

From Player Types to Professional Reward Ecologies: A Bartle-Informed Open-Science Validation Agenda for Software-Role Archetypes

Taller Technologies

Maximiliano Armesto

maximiliano.armesto@tallertechnologies.com

Jan Rosen

jan.rosen@tallertechnologies.com

Martin Gonella

martin.gonella@tallertechnologies.net

Christophe Kolb

christophe.kolb@tallertechnologies.com

July 2026

Abstract

Game-derived motivational taxonomies are increasingly borrowed by organizations, learning platforms, and digital-work systems, frequently with weaker evidentiary discipline than the original behavioral settings require. This article develops a research agenda for translating Richard Bartle’s four-pole taxonomy of player orientations into a testable theory of professional reward ecologies in software work. The focal proposition is that occupations differ in the reward structures they make salient: Technical Leads should show an Achiever-weighted profile, Solution Architects an Explorer-weighted profile, Scrum Masters a Socializer-weighted profile, and Front-End Developers a professionalized Challenger profile corresponding to the historically named Bartle Killer pole, with Achiever and Explorer adjacency. A preliminary forty-case target model, built from a balanced fixed-form Bartle-style instrument, displays the predicted centroids and high role separability. The article treats those observations as a target model for validation rather than a population estimate. It then specifies a preregistered, LinkedIn-enabled, open-science protocol for collecting large-scale item-level evidence; formalizes the measurement model for forced-choice data; defines confirmatory hypotheses, power logic, validity threats, ethics safeguards, and data-release standards; and proposes an inferential architecture that integrates role ecology, psychometrics, organizational economics, and cognitive science. The resulting framework reframes Bartle types as situational reward-orientation coordinates: useful for design and theory when measured transparently, ethically, and with external replication.

Keywords: Bartle taxonomy; player types; professional roles; software work; open science; LinkedIn recruitment; forced-choice psychometrics; organizational economics; motivation; role ecology.

JEL codes: C38; C83; D91; J24; M54.

CCS concepts: Human-centered computing; Empirical studies in HCI; Social and professional topics; Computing occupations.

1 Introduction

Work roles are incentive environments. A role specifies tasks, accountabilities, comparison points, information flows, and audiences for performance. These features shape what kinds of effort feel natural, what kinds of competence become visible, and what kinds of feedback produce meaning. Economists describe this terrain through incentives, screening, complementarities, and task allocation. Cognitive scientists describe adjacent terrain through motivation, affordances, attention, social cognition, and identity. Applied computer scientists encounter the same terrain when digital systems embed dashboards, backlogs, rankings, collaboration rituals, and gamified feedback loops into everyday work. A rigorous account of professional motivation needs all three perspectives.

The present article advances such an account through a deliberately limited and testable extension of Bartle’s player-type taxonomy. Bartle’s original model mapped multiplayer game engagement onto two axes: acting versus interacting and players versus world. The resulting four orientations are Achiever, Explorer, Socializer, and Killer (Bartle, 1996). The model is memorable because it offers a compact ecology of motives: measurable progress, discovery, social interaction, and direct contest. Its weakness arises from the same compression. Later empirical research in online games found multidimensional motivational components that cut across the four boxes and cautioned against treating forced-choice Bartle categories as deep personality essences (Yee, 2006; Hamari & Tuunanen, 2014; Schneider et al., 2016). The best approach is therefore translational, conditional, and empirical: Bartle should serve as a theory-generating coordinate system for observed reward ecologies, then earn validity through external measurement.

This article builds that translation for four software roles. Technical Leads are hypothesized to emphasize Achiever rewards because they are accountable for delivery, quality standards, blockers, technical direction, and operational readiness. Solution Architects are hypothesized to emphasize Explorer rewards because they map constraints, discover feasible architectures, translate customer problems, and search large design spaces. Scrum Masters are hypothesized to emphasize Socializer rewards because they facilitate team process, support self-management, remove impediments, and protect interaction quality. Front-End Developers are hypothesized to emphasize a professionalized Challenger orientation corresponding to Bartle’s historical Killer pole: the legitimate workplace counterpart of direct contest, benchmarked craft, adversarial debugging, visible comparison, interface performance, and winning defined technical contests. In this manuscript, *Challenger* is the occupational label; *Killer* is retained only as the historically established Bartle term.

The empirical starting point is a preliminary target model consisting of forty role-labeled respondents, ten per role, scored on a balanced thirty-item Bartle-style forced-choice form. The target model strongly matches the proposed mapping: Technical Leads center on Achiever, Solution Architects on Explorer, Scrum Masters on Socializer, and Front-End Developers on Challenger/Killer with Achiever and Explorer adjacency. These observations are valuable because they define a high-resolution prior for external validation. They are insufficient as population evidence because the sample is small, score-level, and drawn from a controlled laboratory-style setting. The central contribution of this article is therefore a complete validation agenda: a preregistered, LinkedIn-enabled, open-science data-collection and analysis protocol that can transform a compelling target pattern into a replicable scientific claim.

The article makes five contributions. First, it formalizes *professional reward ecology* as a bridge between game motivation theory and occupational analysis. Second, it refines Bartle’s four poles

into workplace-safe constructs, especially by translating the Killer pole into Challenger behavior grounded in lawful contest and craft. Third, it reports the preliminary target centroids and separation diagnostics with explicit evidentiary boundaries. Fourth, it proposes an item-level psychometric model for forced-choice role data, including role centroids, heterogeneity, and measurement invariance. Fifth, it specifies a public validation design for LinkedIn recruitment that respects platform norms, consent, privacy, open science, and the risk of occupational labeling.

