

Adopting Open Source in 2026: A UTAUT Model Grounded in the FLOSS User Typology

Stefane Fermigier (Abilian) · sf@abilian.com

Draft working paper v0.4, 2026-07-16. Conceptual and research-design: no survey data; every hypothesis is specified and none is estimated.

ABSTRACT

The adoption of open source has been studied with the Unified Theory of Acceptance and Use of Technology (UTAUT), but usually for a representative user, whereas the economics of open source treats adopters as a typology. This paper adapts UTAUT to open-source adoption and grounds its segmentation in the value-network user typology of Jullien, Viseur & Zimmermann (2025). Adopters differ by capability (Innovative users can translate needs into requirements and write code, Frontier users can specify but not implement, No-skill users can do neither) and by need (Control, Lock-out, Adaptable, Low-cost). We keep the four core UTAUT constructs, add a small set of three open-source-specific constructs (community support, transparency and auditability, and strategic alignment carrying sovereignty), and add the two 2026 institutional moderators: regulatory perceived risk (the Cyber Resilience Act and AI Act, mitigated by commercial assurance) and sovereignty. The central design claim is that the driver weights differ across capability segments, so that the acceptance structure varies across segments. The estimation is a multi-group structural model with paired dyadic interviews. Estimation is future work; the contribution of this draft is the model and its segmentation. Whether AI assistance shifts an adopter's capability is a separate question, out of scope here.

1. INTRODUCTION

Open-source adoption sits at the junction of three literatures, and each supplies one piece of this study. Technology-acceptance research supplies the measurement apparatus. UTAUT (Venkatesh et al. 2003) consolidated eight earlier models, from the theory of reasoned action and planned behavior (Ajzen & Fishbein 1970; Ajzen 1991) through the technology acceptance model (Davis 1989) and its hybrids and competitors (Thompson et al. 1991; Taylor & Todd 1995), into four core constructs: performance expectancy, effort expectancy, social influence, and facilitating conditions, with demographic moderators. Its consumer extension added price value, hedonic motivation, and habit (Venkatesh, Thong & Xu 2012), and two decades of applications and meta-analyses have established its reach and catalogued its strains (Dwivedi et al. 2011; Dwivedi et al. 2019; Blut et al. 2022). The lineage lacks an account of who the adopter is: its moderators are demographic, and heterogeneity enters as variance to be absorbed. When the adopter is an organization whose relationship to source code ranges from writing it to being unable to read it, a demographic moderator set has no purchase on the differences that decide the adoption.

Empirical studies of open-source adoption supply the domain content. Organizational surveys, case studies, and reviews document drivers with no analogue in proprietary acceptance: community responsiveness, source access and auditability, exposure to a project's governance, and total cost of ownership under partial self-support (Moon & Baker 2009; Lundell & Gamalielsson 2011; Shaikh & Cornford 2011; Morgan & Finnegan 2014; Yaseen et al. 2019). The component-selection literature shows practitioners weighing community activity, maturity, and documentation when choosing among candidate components (Li et al. 2022; Kuwata et al. 2014), and the economic-impact literature attaches productivity and growth effects to the adoption these micro-decisions aggregate into (Nagle 2017; Blind et al. 2021; Blind & Schubert 2023). The literature is rich in factors and thin in structure: factor

lists vary study by study, samples are conveniences, and where an acceptance framework is used at all, the adopter is a representative user or an ad-hoc demographic split.

The economics of open source supplies the structure. From its founding papers it has treated participants and adopters as heterogeneous by construction: developers signal ability and collect career returns (Lerner & Tirole 2002), firms mix private investment with collective provision (Bonaccorsi & Rossi 2003; von Hippel & von Krogh 2003), and users differ in what they can do with code, from the innovative users of von Hippel (2001) to buyers of sealed products. The theory of FLOSS projects and open-source business-model dynamics (Jullien, Viseur & Zimmermann 2025), which reads a project as a value network in the sense of Morgan, Feller & Finnegan (2013), makes that heterogeneity operational. Adopters are classified by two capacities (translating needs into requirements, developing code) into Innovative, Frontier, and No-skill users, and by need into Control, Lock-out, Adaptable, and Low-cost positions. The typology predicts which services each segment must buy from the market because it cannot self-provide them. This literature lacks measurement: no instrument estimates how the adoption calculus actually differs across those segments.

The 2026 context adds two institutional forces that the earlier acceptance literature did not price. Regulatory liability under the Cyber Resilience Act and the AI Act raises the perceived risk of adopting and shipping open source and creates demand for commercial assurance (APELL 2021; CNLL & Inno³ 2024). Digital sovereignty, pressed by European supervisors and Commission policy (ACM et al. 2026; European Commission 2021a), enters the need for control and lock-in avoidance, most visibly for public buyers. Both enter the model as moderators, and nothing more: a third candidate force, AI assistance, is a strong and separable claim about capability itself, and whether it shifts an adopter's capability is a separate question, out of scope here. This study is the demand-side half of a pair: the supply-side companion ([./././oss-business-models/](https://oss-business-models/)) reads the same value network from the firm's side, and the two share the typology and the 3A-services vocabulary.

