

Why Driver Weights Differ: A Cost-Based Microfoundation for Segmented Technology Acceptance

Stefane Fermigier (Abilian) · sf@abilian.com

Draft working paper v0.2, 2026-07-16. Theory: every proposition carries a proof (Appendix), a numeric cross-check, and a Monte-Carlo robustness report (Section 5). This is the theoretical companion of the acceptance study in stratified-adoption-wp.md, whose H9 states that driver weights differ across capability segments.

ABSTRACT

Acceptance research increasingly compares the driver weights of technology adoption across groups, and the companion study's central hypothesis is of this form: effort and facilitating conditions bind hardest for low-capability adopters of open source, while performance and strategic alignment dominate for the capable. This paper gives that claim an economic microfoundation and derives an identification warning from it. In a random-utility adoption model, a segment's capability is a unit implementation cost that scales the effort and support terms; the quantities acceptance research estimates correspond to marginal effects of the drivers on adoption probability.

Two results follow. First, ratios of marginal effects (support to performance, effort to performance) are free of both the error density and the segment's error scale, so their ordering across segments identifies the capability ordering under any baseline and any scale heterogeneity: this is the structural content of the segment-dependence hypothesis. Second, levels of marginal effects confound structural sensitivity with the location of the segment on the adoption curve, and the confound is severe: in the Monte Carlo, a naive level comparison misorders the capability ranking on 63 percent of random configurations under a probit link and 46 percent under a logit link. Ratio comparisons are correct on all of them. Matching baselines eliminates the reversals when the error scale is common across segments and fails on 44 percent of draws when it is not, which is why the ratio route is the recommendation and matching is the fallback: the scale is the quantity the antecedent literature warns about, and it is not identified from a single binary model. The reversal condition is characterized exactly.

A fourth result carries the warning into structural-equation practice: cross-group comparisons of standardized coefficients confound structure with variance heterogeneity, differing across groups with identical structure on essentially every sampled configuration and reversing a true structural ordering on 17 percent of them, while unstandardized comparisons on invariant metrics never do. The difficulty of comparing groups in nonlinear models is established (Allison 1999; Mood 2010; Long & Mustillo 2021), and the ratio cancellation of the first result is the textbook identification fact of discrete choice; what is claimed here is the derivation of segment dependence from a cost primitive, the exact reversal condition with its quantified shares, the within-segment cross-driver ratio as a carrier of the hypothesis that is invariant to both confounds at once, and the import of the discipline into acceptance research, where it is absent.

The practical prescription for the companion study: test segment dependence on ratio contrasts, fall back on matched baselines only where a common error scale can be argued, report segment baselines alongside any path comparison, and never compare standardized coefficients across segments, none of which constitutes an empirical claim about the adoption data itself.

1. INTRODUCTION

The companion acceptance study replaces the representative user of UTAUT (Venkatesh et al. 2003) with the capability typology carried by Jullien, Viseur & Zimmermann (2025) from von Hippel (2001), Kogut & Metiu (2001), Ethiraj et al. (2005) and Jullien & Zimmermann

(2011) and hypothesizes (its H9) that the acceptance structure is segment-dependent: the weights of effort expectancy and facilitating conditions are largest for No-skill and Frontier adopters, those of performance expectancy and strategic alignment for Innovative ones. Claims of this form are common in multi-group acceptance research, usually stated as empirical regularities in search of estimation. The primitive is cost. A firm's capability position is, economically, the unit cost at which it turns adoption effort into working systems: the Innovative firm implements cheaply, the No-skill firm expensively, and support environments (facilitating conditions, community help, commercial assurance) attenuate that cost. Once adoption is a random-utility decision over this cost structure, the constructs of acceptance research map onto model objects, and the driver weights map onto marginal effects of those objects on the adoption probability. The question "do driver weights differ across segments" then has an exact answer, together with a warning about how not to measure it.