2 Theoretical foundations

2.1 Bartle’s four-pole motivational ecology

Bartle’s original formulation emerged from MUDs, a domain where participants share a persistent virtual environment and create social and strategic consequences for one another. The four ideal-typical orientations can be represented as a two-axis system. Achievers act on the world; Explorers interact with the world; Socializers interact with players; Killers act on players (Bartle, 1996). The framework is elegant because each pole combines an object of attention with an action mode. The model therefore describes a system ecology rather than a mere preference list.

A workplace role gives an actor a recurrent object of attention and a recurrent action mode. The Technical Lead repeatedly acts on delivery systems and codebase standards. The Solution Architect repeatedly interacts with hidden constraints and design spaces. The Scrum Master repeatedly interacts with teams and group process. The Front-End Developer repeatedly acts in public-facing, benchmarked, contestable interface space. Under this interpretation, Bartle profiles are coordinates in a reward ecology: They summarize what a role repeatedly asks a person to notice, pursue, and value.

Figure 1 illustrates the classical axes and overlays preliminary role centroids derived from the target means. The horizontal coordinate is player orientation minus world orientation; the vertical coordinate is action orientation minus interaction orientation. Hence:

$$x_i = (S_i + K_i) - (A_i + E_i), \quad y_i = (A_i + K_i) - (E_i + S_i), \quad (1)$$

where A_i , E_i , S_i , and K_i are the Achiever, Explorer, Socializer, and Killer/Challenger scores for unit i (a respondent or a role centroid). The player axis x_i contrasts the two player-directed poles, Socializer and Killer, against the two world-directed poles, Achiever and Explorer; the action axis y_i contrasts the two acting poles, Achiever and Killer, against the two interacting poles, Explorer and Socializer. The figure is a descriptive projection, rather than a full latent-space model. Its value is conceptual: It shows why the four roles can be interpreted as ecology locations rather than categories with impermeable borders.

2.2 Empirical cautions from game-motivation research

The strongest use of Bartle begins with its limits. Yee’s factor-analytic work on massively multiplayer online role-playing games found ten motivational subcomponents under three broad components: achievement, social, and immersion (Yee, 2006). The findings weakened any claim that the four Bartle categories form mutually exclusive personality classes. Hamari and Tuunanen’s meta-synthesis likewise identified a broader dimensional structure across player-type research, including achievement,

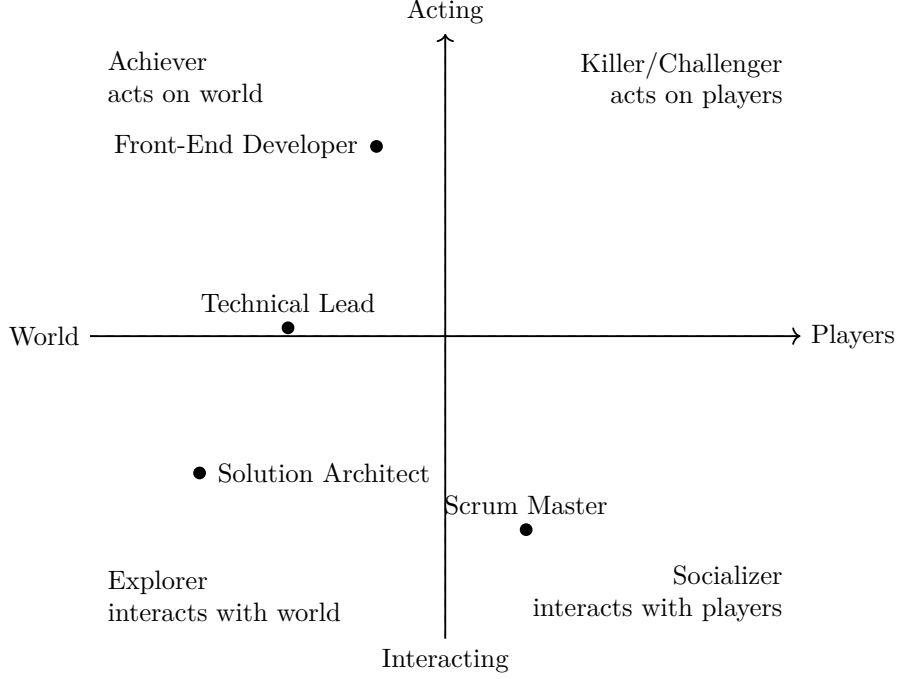


Figure 1: Bartle-derived role-space projection for the preliminary target model. Projection uses target means: $x = (S + K) - (A + E)$ and $y = (A + K) - (E + S)$. Coordinates are divided by 25 for plotting.

exploration, sociability, domination, immersion, intensity, and demographic dimensions (Hamari & Tuunanen, 2014). Schneider and colleagues applied Bartle’s taxonomy in a real-life setting and found empirical links accompanied by blurred profile boundaries and intersections (Schneider et al., 2016).

These findings motivate the research design in this article. Bartle categories are treated as observable coordinates rather than fixed essences. Scores are expected to be multidimensional. Role-level distributions can overlap. External validation needs item-level forced-choice data, convergent measures, measurement-invariance tests, and predeclared analytical rules. The translation to software occupations therefore requires psychometrics and field sampling rather than rhetorical analogy.

2.3 Professional reward ecology

A professional reward ecology is the structured set of affordances, incentives, social comparisons, task constraints, and feedback signals that repeatedly rewards certain orientations in a role. The construct can be written as a tuple:

$$\mathcal{R}_r = \langle T_r, A_r, C_r, F_r, V_r, G_r \rangle, \quad (2)$$

where T_r is the task distribution, A_r the accountability structure, C_r the comparison regime, F_r the feedback cycle, V_r the visibility of output, and G_r the governance context. A Bartle-weighted role

Table 1: Role-ecology hypotheses.