This paper joins the pieces. It adapts UTAUT to organizational open-source adoption with a small construct set, replaces the representative user with the value-network typology as the grouping variable, and adds the two 2026 institutional moderators. Its central claim is that the acceptance structure is segment-dependent (H9). Its test of that claim is disciplined by a companion theory paper ([acceptance-structure-wp.md](#)) that derives the hypothesis from a cost primitive and shows which comparisons identify segment structure and which confound it with baseline adoption levels.

2. RELATED WORK

2.1 The UTAUT lineage and its critique

UTAUT was built as a synthesis: Venkatesh et al. (2003) took the theory of reasoned action (Ajzen & Fishbein 1970), the technology acceptance model (Davis 1989), the motivational model (Davis, Bagozzi & Warshaw 1992), the theory of planned behavior (Ajzen 1991), the combined acceptance and planned-behavior model (Taylor & Todd 1995), the model of PC utilization (Thompson et al. 1991), innovation diffusion theory (Rogers 1995), and social cognitive theory (Bandura 1986), and distilled them into four determinants of behavioral intention and use. UTAUT2 reoriented the model to consumer contexts with price value, hedonic motivation, and habit (Venkatesh, Thong & Xu 2012); TAM3 mapped the antecedent side (Venkatesh & Bala 2008). The framework's portability is its selling point: applications run from ERP training (Keong et al. 2012; Chauhan & Jaiswal 2016) and government ICT (Gupta et al. 2008) through telehealth (Cimperman et al. 2016), mobile learning (Wang, Wu & Wang 2009; Chao 2019), blockchain (Li 2020), and conversational AI (Menon & Shilpa 2023).

The critique is almost as old as the model. Bagozzi (2007) attacked the paradigm itself: a proliferating family of constructs and moderators chasing explained variance, a fragile inten-

tion-behavior link, and a decision-maker with no deliberative interior. Van Raaij & Schepers (2008) showed the unified model's high explained variance depends on stacked moderators and questioned its construct definitions. The meta-analytic record is sobering: Dwivedi et al. (2011) found the moderators rarely used and inconsistently effective; Dwivedi et al. (2019) re-estimated a pared-down model with attitude reinstated; Blut et al. (2022) challenged the model's validity wholesale and catalogued which extensions replicate. This paper draws two design rules from that record. First, parsimony: every added construct must carry a distinct economic channel and a clear prior (Section 3.1). Second, heterogeneity should be theorized before it is estimated. The lineage's own evidence points that way, from acceptance structures that shift with lead-user traits (Mütterlein et al. 2019) to social influence that depends on the referent group (Eckhardt et al. 2009) and on network position (Sykes, Venkatesh & Gosain 2009), yet it stops short of an economic account of who differs and why. The value-network typology supplies one.

2.2 Open-source adoption and selection empirics

Open-source adoption has been surveyed and case-studied since the field professionalized (Fitzgerald 2006). Reviews and organizational studies converge on a driver set that mixes standard IT-adoption factors with open-source specifics: cost and total cost of ownership, quality and flexibility, vendor independence, community support, and skills availability (Moon & Baker 2009; Lundell & Gamalielsson 2011; Daffara 2007; Yaseen et al. 2019). Strategy-flavored studies read adoption as value-network participation: what an organization gains depends on what it can appropriate from the network around a project (West 2003; Morgan, Feller & Finnegan 2013; Morgan & Finnegan 2014). Cost studies warn that the total cost of ownership depends on who operates what: self-support is cheap only for those equipped to perform it (Shaikh & Cornford 2011).

At the component level, selection empirics are the sharpest evidence on what adopters actually weigh. Li et al. (2022) catalogue the factors and metrics practitioners use to select components for integration: community activity, maintenance signals, documentation, license compatibility. Kuwata, Takeda & Miura (2014) formalize community maturity as a quality estimator, and the security literature grounds the auditability argument (Hoepman & Jacobs 2007). Ethiraj et al. (2005), studying software-services firms, established the capability side: client-specific and project-management capabilities are acquired unevenly and matter for performance, which is the microfoundation the typology builds on. Macro studies show the aggregate stakes of these decisions (Nagle 2017; Blind et al. 2021; Blind & Schubert 2023), and the mirror-image supply literature documents how firms position themselves commercially around the projects adopters rely on (Li et al. 2024). None of these studies estimates a validated acceptance model over an economically grounded segmentation, which is the slot this study fills.

