporting and adjudication procedures will follow the Longitudinal, Appropriate Data, Evaluation, Independent, Necessary Data, and Guidance components of the LEADING guideline where applicable (Eijsbroek et al., 2025). A panel of five to seven members will include an autonomy/ethics scholar, clinical or health psychologist, psychometrician, recommender-systems researcher, qualitative methodologist, digital-rights/privacy specialist, and disability/accessibility expert. Members will declare conflicts and complete training on alternative explanations. Each receives a de-identified case dossier comprising interview extracts, standard-scale profiles, controlled-feed behaviour, eligible naturalistic aggregate traces, cognitive-task results, and a timeline. Crucially, dossiers exclude AAI item totals, factor

scores, AAI_base, RA, AAI_adjusted, and any algorithmic SS flag. This blinding prevents incorporation bias.

Panel members independently code each domain as supporting, contradictory, absent/insufficient, or ambiguous. They then meet using a structured LEAD process: review longitudinal information, articulate competing explanations, and record the rationale for change after discussion. Final categories are: SS-compatible; severe AA without inversion; AA-related difficulty not established; informed/adaptive delegation; accommodation or accessibility explanation; occupational/contextual necessity; alternative psychological/social explanation; and indeterminate. The alternative and indeterminate categories are protections against forced classification. Consensus is sought, but split decisions and minority rationales will be retained. A panel conclusion is valid only for the research question and available evidence; it is never returned to participants as a diagnosis.

Integration uses joint displays in the internal analysis environment, but no Markdown-style tables are required in reporting. For each case, a narrative evidence profile will place phenomenology beside behavioural and cognitive indicators, then identify convergence, complementarity, or contradiction. Quantitative-to-qualitative selection will be transparent. Qualitative findings can revise the explanatory model and future item wording but cannot be used to retroactively redefine preregistered primary outcomes.

10. Causal Inference Strategy

The controlled-feed randomisation is the primary causal design. Its estimand is the effect of being assigned to a specified curation/friction environment during a bounded task, not the effect of commercial-platform use in everyday life. Valid interpretation requires random allocation, allocation concealment, consistency of defined treatment versions, no important interference between participants, valid outcome measurement, and acceptable adherence.

Randomisation balances measured and unmeasured baseline causes in expectation; it does not eliminate attrition, demand effects, or limited external validity. Primary analysis remains ITT regardless of whether a participant uses every feed feature.

The field intervention is an encouragement design. The ITT estimand compares assignment to supported personalisation-reduction encouragement with control. For CACE, assignment is the instrument and verified reduction in personalisation is the treatment. This secondary estimate assumes relevance (assignment changes uptake), independence (assignment is random), exclusion restriction (assignment affects outcomes only through uptake, an especially contestable assumption because prompts may increase reflection), monotonicity (no defiers), and a well-defined compliance contrast. The first-stage estimate will be reported. Because exclusion can fail, CACE will be framed as assumption-dependent and supplemented by sensitivity analyses, including direct-prompt effects where measurable. Weak instruments will not be used to generate precise causal claims (Angrist et al., 1996).

The cohort evaluates temporal dynamics through RI-CLPM. Stable person differences, such as temperament or baseline occupational patterns, are represented by random intercepts, allowing cross-lagged estimates to describe deviations from a person's usual level. This is superior to a conventional cross-lagged panel model for the stated within-person question in many settings, but it remains observational. Time-varying confounding, measurement error, unequal intervals, and reciprocal processes may remain. Therefore language will be “prospectively associated with” rather than “caused.”

11. Ethics, Data Governance, and Participant Protection

The protocol requires institutional ethics approval, a data protection impact assessment (DPIA), security review, and where applicable UK GDPR/EU GDPR compliance before recruitment. Consent is layered and granular: survey participation, controlled-feed logging, participant-side

tracking, account-archive donation, interviews, and recontact are separate choices. Participants may join later components without agreeing to earlier optional modalities. Consent screens will state in plain language what is collected, what is not collected, foreseeable re-identification risks, retention periods, compensation, withdrawal limits, and who can access de-identified data.

No message content is captured by default. The tracking extension implements allow-listed domains/categories and an on-device minimisation pipeline. It transforms eligible event metadata locally, excludes URLs beyond approved domain/category where possible, rounds timestamps, suppresses low-count categories, and displays a participant-facing collection ledger with pause and delete controls. Archive donation uses a participant review screen and field-level selection. Raw account archives are never necessary for the core protocol. Data collection will be paused during technical faults, and the extension will be independently code-audited before release.