Role	Primary Bartle-weighted expectation	Reward ecology logic	Expected secondary structure
Technical Lead	Achiever	Delivery ownership, standards, planning, blocker removal, operational readiness, and visible closure reward measurable execution.	Socializer tail from coordination, coaching, and team health.
Solution Architect	Explorer	Solution design, constraint mapping, problem translation, system decomposition, and technical advisory work reward discovery across a design space.	Achiever tail from governance, sign-off, optimization, and accountable solution ownership.
Scrum Master	Socializer	Facilitation, event stewardship, impediment removal, collaboration, and support for self-management reward interaction with teams.	Moderate Achiever from process metrics and delivery-adjacent accountability.
Front-End Developer	Challenger/Killer	Interface craft, visible comparison, adversarial debugging, performance benchmarks, accessibility constraints, and public surface quality reward contestable excellence.	Achiever and Explorer adjacency from implementation closure and edge-case discovery.

The occupational label *Challenger* translates Bartle’s historical *Killer* pole into legitimate professional conduct: competitive craft, benchmark orientation, and direct problem conquest.

hypothesis maps this ecology to a vector of expected motivational salience:

$$\boldsymbol{\mu}_r = (\mu_{rA}, \mu_{rE}, \mu_{rS}, \mu_{rK}), \quad \sum_{t \in \{A, E, S, K\}} \mu_{rt} = 200. \quad (3)$$

The sum of 200 reflects the Bartle-Quotient style scale used in the target instrument, where each forced-choice response awards one point to one of two types and each type appears in the same number of response opportunities. The key hypothesis is comparative: Role ecologies differ in the ordering and spacing of the vector entries.

3 Role hypotheses

Table 1 summarizes the four role mappings. Each mapping is a claim about role affordances rather than innate personality. The same person can exhibit different orientations across systems; the same title can contain heterogeneous work if firms allocate tasks differently. The confirmatory study therefore measures both occupational title and task ecology.

The literature on each role supports the ecology mapping. Technical-lead guidance emphasizes direction, planning, acceptance criteria, blocker removal, incidents, and technical standards (GOV.UK, 2026). Solutions-architect descriptions emphasize customer problem translation, resilient architectures, and collaboration across teams (Amazon Web Services, 2026). The Scrum Guide defines the Scrum Master as accountable for establishing Scrum, enabling team effectiveness, coaching self-management, removing impediments, and ensuring events occur productively (Schwaber &

Table 2: Balanced fixed-form pairing design.

Pairing	Items	A exposure	E exposure	S exposure	K exposure
Socializer–Achiever	5	5	0	5	0
Socializer–Explorer	5	0	5	5	0
Socializer–Killer	5	0	0	5	5
Explorer–Achiever	5	5	5	0	0
Explorer–Killer	5	0	5	0	5
Achiever–Killer	5	5	0	0	5
Total exposure	30	15	15	15	15

Sutherland, 2020). Occupational descriptions of web and digital design work emphasize website and interface layout, functions, navigation, performance, capacity, graphics integration, and client-side implementation (U.S. Bureau of Labor Statistics, 2026). These sources describe the task ecology behind the hypotheses; validation of the motivational mapping requires response data.

4 Preliminary target model

4.1 Instrument structure

The preliminary target model used a balanced fixed-form thirty-item Bartle-style forced-choice instrument. The fixed-form design follows the pairwise Bartle-test logic made available in the Andreassen/Downey item-bank documentation (Andreassen & Downey, 2001). Items were selected in equal counts from the six unordered pairings among the four types: Socializer–Achiever, Socializer–Explorer, Socializer–Killer, Explorer–Achiever, Explorer–Killer, and Achiever–Killer. Each pairing appeared five times. Consequently, each type appeared in fifteen response opportunities, as shown in Table 2.

Let $y_{ij} \in \{A, E, S, K\}$ denote the type selected by respondent i on item j , and let $m_t = 15$ denote the number of opportunities for type t . The normalized score is:

$$B_{it} = 100 \cdot \frac{1}{m_t} \sum_{j=1}^{30} \mathbb{I}(y_{ij} = t), \quad t \in \{A, E, S, K\}. \quad (4)$$

With balanced exposure, the four scores sum to 200:

$$\sum_{t \in \{A, E, S, K\}} B_{it} = 200. \quad (5)$$

This construction improves comparability across respondents because every respondent faces the same type exposure profile.

4.2 Target centroids and dominance outcomes

The preliminary sample contained forty respondents, ten per role. Table 3 presents the target centroids and within-role dispersion. These values are used in this article as design targets for confirmatory sampling and model calibration. Population estimates require item-level external data.

Table 3: Preliminary target-model means and standard deviations on the Bartle-style scale.

Role	n	Mean A	Mean E	Mean S	Mean K	SD A	SD E	SD S	SD K
Technical Lead	10	84.7	41.3	57.3	16.7	8.9	11.7	11.4	9.6
Solution Architect	10	56.7	84.0	38.7	20.7	14.5	14.1	13.6	14.9
Scrum Master	10	49.3	37.3	94.7	18.7	11.4	13.0	6.1	9.8
Front-End Developer	10	58.0	53.3	15.3	73.3	17.8	14.1	10.4	16.0

Table 4: Dominant-type outcomes in the preliminary target model.

Role	A	A/E	A/S	E	K	S
Technical Lead	8	0	2	0	0	0
Solution Architect	1	1	0	8	0	0
Scrum Master	0	0	0	0	0	10
Front-End Developer	2	0	0	1	7	0

The target pattern is coherent. Technical Leads are Achiever-weighted, with a substantial Socializer tail. Solution Architects are Explorer-weighted, with Achiever secondary. Scrum Masters are Socializer-weighted with the cleanest centroid separation. Front-End Developers show Challenger/Killer dominance with Achiever and Explorer adjacency. Table 4 provides a harsher categorical summary based on respondent-level dominant score outcomes.