2.3 The value-network typology and why it supplies the segmentation

Jullien, Viseur & Zimmermann (2025) model FLOSS projects and the business models around them as a value network (Morgan, Feller & Finnegan 2013) organized by two distinctions. On the demand side, a software need is characterized by how much adaptation it requires and how much it must evolve. Weak on both axes yields Low-cost needs (commodity acquisition at the best price); strong adaptability alone yields Adaptable needs (personalized functionality at the best cost); strong upgradability alone yields Lock-out needs (standard functionality that must persist through technology change without capture); strength on both yields Control needs (personalized functionality the organization must keep adaptable over time). On the capability side, two capacities inherited from Ethiraj et al. (2005), translating needs into technological requirements and developing them in code, sort users into Innovative (both capacities), Frontier (specification without development), and No-skill (neither) (Jullien & Zimmermann 2011; Viseur & Jullien 2023; the Innovative label follows von Hippel 2001, the Frontier label Kogut & Metiu 2001).

The typology earns its place as this study’s segmentation on four grounds. It is economic: the capacities are costs, so it predicts which drivers should bind where, and the companion paper turns that prediction into a formal result. It is generative: the theory derives from it which of the 3A services (assistance with needs specification, adaptation, and assurance; Jullien & Zimmermann 2009) each segment must buy, so the demand-side segments connect to supply-side business models. It is external: the segmentation predates this study and rests on two decades of prior work (Jullien & Zimmermann 2009, 2011; Jullien & Roudaut 2012; Viseur & Jullien 2023), so the segments arrive as hypotheses imported from theory. And it is measurable: the two capacities admit direct, behavior-anchored classification items (Section 5.2), where role- or attitude-based segmentations would need the very constructs the model estimates.

3. FRAMEWORK

3.1 UTAUT and the parsimony constraint

UTAUT holds that performance expectancy, effort expectancy, social influence, and facilitating conditions drive behavioral intention, which drives use. It is the right base for open-source adoption: its constructs are validated, its instruments are portable, and its perceptual framing fits a decision made under uncertainty about quality and cost. Its known failure mode is construct proliferation and inflated explained variance when moderators are stacked (Bagozzi 2007; van Raaij & Schepers 2008; Blut et al. 2022), so this paper adds only constructs with a clear prior and folds the rest into the core. The folding rules are explicit. Perceived code quality and customizability are performance expectancy in an organizational referent: they describe how well the adopted software serves the work (items PE1 to PE4, with PE4 carrying customization). Contribution infrastructure, and an open-source program office where one exists, are facilitating conditions: organizational machinery that lowers the cost of acting (FC1 to FC4). Only three additions remain as constructs, each with an economic channel the core cannot carry (Section 3.2), and two institutional moderators enter with their mitigators (Section 4). The full cut list, mapping the source drafts’ 23 hypotheses onto the pruned set, is in the model-reconciliation note (Annex A).

3.2 The three open-source constructs

- **Community support:** the responsiveness and usefulness of a project’s community when an adopter needs help. Open source unbundles the software from the services around it: the license makes the stock free, while the flow of fixes, upgrades, and answers is produced by the community and by vendors selling complements. Community support is the open-source form of a facilitating condition, and it is repeatedly decisive in component selection (Li et al. 2022; Kuwata et al. 2014). Measured by CS1 to CS4.
- **Transparency and auditability:** source access and supply-chain verifiability (SBOM, security posture, governance visibility). This is the construct through which the security and compliance value of open source enters: inspection converts vendor claims into checkable evidence (Hoepman & Jacobs 2007), and the CRA turns transparency into an operational property with compliance value (CNLL & Inno³ 2024). Measured by TA1 to TA4.
- **Strategic and ideological alignment:** the extent to which adopting open source aligns with the organization’s values and its independence from single vendors and non-EU jurisdictions. This carries digital sovereignty; the older individual-level “ideological alignment” is its special case. Measured by SA1 to SA4.

Three further measured variables support the moderator hypotheses without expanding the structural core: regulatory perceived risk (RR1 to RR4), commercial assurance (CA1 to CA3), and price value (PV1 to PV3, the UTAUT2 construct retained as H8’s comparison). The

outcome side pairs behavioral intention (BI1 to BI3) with a formative use-behavior block (UB1 to UB3: breadth, depth, and policy) that anchors stated intention in current practice.

3.3 The value-network typology as segmentation

The segmentation is Jullien, Viseur & Zimmermann's user typology, itself resting on Ethiraj et al. (2005) and Jullien & Zimmermann (2011). Adopters are classified by two capacities, specification (translate a need into a requirement) and development (adapt and fix code), into Innovative, Frontier, and No-skill users, and by need into Control, Lock-out, Adaptable, and Low-cost. The claim the model tests is that the acceptance structure differs by capability segment, because the constructs that bind (effort and facilitating conditions for the lower-capability segments, performance and strategic alignment for the Innovative) are not the same across them.

The gradient the typology formalizes appears in official European statistics mainly as a gap in what they record. The Commission's study of critical digital capacities beyond 2027 reproduces Eurostat indicators of enterprise technology adoption sector by sector. For artificial intelligence it reports use alongside development by a firm's own employees (European Commission 2026, Table 7): 28 percent of chemicals firms use AI against 3 percent developing it in house, 24 against 4 percent in automotive, 26 against 3 percent in electrical equipment, 12 against 1 percent in food and beverage, in textiles, and in tourism alike (use figures 2025, own-development figures 2024). Across the ten sectors the table covers, in-house development runs an order of magnitude below use everywhere except publishing and motion picture, where 56 percent use AI and 12 percent develop it. The cloud row of the same table records the buying side only, defined as "enterprises buying sophisticated cloud computing services, 2025" and running from 27 percent in tourism to 71 percent in pharmaceuticals, with no self-provisioning counterpart anywhere in the instrument. Sectoral adoption statistics of this kind therefore measure acquisition, and for one technology a fraction of production, which leaves most of the capability distinction this paper's segmentation rests on outside their field of view.