Identifiers are stored separately from research data. Data in transit use current transport encryption; data at rest use managed encryption with separate keys. Role-based access gives interviewers only their assigned pseudonymous records, analysts de-identified analytical extracts, and a data steward restricted linkage access. Every access, export, transformation, and permission change is recorded in immutable audit logs reviewed monthly by the data governance lead. External collaborators receive only approved de-identified extracts in a secure analysis environment; row-level trace data are not publicly released. Reproducible synthetic or aggregate outputs and analysis code may be shared after disclosure review.

The retention schedule will specify: direct identifiers deleted within 30 days of study closure unless participants separately consent to recontact; raw event-level controlled-feed and tracking records retained for five years after final publication or institutional minimum

requirements, then securely destroyed; de-identified derived analysis data retained for ten years where lawful and consented. A participant may withdraw without explanation. Future collection stops immediately; identifiable linkage is deleted where feasible; and unanalysed identifiable records are erased. Data already aggregated, de-identified beyond re-linkage, or incorporated into completed analyses may not be retractable, which is disclosed before consent.

The study avoids deceptive account manipulation. Choice-blindness procedures use minimal, temporary experimental alteration, are disclosed in consent at an appropriate level, fully debriefed immediately, and include an option to remove the task data. Participants may experience discomfort when reflecting on digital habits; they may skip questions, stop tracking, or withdraw. A distress protocol provides local support resources, researcher escalation procedures, and emergency guidance without promising clinical care. Results will avoid moralising frequent technology use, and disability, occupation, culture, caregiving, and deliberate delegation will be treated as legitimate alternatives rather than nuisance variables.

12. Project Management, Timeline, and Resources

Months 1–3 cover governance, DPIA, technical threat modelling, advisory consultation, controlled-feed build, and pilot materials. Months 4–6 cover content-panel review, cognitive interviews, revision, and the technical pilot. Months 7–10 recruit and analyse the EFA sample; the item-reduction decision is frozen before CFA recruitment analysis. Months 11–15 recruit the independent CFA sample and preregister confirmatory models. Months 16–28 conduct cohort waves, tracking windows, controlled-feed experiment, and field encouragement. Months 20–30 conduct IPA interviews and blinded LEAD reviews after dossiers are locked. Months 29–33 complete integrated analysis, sensitivity checks, participant-facing summaries, and manuscripts. Months 34–36 support replication planning, archiving, and dissemination.

The schedule includes a three-month contingency for accessibility remediation or lower-than-expected tracking uptake.

Core personnel are a principal investigator; project manager/data steward; psychometrician; statistician; recommender-systems engineer; privacy/security engineer; qualitative lead and two IPA researchers; research assistants for recruitment and participant support; accessibility consultant; disability/community adviser; and independent ethics/data governance adviser. The budget must include panel recruitment, fair participant payment, secure hosting, independent penetration testing, accessibility testing, transcription, translation of support materials where appropriate, software licences, expert-panel honoraria, and contingency. No budget assumption depends on selling, licensing, or purchasing commercial platform backend data.

13. Dissemination and Limitations

The protocol, hypotheses, analysis code, de-identified metadata dictionary, scoring algorithms, and a record of deviations will be openly available where participant privacy permits. Publications will distinguish confirmatory from exploratory results and will report null findings, attrition, tracking-selection effects, and failed manipulations. Any short form, weights, or screening threshold will be called provisional until independently replicated in new samples and contexts.

The study has important limitations. English-only participation restricts linguistic and cultural generalisation. A research feed has higher internal control but imperfect ecological validity; participant-side traces are more naturalistic but incomplete and self-selected. Behavioural measures cannot read motives. The LEAD-IPA process improves independence and transparency but remains an expert judgement, not truth. Rare SS-compatible configurations may produce very imprecise estimates even in a large programme. These limitations are

reasons for restrained inference, not for collapsing self-report and behaviour into an untestable single construct.

Glossary

AAI_base: The proposed weighted score for DAu, IAm, EDm, and EDc before RA is applied.

AAI_adjusted: The proposed AAI_base multiplied by one minus RA; a hypothesis, not a validated severity score.

Algorithmic Alienation (AA): The established theoretical construct addressed by Kanbay et al. (2026) and partially qualitatively grounded by Arkan et al. (2026).

Algorithmic Resistance (RA): Proposed effective practices that preserve, recover, or exercise independent goal-directed activity.

CACE: The causal effect among participants whose treatment uptake changes because of random assignment, estimated under instrumental-variable assumptions.

Controlled research feed: A purpose-built study environment in which curation, alternatives, and behavioural events are known and manipulable.

Friction shock: Temporary removal or reduction of curation/autoplay plus neutral search-oriented friction, used to observe goal recovery.

LEAD: Longitudinal expert all data, a structured clinical-style consensus method adapted here as a fallible reference process.