Preliminary one-way separation diagnostics are large across the four score dimensions. For each dimension t , the effect-size statistic is:

$$\eta_t^2 = \frac{\sum_r n_r (\bar{B}_{rt} - \bar{B}_{.t})^2}{\sum_i (B_{it} - \bar{B}_{.t})^2}. \quad (6)$$

Table 5 reports the target-model diagnostics. The values guide external design; inferential credibility comes from the planned field sample.

A simple leave-one-out role-recovery diagnostic based on the four Bartle-style scores achieved $38/40 = 95\%$ target-model accuracy. The diagnostic should be read as a separability summary rather than a deployable classifier. In an external study, role prediction must be evaluated using held-out samples, calibration curves, preregistered thresholds, and role-ecology covariates.

5 Measurement model for external validation

5.1 Forced-choice item model

The balanced score in Equation 4 is transparent, easy to explain, and useful for role feedback. A confirmatory scientific study should analyze the item-level choices directly. For item j presenting a choice between types $a(j)$ and $b(j)$, a logistic forced-choice model can be written as:

$$\Pr(y_{ij} = a(j)) = \sigma \left[(\theta_{i,a(j)} - \theta_{i,b(j)}) + \delta_j + \gamma_{r(i),j} \right], \quad (7)$$

Table 5: Preliminary target-model role-separation diagnostics.

Scale	$F(3, 36)$	p	η^2
Achiever	13.01	6.5×10^{-6}	0.520
Explorer	25.38	5.3×10^{-9}	0.679
Socializer	97.11	2.6×10^{-17}	0.890
Killer/Challenger	44.99	2.9×10^{-12}	0.789

where $\sigma(z) = (1 + e^{-z})^{-1}$, θ_{it} is respondent i 's latent orientation toward type t , δ_j captures item-level wording bias, and $\gamma_{r(i),j}$ captures role-specific item functioning. Identification can be achieved by imposing $\sum_t \theta_{it} = 0$ or by anchoring one type. Equation 7 is a Bradley–Terry/Thurstone-style representation of pairwise motivational choice (Thurstone, 1927; Bradley & Terry, 1952).

A hierarchical role model then places respondent orientations around role centroids:

$$\boldsymbol{\theta}_i \sim \mathcal{N}(\boldsymbol{\mu}_{r(i)} + \mathbf{X}_i \boldsymbol{\Gamma}, \boldsymbol{\Sigma}_{r(i)}), \quad (8)$$

where $r(i)$ is role, $\mathbf{X}_i$ includes tenure, seniority, industry, country or region, team size, work mode, and self-reported task ecology, and $\boldsymbol{\Sigma}_{r(i)}$ captures within-role heterogeneity. The target hypotheses are restrictions on the ordering of $\boldsymbol{\mu}_r$.

5.2 Mixture structure and blurred boundaries

Game-motivation research suggests that occupational profiles may blur and intersect. The external study should therefore estimate both centroid differences and mixture structure. A role-specific finite mixture can represent subarchetypes:

$$p(\boldsymbol{\theta}_i \mid r) = \sum_{\ell=1}^{L_r} \pi_{\ell r} \mathcal{N}(\boldsymbol{\mu}_{\ell r}, \boldsymbol{\Sigma}_{\ell r}), \quad \sum_{\ell=1}^{L_r} \pi_{\ell r} = 1. \quad (9)$$

For example, Front-End Developers may include a Challenger-dominant craft subgroup, an Achiever-dominant delivery subgroup, and an Explorer-dominant experimentation subgroup. Solution Architects may contain discovery-first and governance-first clusters. Model choice should combine predictive performance, information criteria, posterior predictive checks, and interpretability (McLachlan & Peel, 2000).

5.3 Measurement invariance

External validation also requires invariance analysis. A result is substantively meaningful only when item functioning is comparable across major subgroups. The confirmatory study should test whether item parameters and role centroids remain stable across geography, gender, seniority, industry, organization size, and native language. Differential item functioning can be tested through the $\gamma_{r,j}$ terms in Equation 7, multigroup confirmatory models, and held-out calibration. Where invariance fails, reporting should shift from global role claims to subgroup-specific estimates (Meredith, 1993; Brown & Maydeu-Olivares, 2011).

6 Confirmatory hypotheses

The external study should preregister the following hypotheses. Let μ_{rt} denote the latent or normalized score centroid for role r and type t . Let $\Delta_{r,t,u} = \mu_{rt} - \mu_{ru}$.

H1: Technical Lead Achiever ecology.

Technical Leads have $\mu_{TL,A} > \mu_{TL,E}$, $\mu_{TL,A} > \mu_{TL,S}$, and $\mu_{TL,A} > \mu_{TL,K}$. A secondary condition predicts $\mu_{TL,S} > \mu_{TL,K}$.

H2: Solution Architect Explorer ecology.

Solution Architects have $\mu_{SA,E} > \mu_{SA,A}$, $\mu_{SA,E} > \mu_{SA,S}$, and $\mu_{SA,E} > \mu_{SA,K}$. A secondary condition predicts $\mu_{SA,A} > \mu_{SA,K}$.

H3: Scrum Master Socializer ecology.

Scrum Masters have $\mu_{SM,S} > \mu_{SM,A}$, $\mu_{SM,S} > \mu_{SM,E}$, and $\mu_{SM,S} > \mu_{SM,K}$. A secondary condition predicts a low Challenger/Killer centroid relative to the Front-End Developer centroid.

H4: Front-End Developer Challenger ecology.

Front-End Developers have $\mu_{FE,K} > \mu_{FE,S}$ and a high Challenger/Killer centroid relative to the other three roles. A stricter directional condition predicts $\mu_{FE,K} > \mu_{FE,A}$ and $\mu_{FE,K} > \mu_{FE,E}$, with substantial Achiever and Explorer adjacency.