The predictive logic runs through self-provision. Whatever a segment can produce in-house is cheap for it, so perceptions of the corresponding external supports carry little decision weight; whatever it cannot produce must come from the environment (community, vendor, internal support machinery), so the constructs measuring that environment carry the decision. The theory says the same thing from the supply side: open-source companies exist to sell the 3A services that bridge a buyer's need and its capability, so the services a segment must buy identify the constructs that should bind in its acceptance structure. Table 1 states the predictions cell by cell, with the economic mechanism in one clause each. Two disciplines attach to it. The need and capability dimensions correlate by the theory's own logic (Control needs are pursued mainly by organizations able to exercise control), so the crossing is measured, never assumed, with the need position derived from the TA and SA batteries and the branched items as a secondary cut. And the table is a prediction about relative weights within a segment's calculus: every path is estimated in every segment, and the aggregate capability contrast is what H9 registers.

Capability	Need	Constructs predicted to bind	Economic mechanism
Innovative	Control	PE, SA, TA	the firm can exercise adaptability in-house, so the option value of control and self-verification is realizable
Innovative	Lock-out	SA, TA	exit is cheap for a firm that can port and fix, so lock-in avoidance is a live strategic benefit
Innovative	Adaptable	PE	customization gains dominate where implementation is nearly free

Capability	Need	Constructs predicted to bind	Economic mechanism
Innovative	Low-cost	PV, PE	commodity acquisition at best cost; implementation cost is minor either way
Frontier	Control	SA, CS, CA	control is desired but its exercise must be bought as adaptation services from community or vendor
Frontier	Lock-out	SA, TA, CS	standards and transparency substitute for the development capacity the firm lacks
Frontier	Adaptable	CS, FC, PE	adaptation is outsourced or community-mediated; specification capacity lets the firm evaluate what it buys
Frontier	Low-cost	PV, EE	standard functionality at best cost, with effort pricing the residual integration
No-skill	Control	CA, FC, RR	control is reachable only through an assuring vendor, and regulatory duties fall on bought assurance
No-skill	Lock-out	CA, SI, TA	the firm cannot verify openness claims itself, so it relies on assurance and on what respected peers run
No-skill	Adaptable	FC, CA, EE	every adaptation is a purchased service; available help and ease decide feasibility
No-skill	Low-cost	EE, PV, SI	off-the-shelf acquisition; ease of use, price, and peer defaults decide

Table 1. Predicted binding constructs by capability segment and need position. Construct abbreviations as in Table 2.

Read along the capability rows, the table collapses into H9: the cost-side constructs (EE, FC, and their market substitutes CS and CA) migrate into the binding set as capability falls, while the benefit-side constructs (PE, SA) dominate at the top. Read along the need columns, it locates the moderators: SA binds wherever upgradability matters (Control, Lock-out), the entry point of sovereignty, and RR binds where the adopter cannot self-audit, the site of assurance.

3.4 What the constructs measure: perceptions, agency, fashion

Technology decisions are not all made in the organization's interest, and few are made by a textbook rational actor. Technical buyers choose stacks partly for what they add to a CV (résumé-driven development, Fritzsche et al. 2021), and in open source this careerism has a canonical, partly aligned form: visible contribution is portable reputation (Lerner & Tirole 2002), which firms harvest as employer branding, while hyperscaler certification ecosystems push the same career capital toward incumbent platforms. Decision-makers also buy defensibility (a vendor to hold accountable, per the career-concerns logic of Holmström 1999), follow fashion and cascades (Abrahamson 1991; Bikhchandani et al. 1992), and adopt for legitimacy (DiMaggio & Powell 1983). UTAUT accommodates this because it is a perceptual model: these motives reach the constructs through the respondent's expectancies, and social influence is the individual-level shadow of the institutional pressures.

The accommodation has a price, and this study pays it in two ways. Interpretation: path weights describe the decision calculus of the reporting role, agency and legitimacy terms included. Claims about organizational economics are therefore made where the dyadic design corroborates them, since motive divergence lives between the proposer's and the authorizer's accounts of the same decision. Design: a small set of auxiliary items outside the structural model (an indirect CV-motive item, a defensibility item, a behavioral classifier anchor) serves as robustness covariates for H9 and as a validity check on the capability classifier, without

reinflating the construct set. The taxonomy and the item proposals are in the behavioral-agency note (Annex A).