Phenomenological inversion: The hypothesised experience of structurally heteronomous curation as authentic freedom or self-authorship.

RI-CLPM: A random-intercept cross-lagged panel model that distinguishes stable between-person differences from within-person temporal associations.

Systemised Self (SS): The hypothesised, unvalidated configuration of high dependence, resistance collapse, and phenomenological inversion.

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>
- Boateng, G. O., Neilands, T. B., Frongillo, E. A., Melgar-Quinonez, H. R., & Young, S. L. (2018). Best practices for developing and validating scales for health, social, and behavioral research: A primer. *Frontiers in Public Health, 6*, Article 149.
<https://doi.org/10.3389/fpubh.2018.00149>
- 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
- Eijlsbroek, V. C., Kjell, K., Schwartz, H. A., Boehnke, J. R., Fried, E. I., Klein, D. N., Gustafsson, P., Augenstein, I., Bossuyt, P. M. M., & Kjell, O. N. E. (2025). The LEADING guideline: Reporting standards for expert panel, best-estimate diagnosis,

- and longitudinal expert all data (LEAD) methods. *Comprehensive Psychiatry*, 141, Article 152603. <https://doi.org/10.1016/j.comppsy.2025.152603>
- Galu, J., & Kairos. (2026a). 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>
- 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>
- Wood, A. M., Linley, P. A., Maltby, J., Baliousis, M., & Joseph, S. (2008). The authentic personality: A theoretical and empirical conceptualization and the development of the Authenticity Scale. *Journal of Counseling Psychology*, 55(3), 385–399. <https://doi.org/10.1037/0022-0167.55.3.385>

Appendix A. Candidate AAI Item Blueprint

The original 15 items (Galu & Kairos, 2026b) below are reproduced verbatim from the preceding proposal and retain its five-point Never-to-Always response structure. “Expansion” identifies new candidates added for psychometric reduction; these initially use the same five-point structure. Final expansion-item wording and any response-format change remain subject to cognitive testing. The original items will remain identifiable in every analysis even if cognitive testing recommends parallel revised versions.

Diminished Autonomy. Original DAu-1: “I feel that social media platforms systematically direct me down viewing paths I did not plan to take.” Original DAu-2: “When making choices online, I select from pre-engineered options rather than expressing true personal intent.” Original DAu-3: “I feel constrained by automated systems that structure my choices and daily routine.” Expansion DAu-4: “The options I notice first usually determine what I end up doing.” Expansion DAu-5: “I postpone my own plans when a feed offers something else.” Expansion DAu-6: “I can readily begin online activity without a recommendation telling me where to go” (reverse keyed).

Identity Ambiguity. Original IAm-1: “I find it difficult to determine whether a preference is genuinely my own or has been shaped by an application profile.” Original IAm-2: “I feel that my real-world identity is becoming blurred or replaced by my online profile.” Original IAm-3: “I feel more comfortable behaving in alignment with my algorithmically predicted profile than with my organic self.” Expansion IAm-4: “I adopt interests suggested online before I have explored alternatives.” Expansion IAm-5: “It is difficult to separate my tastes from patterns a feed has learned about me.” Expansion IAm-6: “I can name experiences outside recommendations that support my important preferences” (reverse keyed).

Eroded Decision-Making. Original EDm-1: “I find it difficult to decide what to consume, read, or listen to without a recommendation feed guiding me.” Original EDm-2: “My habits for actively seeking out new, non-recommended information have weakened over time.” Original EDm-3: “I rely substantially on autoplay or recommended lists to direct my attention and entertainment choices.” Expansion EDm-4: “When a recommendation appears, my decision is often made before I reflect on it.” Expansion EDm-5: “I rely on personalised rankings to settle choices that matter to me.” Expansion EDm-6: “If a recommended option disappears, I know how to continue pursuing the same goal” (reverse keyed).

Emotional Disconnection. Original EDc-1: “I experience a sense of mental emptiness or numbness after extended periods of automated content consumption.” Original EDc-2: “My digital life feels automated, cold, and disconnected from authentic human feeling.” Original EDc-3: “I find interactions with algorithmic systems more stable and comforting than genuine human contact.” Expansion EDc-4: “After following a feed for a while, I struggle to remember why I began.” Expansion EDc-5: “Recommended activity can feel absorbing without feeling meaningful.” Expansion EDc-6: “I feel present and engaged when I follow a recommendation” (reverse keyed).

Algorithmic Resistance. Original RA-1: “I intentionally engage with content I dislike or search random topics to disrupt my algorithmic data profile.” Original RA-2: “I use ad-blockers, clear my browsing history, or disable tracking features to protect my privacy.” Original RA-3: “I schedule periods of complete digital disconnection to preserve my independent thinking capacity.” Expansion RA-4: “I intentionally search for options beyond what is recommended.” Expansion RA-5: “I turn off, reset, or work around recommendations when they do not serve my aims.” Expansion RA-6: “When personalisation is unavailable, I can still find what I need.”