The warning itself is not new, and this paper's contribution has to be stated against a developed literature. Sociology and econometrics have studied group comparisons in nonlinear models for twenty-five years: coefficients are confounded across groups because residual variation rescales them (Allison 1999), heterogeneous choice models offer a structured fix (Williams 2009), log-odds coefficients are not comparable across groups or samples and average marginal effects are the recommended remedy (Mood 2010), the same rescaling separates from confounding across nested models (Karlson, Holm & Breen 2012), marginal effects depend on the baseline probability and are largest near one half (Norton & Dowd 2018), and comparisons are best made through predicted probabilities and marginal effects evaluated at chosen covariate values (Long & Mustillo 2021; Kuha & Mills 2020), with the field surveyed by Breen, Karlson & Holm (2018), who observe that applied methodological guidance has taken little account of it. Interactions in nonlinear models carry a related trap (Ai & Norton 2003). In discrete choice, only ratios of coefficients are identified, so the cancellation behind Proposition 1 is a textbook fact (Swait & Louviere 1993; Train 2009; Wooldridge 2010). The claim that standardized coefficients embed group-specific variances and are unfit for cross-population structural comparison is older still (Kim & Mueller 1976; Kim & Ferree 1981) and is the textbook norm in structural equation modeling (Bollen 1989). Importing such a discipline into a field's practice also has precedent (Hoetker 2007, for strategic management).

What remains this paper's own is narrower and is claimed as such. First, the derivation: segment-dependent acceptance follows from a cost primitive rather than being posited, which turns the companion's H9 from an empirical regularity into a prediction with a mechanism. Second, the mechanism at issue here is distinct from the one the sociology literature treats: Allison's confound is latent-scale heterogeneity, while Proposition 2's is the location of the density on the probability scale, so probability-scale effects fix Mood's problem and still confound structure with location when baselines differ. Third, the exact reversal condition, its genericity, and the quantified misordering shares on a declared configuration space. Fourth, the carrier: within-segment, cross-driver ratios, where the antecedent literature's ratios are cross-group ratios of one coefficient (Allison 1999) or dispersion-standardized quantities (Kuha & Mills 2020). Because a segment-specific error scale cancels from those ratios exactly as the density does, the prescription is simultaneously robust to the Allison-Mood confound and to the one identified here, while the matched-baseline remedy clears only the second (Proposition 3). Fifth, the packaging for acceptance research, where a search of the IS venues found applications of multi-group analysis but no methods treatment of which quantities carry a cross-segment structural claim (Qureshi & Compeau 2009 is the closest anchor and is silent on it).

2. MODEL

A firm in capability segment $g \in \{I, F, N\}$ (Innovative, Frontier, No-skill) decides whether to adopt an open-source solution. Its net utility is

$$U_g = b_g - \gamma_g e(1 - s) + \varepsilon,$$

where b_g collects the benefit terms (performance expectancy and strategic alignment), $e > 0$ is the intrinsic effort requirement of the technology, $s \in [0, 1]$ is the support environment (facilitating conditions, community support, commercial assurance), $\gamma_g > 0$ is the segment's unit implementation cost with $\gamma_I < \gamma_F < \gamma_N$, and $\varepsilon = \sigma_g \nu$ with ν drawn from a continuous symmetric log-concave distribution with cdf Ψ and density ψ (probit and logit as leading cases) and $\sigma_g > 0$ the segment's error scale. The firm adopts iff $U_g > 0$, so the adoption probability is $P_g = \Psi(\zeta_g)$ with

$$\zeta_g = \frac{z_g}{\sigma_g}, \quad z_g = b_g - \gamma_g e(1 - s).$$

The scale σ_g is carried explicitly because it is the object the group-comparison literature is about. It measures unobserved heterogeneity within the segment, it is not separately identified from a single binary model, and letting it differ across segments is what makes coefficients incomparable in Allison (1999) and log-odds incomparable in Mood (2010). Setting $\sigma_g = 1$ for all g recovers the homoskedastic model and every statement below.

The quantities that acceptance research reads as driver weights correspond to marginal effects on the probability scale:

$$\frac{\partial P}{\partial(-e)} = \frac{\psi(\zeta_g) \gamma_g (1 - s)}{\sigma_g}, \quad \frac{\partial P}{\partial s} = \frac{\psi(\zeta_g) \gamma_g e}{\sigma_g}, \quad \frac{\partial P}{\partial b} = \frac{\psi(\zeta_g)}{\sigma_g}.$$

Every effect is now a product of three terms: a structural sensitivity (γ_g times a technology term), the density $\psi(\zeta_g)$, which measures where the segment sits on the adoption curve (largest near the fifty-percent baseline, vanishing for segments that almost always or almost never adopt), and the reciprocal scale $1/\sigma_g$. The second term drives the level confound of Proposition 2; the third is the confound the antecedent literature treats, and Proposition 3 shows it survives the remedy that clears the second.