3.5 Constructs, measurement, and the hypotheses they feed

Table 2 summarizes the measurement side. All battery identifiers refer to the English master questionnaire (Annex A), where every item is worded, tagged with its provenance, and budgeted; all Likert batteries are 7-point with a stated referent and a “cannot say” opt-out.

Construct	Definition	Measurement source	Mode	Feeds
Performance expectancy (PE)	perceived improvement in the organization’s work, quality and customizability included	PE1–PE4, adapted from Venkatesh et al. (2003)	reflective	H1; denominator of the H9 ratio contrasts
Effort expectancy (EE)	perceived ease of learning and operating the adopted open source	EE1–EE4, adapted from Venkatesh et al. (2003)	reflective	H2; H9
Social influence (SI)	perceived expectation of valued peers and communities that the organization use open source	SI1–SI3, adapted from Venkatesh et al. (2003)	reflective	H3
Facilitating conditions (FC)	perceived resources, skills, compatibility, and available help	FC1–FC4, adapted from Venkatesh et al. (2003)	reflective	H4; H7 mitigation; H9
Community support (CS)	responsiveness and usefulness of the adopted projects’ communities	CS1–CS4, written for this study	reflective	H5
Transparency and auditability (TA)	source access and supply-chain verifiability (SBOM, governance)	TA1–TA4, written for this study	reflective	H6
Strategic and sovereignty alignment (SA)	alignment with values and independence from single vendors and non-EU jurisdictions	SA1–SA4, written for this study	reflective	H8
Regulatory perceived risk (RR)	perceived risk from CRA and AI Act duties	RR1–RR4, written for this study	reflective	H7
Commercial assurance (CA)	availability and weight of commercial backing for adopted components	CA1–CA3, written for this study	reflective	H7 mitigation
Price value (PV)	perceived total-cost-of-ownership value	PV1–PV3, adapted from Venkatesh, Thong & Xu (2012)	reflective	H8 comparison
Behavioral intention (BI)	intention to adopt more open source	BI1–BI3, adapted from Venkatesh et al. (2003)	reflective	outcome of H1–H9

Construct	Definition	Measurement source	Mode	Feeds
	over the next twelve months			
Use behavior (UB)	current breadth, depth, and policy stance of open-source use	UB1–UB3, written for this study	formative	intention-behavior check

Table 2. Constructs, definitions, measurement, and the hypotheses each feeds.

The capability classifier (C1a, C1b, C2a, C2b) and its behavioral anchor AUX3 sit outside the structural model as the grouping variable and its validation (Section 5.2). AUX1 and AUX2, the agency and defensibility items of Section 3.4, are registered as robustness covariates only and never enter as constructs.

4. THE MODEL AND HYPOTHESES

The pruned hypothesis set follows (full derivation and cut list in the model-reconciliation note, Annex A). Each block carries its mechanism: the argument for the sign comes from cost, information, risk, or agency, and the citations locate the mechanism.

Core UTAUT.

In the random-utility reading the companion paper makes exact, an organization adopts when perceived benefit exceeds perceived cost plus a stochastic term, and the four core constructs are the perceptual shadows of that inequality’s parts. Performance expectancy proxies the gross benefit: productivity, quality, and customization gains from the adopted software and its evolution. Effort expectancy proxies the implementation cost: what the technology demands times the price at which the organization converts effort into working systems. Facilitating conditions proxy the support environment that attenuates the cost term: skills, compatible systems, contracted help. Social influence carries information and legitimacy: under quality uncertainty, peer adoption is informative (the cascade logic of Bikhchandani et al. 1992), and field expectations confer legitimacy on the choice (DiMaggio & Powell 1983), so the construct raises expected benefit and lowers perceived risk through channels the other three do not carry. All four signs are positive; the segmentation claim (H9) is about their relative weights.

- H1. Performance expectancy raises adoption intention.
- H2. Effort expectancy raises adoption intention.
- H3. Social influence raises adoption intention.
- H4. Facilitating conditions raise adoption intention.

Open-source constructs.

Community support is priced as assistance the adopter did not buy. An active, responsive community lowers the expected cost of the incidents and questions that adoption will generate, shrinks the variance around resolution times, and its visible activity doubles as a signal of the project’s sustainability (Kuwata et al. 2014). Selection studies find community responsiveness among the decisive criteria when practitioners choose components (Li et al. 2022). The sign is positive, and by the self-provision logic of Table 1 the construct should carry most weight where adopters consume assistance they cannot produce.

Transparency and auditability price verification. Under closed source, security and quality are credence attributes: the buyer bears an information asymmetry it can reduce only by trusting the vendor’s claims. Source access and supply-chain evidence convert part of that asymmetry into verifiable fact at a cost the adopter controls (Hoepman & Jacobs 2007). Regulation raises the value of the conversion: liability regimes demand evidence, and auditability lowers the cost of producing it (CNLL & Inno³ 2024). Since verification is worth

most when the price of being wrong is highest, the construct should load more strongly where security stakes and regulatory exposure are high, which is the interaction H6 registers.

- H5. Community support raises adoption intention.
- H6. Transparency and auditability raise adoption intention, more strongly for security-conscious and regulated adopters.