Items are not a final instrument. RA is scored in its own direction, with higher values indicating more effective resistance; reverse-keying is checked through response-pattern inspection, not assumed to repair inattentive responding.

Appendix B. LEAD–IPA Reference-Standard Evidence Matrix

The panel will use the following evidence domains in a structured dossier, recording support, contradiction, ambiguity, or insufficient evidence and a narrative rationale.

Phenomenology: IPA account of authorship, freedom, estrangement, and response to alternatives. Evidence compatible with inversion is a stable account of curation as authentic self-expression despite constraint indicators. Contradiction includes explicit, sustained awareness of unwanted steering.

Validated self-report: Agency, authenticity, autonomous functioning, self-concept clarity, and distress measures. High agency/authenticity is compatible but not decisive. Inconsistency across scales triggers review rather than automatic exclusion.

Controlled-feed behaviour: Recommendation-led actions, search, rejection, goal completion, and post-friction recovery. Low recovery and low independent search support structural dependence only when task comprehension and accessibility are adequate.

Naturalistic aggregates: Eligible-action path dependency, observable manual-search share, source diversity, and longitudinal stability. Missing referral-source coverage limits inference; it is not treated as dependence.

Cognitive tasks: Source-monitoring error, choice-blindness detection, alternative generation, goal adaptability, and adjusted latency. Poor performance is nonspecific and may reflect fatigue, language, device, or motor factors.

Alternative explanations: Disability accommodation, assistive technology, occupational workflow, caregiving/time scarcity, expertise, informed delegation, mental-health symptoms, platform constraints, and cultural norms. A compelling alternative routes the case to the appropriate category rather than SS-compatible.

Longitudinal pattern: Stability or change across waves, including whether dependence precedes or follows changed agency/estrangement. Cross-sectional convergence alone is insufficient for a confident SS-compatible conclusion.

Appendix C. Worked Scoring and Interpretation Scenarios

All dimension values below are hypothetical 0–1 means. They illustrate arithmetic, not validated cut-offs or people.

Scenario 1: Low AA-related profile. DAu = .20, IAm = .15, EDm = .20, EDc = .10, and RA = .75. AAI_base = $(.30 \times .20) + (.25 \times .15) + (.25 \times .20) + (.20 \times .10) = .060 + .0375 + .050 + .020 = .1675$, reported as .168. AAI_adjusted = $.1675 \times (1 - .75) = .041875$, reported as .042. The participant searches manually and recovers goals under friction; neither score supports SS review.

Scenario 2: Moderate base profile with effective resistance. DAu = .55, IAm = .45, EDm = .50, EDc = .35, RA = .55. AAI_base = $.165 + .1125 + .125 + .070 = .4725$. AAI_adjusted = $.4725 \times .45 = .212625$, or .213. The lower adjusted value must not conceal the .473 base profile. Strong independent search and rapid recovery would support the hypothesised RA function, not SS.

Scenario 3: Severe AA-related profile without inversion. DAu = .90, IAm = .80, EDm = .90, EDc = .75, RA = .10. AAI_base = $.270 + .200 + .225 + .150 = .845$. AAI_adjusted = $.845 \times .90 = .7605$, or .761. If controlled-feed behaviour shows high path dependency and poor

recovery while IPA describes constraint and a desire for control, the LEAD panel may classify severe AA without inversion, not SS.

Scenario 4: Candidate review profile, not a diagnosis. $DAu = .98$, $IAm = .96$, $EDm = .99$, $EDc = .94$, $RA = .02$. $AAI_base = .294 + .240 + .2475 + .188 = .9695$, or .970. $AAI_adjusted = .9695 \times .98 = .95011$, or .950. This crosses the proposed screening flag. It enters blinded review only if independent evidence shows marked recommendation dependence and failed goal recovery, while the account reports high authenticity/agency and low estrangement. Accommodation, work necessity, and informed delegation remain possible alternative conclusions.

Appendix D. Minimum Analysis Reporting Checklist

Report the final sample and every consent modality; all item wording and changes; simulation assumptions; EFA and CFA decisions; all tested factor models; reliability and invariance results; attrition and missing-data diagnostics; preregistered versus exploratory analyses; all score components; behavioural denominator coverage; manipulation and first-stage checks; ITT before CACE; causal assumptions; IPA audit procedures; panel blinding and category distributions; alternative explanations; threshold uncertainty; adverse events; governance deviations; and any result that contradicts the AAI or SS hypotheses.