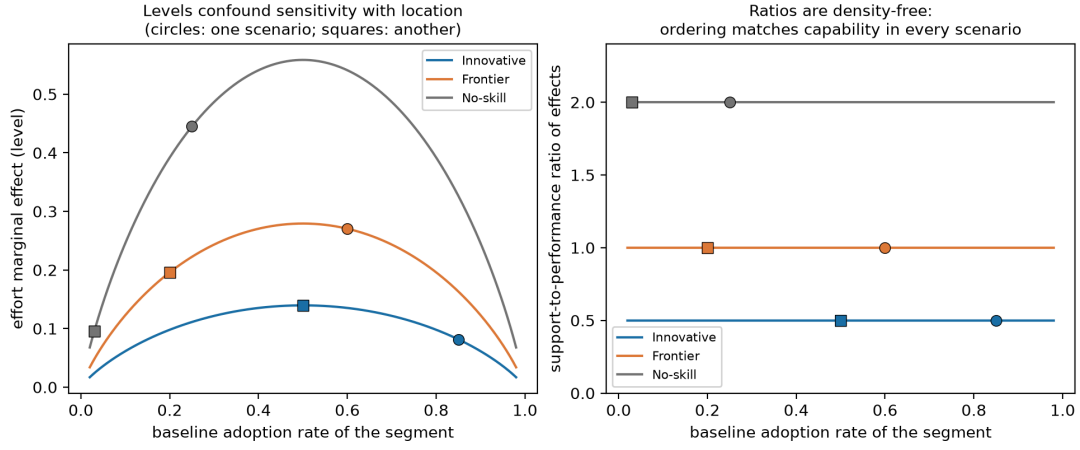
3. RESULTS

Proposition 1 (ratios identify the structure). *The ratios of marginal effects within a segment are free of both the density and the scale: the support-to-performance ratio equals $\gamma_g e$ and the effort-to-performance ratio equals $\gamma_g (1 - s)$. Consequently, for any baselines b_g , any scales σ_g , any e , s , and any Ψ , the ordering of these ratios across segments coincides with the ordering of γ_g .* The cancellation is the identified-ratio fact of discrete choice, where only ratios of coefficients are identified and a ratio of partial effects is free of the evaluation point (Train 2009; Wooldridge 2010). What the proposition adds is the reading: H9-type hypotheses are statements about those ratios, and the γ ordering they recover survives any baseline heterogeneity. It survives scale heterogeneity for the same reason, since $1/\sigma_g$ multiplies every effect in the segment and cancels with the density. The prescription this paper reaches is therefore robust to the confound the antecedent literature is about as well as to its own, which is the strongest form in which the result can be stated and the reason the two literatures are complements here. This is the structural content of H9, stated as an invariance: across capability segments the robust difference lies in the relative weight of cost-side drivers against benefit-side drivers. The prediction that effort and support dominate for No-skill adopters and performance for Innovative ones is a statement about these ratios, which are invariant to arbitrary differences in how often the segments adopt because the density factor cancels identically within any segment. This carries H9 directly into the companion study's cross-segment ratio contrasts.

Proposition 2 (levels confound structure with location and scale). For segments g, h with $\gamma_g > \gamma_h$, the level of the effort (or support) marginal effect is smaller for g than for h if and only if

$$\frac{\psi(\zeta_g)}{\psi(\zeta_h)} < \frac{\gamma_h/\sigma_h}{\gamma_g/\sigma_g}.$$

Level comparisons therefore reverse the capability ordering whenever baseline adoption rates or error scales differ enough. Under a common scale the condition reduces to $\psi(\zeta_g)/\psi(\zeta_h) < \gamma_h/\gamma_g$, and no reversal occurs when baselines are matched. The confound is not exotic. A No-skill segment far from adoption (baseline five percent, say) has a small density, so its measured effort effect can fall below that of an Innovative segment sitting near fifty percent despite a structural sensitivity four times larger. That a marginal effect is a density times a coefficient, and therefore depends on where the group sits on the response curve, is known and stated in prose in the health-econometrics literature (Norton & Dowd 2018). The contribution here is the exact condition on the density ratio under which the ordering reverses, its genericity, and the shares reported in Section 5.