2026 institutional moderators.

Regulatory perceived risk is an expected-cost and ambiguity term. The CRA and the AI Act attach duties and liability to placing software on the market, and for open-source components the allocation of those duties across producers, integrators, and stewards has been the object of prolonged uncertainty in the sector (APELL 2021). Unresolved allocation reads as ambiguity, which decision-makers price above measurable risk of the same magnitude. The two mitigators are the market response and the internal response. Commercial assurance is the market response: a counterparty that absorbs part of the compliance burden and, per the defensibility channel of Section 3.4, gives the decision-maker an accountable party (Holmström 1999). Facilitating conditions are the internal response: compliance work is largely a fixed cost, cheaper per adoption for organizations with the skills and machinery already in place. The construct measures perceived risk, and the hypothesis says nothing about whether the regulation's actual burden or protection matches the perception.

Sovereignty is an option value that price value does not carry. A total-cost comparison prices the current contract; dependence on a single vendor or a non-EU jurisdiction is exposure to future price rises, unilateral term changes, and extraterritorial legal reach, none of which appear in this year's invoice. Open source lowers exit costs and widens switching options, so the strategic-alignment construct captures a hedging benefit realized later. For European public buyers an institutional mandate now sits on top of it: supervisory authorities call for reduced dependence and collective demand (ACM et al. 2026), and Commission policy makes openness a stated objective (European Commission 2021b). Representative benchmarking shows sovereignty attitudes are strong while stated willingness to pay is thin (Bitkom 2026), which is why H8 is a relative-weight claim about drivers of intention within respondents and no willingness-to-pay claim is made.

- H7. Regulatory perceived risk (CRA, AI Act) lowers adoption intention, mitigated by commercial assurance and facilitating conditions.
- H8. Strategic and sovereignty alignment predicts adoption intention over price value for European public-sector buyers.

Segment structure (the contribution).

The companion paper derives H9 from a cost primitive. A segment's capability is the unit cost at which it converts adoption effort into working systems. That unit cost scales the effort and support terms of the adoption decision, so the marginal effects of the cost-side drivers rise as capability falls, and the theory's ordering of segments becomes an ordering of driver-weight ratios (its Proposition 1). For the Innovative segment the cost side is cheap, so the decision margin moves with the benefit terms: performance, and the strategic value of control and lock-in avoidance that only adopters able to adapt and audit code can fully exercise (Table 1). The companion also disciplines the test: ratios of driver weights identify the capability ordering under any baseline heterogeneity, while level comparisons confound structure with each segment's location on the adoption curve, misordering the segments on 62.8 percent (probit) and 46.4 percent (logit) of sampled configurations in its Monte Carlo (Section 5.4).

- H9. The driver weights differ across capability segments: effort expectancy and facilitating conditions bind hardest for the No-skill and Frontier segments, while performance expectancy and strategic alignment dominate for the Innovative segment.

Hypothesis	Sign	Segment prediction	Test
H1 (PE)	positive	weight largest for Innovative, relative to cost-side drivers	pooled directional SEM path
H2 (EE)	positive	binds hardest for No-skill and Frontier	pooled path; H9 contrast
H3 (SI)	positive	strongest where internal evaluation capacity is scarce (No-skill)	pooled path
H4 (FC)	positive	binds hardest for No-skill and Frontier	pooled path; H9 contrast
H5 (CS)	positive	carries most weight for Frontier adopters	pooled path
H6 (TA)	positive	stronger for security-conscious and regulated adopters	pre-specified interaction
H7 (RR)	negative	mitigated by CA and FC; No-skill most exposed	pre-specified interactions
H8 (SA over PV)	relative	European public-sector buyers	nested comparison, public-sector subsample
H9 (structure)	directional	EE and FC weights larger, relative to PE, as capability falls	two-group directional contrast; continuous moderation; ratio contrasts (Section 5.4)

Table 3. Hypothesis summary. Directional tests at alpha .05; the H9 family carries a Holm correction.

5. RESEARCH DESIGN (NO DATA)

The design artifacts are the survey instrument with its analysis plan, the English master questionnaire with item wording and provenance, the pre-registration draft (frozen after the pilot and before any full-sample response), and the computed power analysis; all are listed in Annex A. This section states the design and its defense; no estimate appears anywhere in it.

5.1 Instrument and sampling frame

A dedicated instrument is needed because the official evidence base carries no measurement of the decision itself. The Commission’s deployment study fields a market survey of firms ($n = 380$, computer-assisted telephone interviewing) on investment volumes and self-assessed competitiveness, and asks non-adopters nothing about why they have not adopted. Its separate stakeholder survey ($n = 231$) ranks “Reliance on non-EU technologies or suppliers” thirteenth of fourteen deployment barriers, below “Insufficient demand for the technology in Europe” (European Commission 2026). Effort expectancy, facilitating conditions, and social influence have no counterpart in that evidence base, which records volumes and outcomes and leaves the calculus behind them unobserved. Two features of the market survey also bound what its numbers can support. Micro enterprises under ten employees make up 40 percent of its sample and small enterprises of ten to forty-nine a further 40 percent, so its firm-level averages do not describe European firms generally, and per-technology cells fall to roughly eight respondents for quantum and for interoperability.