H5: Role separability.

A preregistered held-out classifier using only item-derived orientations predicts the four role labels above the 25% chance baseline and maintains acceptable calibration. The classifier is a validation diagnostic rather than a hiring or selection tool.

H6: Role-ecology mediation.

Task-ecology measures mediate part of the association between occupational title and Bartle-style scores. Titles are expected to matter less after explicit measures of task accountability, visibility, social facilitation, discovery demand, and benchmark exposure enter the model.

H7: Measurement invariance.

Core item parameters remain stable enough across major demographic and professional subgroups to support pooled interpretation. Subgroup-specific reporting is required where stability is weak.

7 LinkedIn-enabled open-science validation protocol

7.1 Recruitment architecture

LinkedIn is an appropriate recruitment environment because the sampling frame is occupational and because participants can recognize the relevance of a study about software roles. A compliant design treats LinkedIn as a dissemination channel and uses a separate consented survey for data collection. The study should avoid automated extraction, comment harvesting, profile scraping, or hidden observational capture. Recruitment can use public posts, professional groups, direct invitations within platform norms, alumni networks, and snowball sharing. Each post should direct

Table 6: Open-science LinkedIn validation workflow.

Stage	Main actions	Open-science artifact
Preregistration	Declare hypotheses, exclusions, power targets, role definitions, scoring, item model, invariance tests, and feedback boundaries.	Timestamped protocol, analysis plan, survey instrument.
Pilot	Recruit about 20 respondents per role to test comprehension, completion time, item balance, and data-quality rules.	Pilot codebook, item diagnostics, revision log.
Calibration	Recruit a large diverse sample to estimate item parameters, wording bias, role centroids, and mixture structure.	De-identified item data, analysis scripts, reproducible report.
Confirmation	Recruit an independent sample or holdout cohort to test the preregistered hypotheses and classifier calibration.	Locked confirmatory report, model objects, sensitivity analyses.
Replication	Invite external labs and practitioner communities to reuse the instrument under the same codebook.	Repository release, data dictionary, versioned materials.

participants to a survey landing page containing the study purpose, eligibility criteria, consent form, estimated effort, data-use terms, debriefing, and contact information.

Eligibility criteria should include: age at least eighteen years; current employment or recent employment within the past twelve months as a Technical Lead, Solution Architect, Scrum Master, or Front-End Developer; at least six months of role tenure; English proficiency sufficient for the instrument; and consent for de-identified research use. The survey should collect occupational title, role tenure, seniority, region, industry, organization size, team size, work mode, technology stack, and a short task-ecology inventory. Optional verification can be handled through self-attested profile links stored separately from survey data and deleted after quality checks.

Table 6 summarizes the workflow.

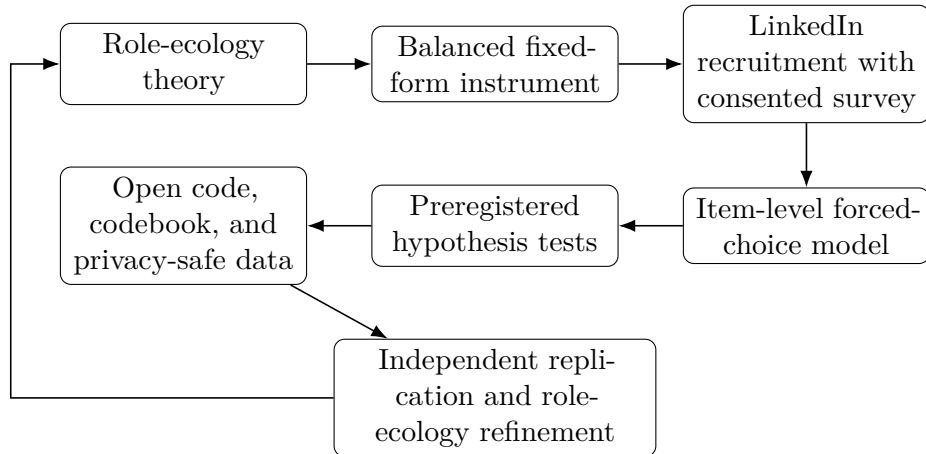


Figure 2: Validation pipeline from role-ecology theory to open replication.

7.2 Sample-size logic

The target model displays large role differences. A defensible external study should power for substantially smaller effects. For a pairwise standardized centroid difference $d_{\min}$, a conventional approximation for equal group sizes is:

$$n_r \geq \left\lceil \frac{2 \left(z_{1-\alpha/2} + z_{1-\beta} \right)^2}{d_{\min}^2} \right\rceil. \quad (10)$$

With $\alpha = .05$, $1 - \beta = .90$, and $d_{\min} = 0.35$, Equation 10 yields roughly 172 respondents per role for a single focal contrast. Because the validation program includes multiple roles, subgroup analyses, invariance tests, and holdout classification, a stronger confirmatory target is $n = 300$ per role, yielding $N = 1,200$ for the main external sample after pilot work. A staged design can use $N \approx 80$ for pilot testing, $N \approx 600$ for calibration, and $N \approx 1,200$ for confirmation.

7.3 Bias controls

LinkedIn recruitment introduces self-selection, network clustering, professional branding, and role-title heterogeneity. The validation protocol should measure and model these threats. Recommended safeguards include stratified recruitment by region, industry, seniority, organization size, and work mode; referral-source tracking through anonymous campaign codes; duplicate screening through hashed device and email checks; attention checks embedded in the survey; role-definition checks based on concrete tasks; preregistered exclusions; and early-versus-late response comparisons. Weighting can adjust descriptive estimates to target margins when credible occupational benchmarks are available. Predictive and inferential claims should be reported under both unweighted and weighted specifications.