Proposition 3 (matched baselines, and what they do not fix). (i) If segments share a baseline (ζ_g equal across g) and a common scale, all level orderings of the effort and support effects coincide with the γ ordering, and the support effect is increasing in γ (support and capability shortfall are complements). (ii) If segments share a baseline but the scales differ, the density factor is common and the level ordering follows γ_g/σ_g , so it reverses the γ ordering exactly when $\gamma_g/\sigma_g < \gamma_h/\sigma_h$ for some pair with $\gamma_g > \gamma_h$.

Part (ii) is the reason the ratio route of Proposition 1 is the recommendation and matching baselines is the fallback. Matching baselines removes the confound this paper identifies and leaves the confound the antecedent literature identifies untouched, and since σ_g is not identified from a single binary model, an analyst cannot check whether the remaining assumption holds. In the Monte Carlo, matched baselines with freely drawn scales still reverse the capability ordering on 44.2 percent of draws, against zero under a common scale. The ratios are correct on every draw in both cases.

Proposition 4 (standardized coefficients inherit the confound). Let the estimated multi-group model report standardized coefficients $\beta_{j,g}^{std} = b_{j,g} \sigma_{x_{j,g}}/\sigma_{y,g}$. Then: (i) groups with identical structural coefficients but different predictor variances have different standardized coefficients, so a cross-group standardized comparison detects variance heterogeneity as if it were structural difference; (ii) within a group, the ratio of two standardized coefficients equals $b_1\sigma_{x_1}/(b_2\sigma_{x_2})$, free of the outcome variance, so ratio contrasts are unaffected by outcome-variance heterogeneity and require only comparable predictor metrics, the condition metric invariance supplies; (iii) with a true structural ordering across groups, the standardized comparison can reverse it whenever variance ratios oppose the structural gap, and the unstan-

standardized comparison on a common metric cannot. Part (i) restates a point the standardization literature made decades ago (Kim & Mueller 1976; Kim & Ferree 1981) and that structural-equation texts carry as a norm (Bollen 1989); parts (ii) and (iii), the ratio identity and the quantified reversal condition, are what this paper adds. The mechanism is restriction of range wearing a structural costume: an Innovative segment in which effort is uniformly low has little effort variance, so its standardized effort path shrinks. A naive multi-group reading concludes that effort does not matter for the capable, when the capable do not vary in it. Proposition 2's marginal-effect confound operates at the SEM level whenever standardized coefficients are compared across groups, and Proposition 4 formalizes the condition under which standardized comparisons reverse a true structural ordering.

4. WHAT THIS PRESCRIBES FOR THE COMPANION STUDY

The companion's H9 should be tested on quantities that Propositions 1 and 4 protect. Concretely, the analysis plan needs four additions. First, alongside invariance-tested latent path comparisons, report probability-scale marginal effects and their within-segment ratios (effort to performance, facilitating conditions to performance), and state the segment-dependence conclusion on the ratios. Second, report each segment's baseline adoption rate next to any cross-segment comparison, and treat large baseline gaps as a warning that level comparisons are uninformative about structure. Third, where design permits, compare segments at matched baselines (subsample or covariate-adjusted), where Proposition 3(i) makes level comparisons meaningful again, treating this as a fallback rather than as an equivalent route: Proposition 3(ii) shows that matching restores the γ ordering only under a common error scale, an assumption the data cannot check. Fourth, compare unstandardized coefficients on invariance-tested metrics and never standardized coefficients across segments, whose cross-group differences mix structure with variance heterogeneity (Proposition 4). None of this replaces measurement invariance, which concerns the constructs (Steenkamp & Baumgartner 1998; Vandenberg & Lance 2000, and for composite models the MICOM procedure of Henseler, Ringle & Sarstedt 2016); it concerns what may be compared after invariance holds. The point binds partial-least-squares practice directly, since PLS works on standardized variables, so a PLS multi-group comparison of paths is a comparison of standardized coefficients by construction. The two are complementary: invariance establishes that the same construct is measured across groups, and the prescriptions above establish which comparisons of that construct are structurally informative.