The survey is Module A of a three-module instrument (architecture note in Annex A): the acceptance core measured here, a demand-pooling module feeding a companion empirical paper, and a short context block. Module order is fixed with A first, so the acceptance measurement is never preceded by pooling or sovereignty-trend items; item order is randomized within each battery; one instructed-response attention check (AC1) is embedded. The user path is about 79 items and targets 15 minutes; the provider path targets about 17.5 minutes with its mirror blocks.

The frame is European organizations that use, evaluate, or provide open source, reached through the CNLL, APELL, OW2, and Eclipse networks. The defense of this frame is that the estimand requires it. The population of interest is scarce, unenumerable, and unreachable by register-based sampling at any budget this study commands. Association channels reach exactly the organizations whose adoption calculus the model describes, including enough Innovative and Frontier adopters for a capability contrast to be estimable at all. The price is selection: the sample is a specialty sample of the open-source-engaged, so no population share is ever reported, all claims are within-sample and structural (which drivers bind where), and comparisons to representative benchmarks are labeled approximate context (Bitkom 2026). Quotas are set on the capability segments, with a floor of 100 completes per segment and a target of 450 or more in total (Section 5.5); fieldwork closes at the earlier of the target or twelve weeks, with non-response documented. Exclusion rules are pre-registered: failing AC1, completion under five minutes, straight-lining (zero variance across 15 or more consecutive Likert items), and duplicate organization-responder pairs, keeping the more senior role for the organization-level analysis.

5.2 The capability classifier and its behavioral validation

The grouping variable comes from four classification items measuring the typology's two capacities for the adopting team, each answered yes, partly, or no: evaluating and choosing among components (C1a) and writing an implementable specification (C1b) for the specification capacity; deploying and integrating against an API (C2a) and diagnosing and fixing defects in source, or porting across a breaking change (C2b), for the development capacity. A capacity counts as high when both items are yes, or one yes and one partly; Innovative means both capacities high, Frontier specification high and development low, No-skill both low. The items are concrete capacities instead of self-ratings, which blunts the main threat: overconfident self-assessment inflating capability (Section 3.4).

The classifier also carries a behavioral validation. AUX3 asks whether the team has had a patch or contribution accepted by an open-source project it did not lead in the past two years. An accepted upstream contribution is observable, hard to misremember, and close to the definition of the development capacity, so the Innovative assignment is checked against it: a high rate of Innovative-classified teams with no accepted contribution flags inflation, and the classification's sensitivity to reassigning them is reported. The same capacities, scored continuously, feed the co-primary moderation test (Section 5.4).

5.3 Measurement model and the invariance ladder

Estimation begins with a confirmatory factor analysis over the eleven reflective constructs: standardized loadings, composite reliability, average variance extracted, and HTMT discriminant validity are reported for every construct, and a construct that fails is dropped or respecified with the change documented. The formative use-behavior block is evaluated by indicator weights and collinearity, separately from the reflective machinery.

Cross-segment comparison then has to earn its license through the invariance ladder (Steenkamp & Baumgartner 1998). Configural invariance (the same factor pattern in every segment) establishes that the constructs exist in the same form across segments and licenses qualitative statements only. Metric invariance (equal loadings) makes a unit of a latent construct mean the same thing in every segment; it is the requirement for comparing structural paths, so it is the gate H9 must pass. Scalar invariance (equal intercepts) additionally licenses latent mean comparisons, which the descriptive outputs by segment need. If full metric invariance fails, partial invariance is the pre-registered fallback: comparisons are restricted to the constructs whose loadings hold, and the restriction is stated wherever the results appear. Without this ladder a cross-segment path difference is uninterpretable, since it could reflect different acceptance structures or merely different measurement, and nothing in the data separates the two.

5.4 Structural estimation and the H9 discipline

The structural model estimates H1 to H8 pooled, with H6 to H8 as pre-specified interactions and directional tests at alpha .05, then by capability segment. H9 is registered with two co-primary tests. The first is a two-group directional contrast, Innovative against the rest, on the effort-expectancy and facilitating-conditions paths, tested by a Wald z on the coefficient difference; two groups need one boundary instead of two, and the contrast uses the directional prior. The second is continuous-capability moderation: the classifier's capacities, scored continuously, interact with EE and FC in the pooled model, which uses all respondents on one parameter and is expected to be more efficient than any split. A Holm correction applies within the H9 family, and both tests are reported together regardless of agreement.