7.4 Ethics and feedback design

The study creates potential labeling risk, especially around the historical Killer term. Participant-facing materials should use Challenger for the workplace analogue and explain that scores describe situational reward salience rather than character, morality, or employability. Individual feedback should be developmental, optional, and non-evaluative. The survey should state that results are unsuitable for hiring, promotion, termination, or surveillance decisions.

Privacy controls should include data minimization, pseudonymous respondent identifiers, separated contact information, encrypted storage, documented retention periods, withdrawal procedures, and suppression rules for small cells in public datasets. LinkedIn content should stay outside the research dataset unless a participant provides explicit consent for a separate verification process. The platform should function as recruitment infrastructure, while the scientific dataset consists of consented survey responses. This design aligns with open research norms through transparency while protecting participants from profile exposure and occupational misuse (LinkedIn Corporation, 2025; Griffiths et al., 2025).

7.5 Open-science release

The validation should follow transparent and reusable research practices associated with the TOP framework and FAIR principles: preregistered hypotheses, public materials, public code, de-identified data, sensitivity analyses, and a versioned codebook (Center for Open Science, 2026; Wilkinson et al., 2016). Data release should make artifacts findable, accessible, interoperable, and reusable through metadata, repository records, variable definitions, and licensing where consent permits. Transparency should also include reporting of recruitment scripts, campaign dates, exclusion counts, missingness, item revisions, and all deviations from the preregistered plan. The project should deposit materials in a durable repository such as OSF, an institutional repository, or another platform that supports persistent identifiers and versioning. When privacy risks restrict raw data release, the project should release a synthetic dataset, aggregate matrices, and code that reproduces analyses on approved-access data.

8 Analytical plan

8.1 Descriptive estimation

The first external report should estimate role-specific means, medians, variances, correlations, and dominance patterns for the normalized scores. Visualizations should include centroid profiles, density plots by role and type, pairwise role separation plots, and the Bartle-axis projection from Figure 1. Descriptive tables should report bootstrap confidence intervals and robust alternatives because forced-choice distributions can be skewed and bounded.

8.2 Confirmatory tests

Confirmatory tests should use multilevel models of the normalized scores and item-level forced-choice models. For score-level reporting, a multivariate model can be written as:

$$\mathbf{B}_i = \boldsymbol{\alpha} + \mathbf{R}_i\boldsymbol{\beta} + \mathbf{X}_i\boldsymbol{\Gamma} + \boldsymbol{\varepsilon}_i, \quad \boldsymbol{\varepsilon}_i \sim \mathcal{N}(\mathbf{0}, \boldsymbol{\Omega}), \quad (11)$$

where $\mathbf{B}_i = (B_{iA}, B_{iE}, B_{iS}, B_{iK})$ and $\mathbf{R}_i$ is a role indicator matrix. Planned contrasts test H1–H4. False-discovery control should be applied to secondary contrasts (Benjamini & Hochberg, 1995).

For item-level inference, Equation 7 provides the main model. The item-level model estimates latent orientation differences, role centroids, wording biases, and role-specific item functioning. Posterior predictive checks or parametric bootstraps should compare simulated response patterns with observed dominance counts, score distributions, and role separability.

8.3 Predictive validation

Role prediction should be evaluated as a diagnostic of construct separability. The preregistered pipeline should split data into calibration and holdout sets or use nested cross-validation. Candidate models can include nearest-centroid classifiers, regularized multinomial logistic regression, random forests, and Bayesian hierarchical classifiers. Evaluation should include accuracy, macro-F1, balanced accuracy, confusion matrices, calibration curves, and decision-curve analysis. Any predictive model must be explicitly barred from individual employment decisions.

8.4 Convergent and discriminant validity

A strong external study should include convergent measures. Achiever scores should correlate with preference for goal closure, scoreboards, and measurable progress. Explorer scores should correlate with curiosity, systems thinking, and tolerance for ambiguity. Socializer scores should correlate with facilitation orientation, prosocial motivation, and collaborative role identity. Challenger scores should correlate with competitive craft, benchmark orientation, and comfort with direct contest. Discriminant validity requires that the Bartle-style dimensions add explanatory value beyond broad personality measures, job seniority, and generic work engagement.

8.5 Task-ecology mediation

The most important theoretical test is whether task ecology mediates the title effect. A task-ecology inventory should measure delivery accountability, discovery demand, team-facilitation load, benchmark exposure, output visibility, incident pressure, design ambiguity, and cross-functional coordination. Mediation can be estimated as:

$$\mathbf{M}_i = \mathbf{a}_0 + \mathbf{R}_i \mathbf{a} + \mathbf{X}_i \mathbf{c} + \mathbf{u}_i, \quad (12)$$

$$B_{it} = b_{0t} + \mathbf{R}_i \mathbf{b}_{Rt} + \mathbf{M}_i \mathbf{b}_{Mt} + \mathbf{X}_i \mathbf{b}_{Xt} + e_{it}, \quad (13)$$

where $\mathbf{M}_i$ is the vector of measured task-ecology features. The decisive question is whether role titles predict Bartle-style scores mainly because titles carry task ecologies. If so, the theory becomes more general than the four chosen occupations.

9 Implications

9.1 Organizational economics

The framework links job design to motivation through role-specific reward salience. A firm that assigns delivery ownership, public incident responsibility, and throughput dashboards to Technical Leads is selecting and producing Achiever-oriented behavior. A firm that gives architects broad discovery authority and customer-problem translation produces Explorer salience. A firm that gives Scrum Masters facilitation legitimacy produces Socializer salience. A firm that makes front-end output visible, benchmarked, and tightly constrained produces Challenger salience. These patterns can affect selection, retention, task fit, team composition, and the subjective cost of effort.