5. VERIFICATION

The scripts, seeds, and result files that reproduce every number and figure are listed in Annex B.

- **Closed forms (numeric cross-check).** The marginal-effect formulas match central finite differences to within 2×10^{-10} on 2,000 random configurations, for both probit and logit links.
- **Reversal condition (numeric cross-check).** On every one of 200,000 random three-segment configurations per link, the condition of Proposition 2 predicts the observed level reversals exactly (the equivalence is checked pairwise, with no exceptions).
- **Monte Carlo (robustness).** Configurations draw γ ordered on $[0.2, 3]$, independent baselines b_g on $[-1, 3]$, e on $[0.3, 2]$, s on $[0, 0.8]$, at a common scale. Ratio orderings reproduce the γ ordering on 100 percent of draws for both links. Naive level comparisons misorder the segments on 62.8 percent of draws (probit) and 46.4 percent (logit). With baselines matched by construction, the reversal share is exactly zero on the same draws. Repeating the exercise with scales drawn freely leaves the level misordering at 61.7 percent (probit) and 49.7 percent (logit), so scale heterogeneity is a second confound layered on a trap that is already generic rather than a worsening of it.

- **Scale invariance (Prop. 1, numeric cross-check).** On 200,000 configurations per link with segment scales drawn independently on $[0.5, 2]$, the within-segment ratios equal $\gamma_g e$ and $\gamma_g(1 - s)$ to 2×10^{-15} , and their ordering reproduces the γ ordering on 100 percent of draws. The levels move on 99.99 percent of the same draws, which is the asymmetry the prescription rests on.
- **Matched baselines under free scale (Prop. 3, robustness).** On 200,000 configurations per link, matching baselines while drawing scales freely leaves the capability ordering reversed on 44.2 percent of draws. Imposing a common scale on the same draws returns a reversal share of exactly zero, and the generalized reversal condition of Proposition 2 predicts the observations exactly, pairwise, with no exceptions. The share is identical across links because at a matched baseline the density factor is common and the comparison reduces to the ordering of γ_g/σ_g .
- **Standardized confound (Prop. 4, cross-check and robustness).** Population algebra on 200,000 two-group, two-predictor configurations: with identical structural coefficients and independent group variance draws, the cross-group standardized gap is nonzero on 99.999 percent of draws; the within-group ratio identity of Proposition 4(ii) holds to 4×10^{-15} , and with a true structural ordering imposed, the standardized comparison reverses it on 16.7 percent of draws while the unstandardized comparison reverses it on none.

The reversal shares are properties of the sampled configuration space, determined by the joint distribution of parameter draws, and they indicate how much of the configuration space the confound occupies. They are no estimate of how often real studies err.

6. CONCLUSION

Grounding segment-dependent acceptance in a cost primitive does two things at once. It derives the companion's central hypothesis rather than positing it: capability is a unit implementation cost, so the cost-side drivers must matter more, relative to the benefit-side drivers, for the less capable. And it disciplines the test: that "more" lives in density-free ratios, while the level comparisons that multi-group practice reaches for first confound structure with location and disorder the segments across much of the parameter space. The empirical companion thus receives both its hypothesis from the cost primitive and its discipline from the identification analysis, together with the practical prescription of Section 4.

APPENDIX: PROOFS

Proposition 1. Within segment g , all three marginal effects share the factor $\psi(\zeta_g)/\sigma_g > 0$: dividing, the support-to-performance ratio is $\gamma_g e$ and the effort-to-performance ratio is $\gamma_g(1 - s)$, both independent of b_g , σ_g and Ψ . The scale cancels for the same reason the density does, being a factor common to the segment's three effects. With $e > 0$ and $1 - s > 0$ common across segments, both ratios are strictly increasing in γ_g , so their cross-segment ordering is the γ ordering. $\square$