The model also predicts mismatch costs. A person with high Explorer salience may struggle in a Technical Lead role dominated by short-cycle delivery scoreboards. A person with high Achiever salience may find a Scrum Master role frustrating if success arrives through indirect facilitation and group self-organization. A person with high Socializer salience may experience front-end benchmark competition as unrewarding. These are hypotheses about work design and subjective reward rather than prescriptions for exclusion.

9.2 Cognitive science and human-computer interaction

Digital work systems increasingly mediate professional attention. Backlog tools, dashboards, code-review systems, design systems, CI/CD pipelines, incident systems, and analytics panels encode

what counts as progress, discovery, collaboration, and contest. The Bartle-derived framework gives HCI researchers a compact way to reason about motivational affordances. A dashboard that emphasizes closed tickets and velocity amplifies Achiever cues. A dependency graph and sandbox environment amplify Explorer cues. A ritual board and facilitation prompts amplify Socializer cues. A performance leaderboard and public visual-diff system amplify Challenger cues.

This framing also cautions designers against one-size gamification. The same mechanic can motivate one role and irritate another. Points-badges-leaderboards fit some Achiever and Challenger ecologies, while discovery maps, collaborative rituals, and self-directed mastery cues fit different ecologies. Self-determination theory suggests that autonomy, competence, and relatedness remain central across contexts (Ryan & Deci, 2000). Bartle-derived cues should therefore support autonomous competence and relatedness rather than manipulate behavior through shallow rewards.

9.3 Management practice

Practitioners can use the framework as a diagnostic language. It can help teams ask: which rewards does this role make salient; which cues does our tooling amplify; which misalignments create fatigue; which role transitions require motivational support; and which team rituals overemphasize one orientation at the expense of others. The framework should support coaching, role clarity, and work redesign. It should never support stereotyping or deterministic talent sorting.

10 Boundaries of inference

Several boundaries are central. The preliminary target model contains forty respondents and should be treated as a high-resolution hypothesis generator. Its score-level summaries support design of the external study, while item-level field data are required for psychometric validity. Forced-choice instruments create ipsative dependencies; item-level modeling and balanced exposure reduce ambiguity. LinkedIn recruitment can produce selection bias; stratification, campaign tracking, weighting, and replication reduce that risk. Occupational titles vary across firms; task-ecology measures are necessary to interpret title effects. The Challenger translation requires care because the historical Killer label can be misread in workplaces; participant-facing feedback should use Challenger and explain the construct developmentally.

The broader conceptual boundary is equally important. Bartle-derived role profiles describe reward salience in systems. They leave virtue, intelligence, productivity, personality disorder, leadership potential, and employment worth outside their scope. A mature validation program should evaluate where the framework predicts observed experience and where richer motivational models perform better.

11 Conclusion

This article develops a scientific path from an elegant game taxonomy to a testable theory of professional reward ecologies. The preliminary target model displays a striking pattern: Achiever-weighted Technical Leads, Explorer-weighted Solution Architects, Socializer-weighted Scrum Masters, and Challenger/Killer-weighted Front-End Developers with Achiever and Explorer adjacency. The

pattern is theoretically coherent and practically useful. Its evidentiary future depends on external item-level validation.

The proposed LinkedIn-enabled open-science program offers that path. It combines preregistered hypotheses, balanced forced-choice measurement, hierarchical item modeling, mixture analysis, invariance testing, ethical recruitment, privacy-preserving data release, and explicit safeguards against workplace misuse. The result is a framework that treats professional roles as motivational ecologies. When validated with transparent data, it can contribute to organizational economics, cognitive science, HCI, and management practice by showing how work systems channel the same human motives that game worlds made visible.

Ethics statement

The proposed external study requires institutional ethics review before data collection. Recruitment should use voluntary, consented participation; data collection should occur through a dedicated survey; LinkedIn should serve only as a recruitment channel; and participant-facing feedback should avoid stigmatizing labels. Public datasets should be de-identified, with small-cell protections and withdrawal procedures described in the consent materials.

Data and materials availability

The preliminary target centroids and tables needed to understand this manuscript are embedded in the article. The planned external validation should release the preregistration, survey materials, codebook, de-identified item-level data where consent and privacy allow, analysis scripts, and reproducible reports in a persistent repository. When raw data release creates privacy risk, the project should release synthetic data, aggregate matrices, and code suitable for controlled-access replication.

Funding statement

No external funding is declared for the manuscript as written.

Competing interests statement

The authors declare no competing interests.

A Preregistered variable schema

Table 7 provides a suggested minimum codebook for the LinkedIn-enabled validation study.

Table 7: Minimum validation-study variable schema.

Domain	Variable examples	Purpose
Role identity	Role title, role family, self-described responsibilities, six-month tenure eligibility	Distinguish nominal title from actual task ecology.
Professional context	Seniority, years of experience, industry, organization size, team size, employment arrangement	Estimate moderators and improve weighting.
Recruitment trace	Campaign code, source type, referral indicator, recruitment date	Assess LinkedIn network clustering and self-selection.
Bartle item data	Item identifier, pair shown, selected option, response time	Estimate item parameters, wording effects, and latent orientations.
Task ecology	Delivery accountability, discovery demand, facilitation load, benchmark exposure, public output visibility, incident pressure	Test mediation between role title and motivational profile.
Convergent constructs	Curiosity, prosocial motivation, competitiveness, goal orientation, work engagement	Establish convergent and discriminant validity.
Data quality	Attention checks, duplicate hash, completion time, missingness indicators	Apply preregistered exclusion and sensitivity rules.
Consent and privacy	Consent version, withdrawal code, optional verification flag stored separately	Maintain ethical data governance and auditability.

B Recommended participant feedback language

Participant feedback should be developmental and role-ecological. A concise template follows.

Your profile summarizes which reward cues you tended to choose in this survey: measurable progress, discovery, social facilitation, or challenge. The result describes preferences in this task context and may vary across roles, teams, and career stages. It is designed for reflection and research. It should never be used for hiring, promotion, termination, or surveillance.