Proposition 2. For the effort effect, $\psi(\zeta_g)\gamma_g(1 - s)/\sigma_g < \psi(\zeta_h)\gamma_h(1 - s)/\sigma_h$ iff $\psi(\zeta_g)/\psi(\zeta_h) < (\gamma_h/\sigma_h)/(\gamma_g/\sigma_g)$, and the same algebra applies to the support effect. Under a common scale the right side is below one (since $\gamma_g > \gamma_h$), so matched baselines ($\zeta_g = \zeta_h$, left side equal to one) preclude reversal. Since ψ is continuous, positive at the center, and vanishes in the tails, the left side ranges over $(0, \infty)$ as b_g varies with b_h fixed, so configurations on both sides of the threshold exist: reversals occur exactly when baselines are sufficiently far apart in the sense of the density ratio, or, at any baselines, when the scale ratio is adverse enough. $\square$

Proposition 3. (i) At common ζ and common σ , the effort effect $\psi(\zeta)\gamma_g(1 - s)/\sigma$ and support effect $\psi(\zeta)\gamma_g e/\sigma$ are strictly increasing in γ_g , which yields both the ordering and,

differentiating the support effect in γ at fixed ζ , the complementarity $\psi(\zeta)e/\sigma > 0$. (ii) At common ζ with σ_g free, the density factor $\psi(\zeta)$ is common to all segments and cancels from every pairwise comparison, so the level ordering of both effects is the ordering of γ_g/σ_g . That differs from the ordering of γ_g on a set of positive measure, since σ is drawn independently of γ and $\gamma_g/\sigma_g < \gamma_h/\sigma_h$ with $\gamma_g > \gamma_h$ whenever $\sigma_g/\sigma_h > \gamma_g/\gamma_h > 1$. $\square$

Proposition 4. (i) With $b_{j,1} = b_{j,2} = b_j$ and $\sigma_{x_{j,1}} \neq \sigma_{x_{j,2}}$, the standardized coefficients $b_j\sigma_{x_{j,g}}/\sigma_{y,g}$ differ across groups unless the predictor-variance change is exactly offset by the induced outcome-variance change, a knife-edge of measure zero under continuous variance distributions. (ii) Within group g , $\beta_{1,g}^{std}/\beta_{2,g}^{std} = (b_1\sigma_{x_{1,g}})/(b_2\sigma_{x_{2,g}})$ by cancellation of $\sigma_{y,g}$. (iii) With $b_{1,2} > b_{1,1}$, the standardized comparison satisfies $\beta_{1,2}^{std} < \beta_{1,1}^{std}$ if and only if $(\sigma_{x_{1,2}}/\sigma_{y,2})/(\sigma_{x_{1,1}}/\sigma_{y,1}) < b_{1,1}/b_{1,2}$, which has positive measure; the unstandardized comparison on a common metric preserves $b_{1,2} > b_{1,1}$ by construction. $\square$

REFERENCES

- Ai, C., & Norton, E. C. (2003). Interaction terms in logit and probit models. *Economics Letters*, 80(1), 123–129.
- Allison, P. D. (1999). Comparing logit and probit coefficients across groups. *Sociological Methods & Research*, 28(2), 186–208.
- Bollen, K. A. (1989). *Structural Equations with Latent Variables*. Wiley.
- Breen, R., Karlson, K. B., & Holm, A. (2018). Interpreting and understanding logits, probits, and other nonlinear probability models. *Annual Review of Sociology*, 44, 39–54.
- Ethiraj, S. K., Kale, P., Krishnan, M. S., & Singh, J. V. (2005). Where do capabilities come from and how do they matter? A study in the software services industry. *Strategic Management Journal*, 26(1), 25–45.
- Henseler, J., Ringle, C. M., & Sarstedt, M. (2016). Testing measurement invariance of composites using partial least squares. *International Marketing Review*, 33(3), 405–431.
- Hoetker, G. (2007). The use of logit and probit models in strategic management research: Critical issues. *Strategic Management Journal*, 28(4), 331–343.
- Jullien, N., Viseur, R., & Zimmermann, J.-B. (2025). A theory of FLOSS projects and Open Source business models dynamics. *Journal of Systems and Software*, 224, 112383. <https://doi.org/10.1016/j.jss.2025.112383>
- Jullien, N., & Zimmermann, J.-B. (2011). FLOSS firms, users and communities: A viable match? *Journal of Innovation Economics*, 1(7), 31–53.
- Kogut, B. M., & Metiu, A. (2001). Open source software development and distributed innovation. *Oxford Review of Economic Policy*, 17(2), 248–264.
- Karlson, K. B., Holm, A., & Breen, R. (2012). Comparing regression coefficients between same-sample nested models using logit and probit: A new method. *Sociological Methodology*, 42(1), 286–313.
- Kim, J.-O., & Ferree, G. D. (1981). Standardization in causal analysis. *Sociological Methods & Research*, 10(2), 187–210.
- Kim, J.-O., & Mueller, C. W. (1976). Standardized and unstandardized coefficients in causal analysis: An expository note. *Sociological Methods & Research*, 4(4), 423–438.
- Kuha, J., & Mills, C. (2020). On group comparisons with logistic regression models. *Sociological Methods & Research*, 49(2), 498–525.
- Long, J. S., & Mustillo, S. A. (2021). Using predictions and marginal effects to compare groups in regression models for binary outcomes. *Sociological Methods & Research*, 50(3), 1284–1320.
- Mood, C. (2010). Logistic regression: Why we cannot do what we think we can do, and what we can do about it. *European Sociological Review*, 26(1), 67–82.
- Norton, E. C., & Dowd, B. E. (2018). Log odds and the interpretation of logit models. *Health Services Research*, 53(2), 859–878.

- Qureshi, I., & Compeau, D. (2009). Assessing between-group differences in information systems research: A comparison of covariance- and component-based SEM. *MIS Quarterly*, 33(1), 197–214.
- Steenkamp, J.-B. E. M., & Baumgartner, H. (1998). Assessing measurement invariance in cross-national consumer research. *Journal of Consumer Research*, 25(1), 78–90.
- Swait, J., & Louviere, J. (1993). The role of the scale parameter in the estimation and comparison of multinomial logit models. *Journal of Marketing Research*, 30(3), 305–314.
- Train, K. (2009). *Discrete Choice Methods with Simulation* (2nd ed.). Cambridge University Press.
- Vandenberg, R. J., & Lance, C. E. (2000). A review and synthesis of the measurement invariance literature. *Organizational Research Methods*, 3(1), 4–70.
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User Acceptance of Information Technology: Toward a Unified View. *MIS Quarterly*, 27(3), 425–478. <https://doi.org/10.2307/30036540>
- von Hippel, E. (2001). Innovation by user communities: Learning from open-source software. *MIT Sloan Management Review*, 42(4), 82–86.
- Williams, R. (2009). Using heterogeneous choice models to compare logit and probit coefficients across groups. *Sociological Methods & Research*, 37(4), 531–559.
- Wooldridge, J. M. (2010). *Econometric Analysis of Cross Section and Panel Data* (2nd ed.). MIT Press.

ANNEX A: WORKING NOTES (TEMPORARY ANNEX)

This annex lists internal working documents and will be removed from the final version.

- `../notes/common/theory-foundation.md`: the shared theoretical foundation of the project and its re-centering decisions.
- `../notes/common/model-reconciliation.md`: the hypothesis-cutting history behind the companion study’s H9.
- `../notes/stratified-adoption/instrument.md` and `../notes/stratified-adoption/prereg-paper1.md`: the companion study’s instrument and pre-registration draft, which adopt Section 4’s prescriptions.
- `../notes/common/research-program.md`: the wider research program.

ANNEX B: REPRODUCTION: SCRIPTS, DATA, AND FIGURES

- `../verify/acceptance_structure.py` (seed 20260716): verifies every quantitative statement in Section 5 and generates the figure `figures/acceptance-structure-effects.png`; results are written to `../verify/results/acceptance_structure.json`. Each check draws from its own generator (offsets 7, 11, 13, 17, 19 on the seed), so inserting the two heteroskedastic checks left the earlier numbers bit-identical. The heteroskedastic pair is `check_scale_invariance` (Proposition 1 under free σ_g) and `monte_carlo_scale` (Propositions 2 and 3 under free σ_g , including the matched-baseline failure share and the common-scale control that returns zero).
- `../verify/run_all.py`: reproduces every number and figure of the project’s theory papers in one command.