C Illustrative validation-analysis checklist

1. Freeze the preregistration and versioned instrument before launch.

2. Run pilot comprehension checks and document every item revision.
3. Confirm equal type exposure across the fixed form.
4. Inspect response-time distributions, duplicates, missingness, and attention checks.
5. Estimate score-level descriptive statistics by role.
6. Estimate the forced-choice item model and assess item-level fit.
7. Test H1–H4 as planned contrasts.
8. Test H5 through holdout classification and calibration.
9. Test H6 through task-ecology mediation.
10. Test H7 through differential item functioning and multigroup invariance.
11. Release the reproducible report, codebook, analysis code, and privacy-preserving data artifacts.

References

- Andreasen, E. S., & Downey, B. (2001). The MUD personality test. *The MUD companion*, 1, 33–35.
- Amazon Web Services. (2026). *Solutions Architect*. <https://www.amazon.jobs/content/en/teams/amazon-web-services/solutions-architect>
- Bartle, R. A. (1996). Hearts, clubs, diamonds, spades: Players who suit MUDs. *Journal of MUD Research*, 1(1), 19. <https://mud.co.uk/richard/hcds.htm>
- Benjamini, Y., & Hochberg, Y. (1995). Controlling the false discovery rate: A practical and powerful approach to multiple testing. *Journal of the Royal Statistical Society: Series B*, 57(1), 289–300. <https://doi.org/10.1111/j.2517-6161.1995.tb02031.x>
- Bradley, R. A., & Terry, M. E. (1952). Rank analysis of incomplete block designs: I. The method of paired comparisons. *Biometrika*, 39(3/4), 324–345. <https://doi.org/10.2307/2334029>
- Brown, A., & Maydeu-Olivares, A. (2011). Item response modeling of forced-choice questionnaires. *Educational and Psychological Measurement*, 71(3), 460–502. <https://doi.org/10.1177/0013164410375112>
- Center for Open Science. (2026). *Transparency and Openness Promotion guidelines*. <https://www.cos.io/initiatives/top-guidelines>
- Griffiths, M., Bloyce, D., & Law, R. (2025). LinkedIn as a research participant recruitment tool in qualitative research: A case study. *Qualitative Research Journal*. <https://doi.org/10.1108/QRJ-04-2024-0085>
- GOV.UK. (2026). *Tech lead: Role responsibilities*. <https://docs.publishing.service.gov.uk/manual/tech-lead-responsibilities.html>
- Hamari, J., & Tuunainen, J. (2014). Player types: A meta-synthesis. *Transactions of the Digital Games Research Association*, 1(2). <https://doi.org/10.26503/todigra.v1i2.13>
- LinkedIn Corporation. (2025). *LinkedIn User Agreement*. Effective November 3, 2025. <https://www.linkedin.com/legal/user-agreement>
- McLachlan, G., & Peel, D. (2000). *Finite mixture models*. Wiley. <https://doi.org/10.1002/0471721182>
- Meredith, W. (1993). Measurement invariance, factor analysis and factorial invariance. *Psychometrika*, 58(4), 525–543. <https://doi.org/10.1007/BF02294825>

- Ryan, R. M., & Deci, E. L. (2000). Self-determination theory and the facilitation of intrinsic motivation, social development, and well-being. *American Psychologist*, 55(1), 68–78. <https://doi.org/10.1037/0003-066X.55.1.68>
- Schneider, M. O., Moriya, E. T. U., Silva, A. V. da, & Néto, J. C. (2016). Analysis of player profiles in electronic games applying Bartle’s taxonomy. *Proceedings of SBGames 2016, Art & Design Track*, 617–623. <https://www.sbgames.org/sbgames2016/downloads/anais/157721.pdf>
- Schwaber, K., & Sutherland, J. (2020). *The Scrum guide: The definitive guide to Scrum*. <https://scrumguides.org/docs/scrumguide/v2020/2020-Scrum-Guide-US.pdf>
- Thurstone, L. L. (1927). A law of comparative judgment. *Psychological Review*, 34(4), 273–286. <https://doi.org/10.1037/h0070288>
- U.S. Bureau of Labor Statistics. (2026). *Web developers and digital designers*. Occupational Outlook Handbook. <https://www.bls.gov/ooh/computer-and-information-technology/web-developers.htm>
- Wilkinson, M. D., Dumontier, M., Aalbersberg, I. J., Appleton, G., Axton, M., Baak, A., Blomberg, N., Boiten, J.-W., da Silva Santos, L. B., Bourne, P. E., Bouwman, J., Brookes, A. J., Clark, T., Crosas, M., Dillo, I., Dumon, O., Edmunds, S., Evelo, C. T., Finkers, R., Gonzalez-Beltran, A., Gray, A. J. G., Groth, P., Goble, C., Grethe, J. S., Heringa, J., ’t Hoen, P. A. C., Hooft, R., Kuhn, T., Kok, R., Kok, J., Lusher, S. J., Martone, M. E., Mons, A., Packer, A. L., Persson, B., Rocca-Serra, P., Roos, M., van Schaik, R., Sansone, S.-A., Schultes, E., Sengstag, T., Slater, T., Strawn, G., Swertz, M. A., Thompson, M., van der Lei, J., van Mulligen, E., Velterop, J., Waagmeester, A., Wittenburg, P., Wolstencroft, K., Zhao, J., & Mons, B. (2016). The FAIR Guiding Principles for scientific data management and stewardship. *Scientific Data*, 3, 160018. <https://doi.org/10.1038/sdata.2016.18>
- Yee, N. (2006). Motivations for play in online games. *CyberPsychology & Behavior*, 9(6), 772–775. <https://doi.org/10.1089/cpb.2006.9.772>