The companion theory paper (acceptance-structure-wp.md) contributes an identification discipline that the analysis plan adopts wholesale. In its model, every driver's marginal effect is the product of a structural sensitivity and the segment's density at its baseline adoption rate, so level comparisons across segments confound structure with location on the adoption curve. The confound is generic within the sampled space: in the companion's Monte Carlo, naive level comparisons misorder the capability ranking on 62.8 percent of random configurations under a probit link and 46.4 percent under a logit link, while ratios of marginal effects (effort to performance, facilitating conditions to performance) are free of the density and reproduce the true ordering on 100 percent of draws, and matching baselines eliminates the reversals entirely. The plan therefore states the H9 conclusion on ratio contrasts (EE and FC relative to PE) or at matched baseline adoption rates, reports each segment's baseline adoption rate next to every comparison, and treats level comparisons across segments with large baseline gaps as uninformative about structure. Invariance (Section 5.3) concerns whether the constructs are comparable; the ratio discipline concerns which comparisons of them carry structural meaning. Both must hold.

5.5 Power, quotas, and the binding constraint

Power was computed, never assumed (2,000 replications per cell, alpha .05; scripts, seeds, and result files in Annex B). The H9 contrast is approximated as a difference in a standardized path coefficient between two segments, estimated by per-group regression on composites with reliability 0.85 and inter-predictor correlation 0.40. For a true standardized difference of 0.25, 80 percent power needs 216 per group for the directional test and 272 two-tailed; at a difference of 0.15, power stays out of reach at feasible sizes (0.51 two-tailed and 0.64 directional at 400 per group). Along the directional curve at 0.25, power is 0.52 at 100 per group, 0.68 at 150, and 0.78 at 200. The scenario effect sizes are planning assumptions with no empirical anchor in this population, and are flagged as such.

The design accepts the consequences. The three-group pairwise H9 test would need 650 to 800 respondents while association distribution is expected to yield 150 to 300. Pre-registering it as the primary test would be planning to fail, so the two-group directional contrast and the continuous moderation are co-primary. The pre-registration commits to reporting achieved sensitivity (the minimum detectable difference at the realized N) whenever the realized sample falls short of 216 per group. Quota guidance is a floor of 100 completes per capability segment and a target of 450 or more in total. If totals land near 150 to 300, the paper reports the pooled paths (H1 to H8), the continuous moderation, and segment sensitivity bounds, and defers strong H9 claims to a second wave. For context, the shared fieldwork's companion paper tests proportion differences that need 87 per cell under its strong-sorting scenario, so the fieldwork delivers value at sample sizes where the H9 contrast is still underpowered.

5.6 The dyadic protocol

A subsample of organizations contributes paired interviews: a proposer (who writes, operates, or evaluates) and an authorizer (who owns budget, strategy, or governance) from the same organization, around one concrete adoption decision, recruited through a closing item. The anchor question asks why the authorizer said yes or no to the proposer's specific

proposal. The dyad serves three purposes. It cross-checks the segment findings against organizational accounts of real decisions. It is the built-in detector for the agency confound of Section 3.4: motive divergence lives between the proposer’s and the authorizer’s accounts of the same decision, and where the two disagree about why the proposal was made, both accounts are recorded as data. And it disciplines interpretation, because claims about organizational economics are made where the dyad corroborates them. The interviews are analyzed after the quantitative results are fixed, as corroboration, never as a source of new hypotheses for the same data.

6. PLANNED ANALYSES AND REPORTING

The tables below are shells: they fix, before any data exist, what will be reported and in what form. Every cell marked “to be estimated” will be filled from the registered analyses, and the shells are part of the commitment: a construct or contrast that fails does so in public, in the same table it would have occupied had it succeeded.

Construct	Items	Loading range	Composite reliability	AVE	Max HTMT
PE	PE1–PE4	to be estimated	to be estimated	to be estimated	to be estimated
EE	EE1–EE4	to be estimated	to be estimated	to be estimated	to be estimated
SI	SI1–SI3	to be estimated	to be estimated	to be estimated	to be estimated
FC	FC1–FC4	to be estimated	to be estimated	to be estimated	to be estimated
CS	CS1–CS4	to be estimated	to be estimated	to be estimated	to be estimated
TA	TA1–TA4	to be estimated	to be estimated	to be estimated	to be estimated
SA	SA1–SA4	to be estimated	to be estimated	to be estimated	to be estimated
RR	RR1–RR4	to be estimated	to be estimated	to be estimated	to be estimated
CA	CA1–CA3	to be estimated	to be estimated	to be estimated	to be estimated
PV	PV1–PV3	to be estimated	to be estimated	to be estimated	to be estimated
BI	BI1–BI3	to be estimated	to be estimated	to be estimated	to be estimated

Table 4. Measurement quality (shell). The formative UB block is evaluated by indicator weights and collinearity and reported separately.

Level	Equality constraint	What passing licenses	Fit and change in fit	Decision
Configural	same factor pattern in all segments	constructs exist in the same form; qualitative comparison	to be estimated	to be estimated
Metric	loadings equal across segments	structural paths comparable (the H9 gate)	to be estimated	to be estimated
Scalar	intercepts equal across segments	latent means comparable (segment descriptives)	to be estimated	to be estimated
Partial (fallback)	documented subset of constraints	comparison restricted to passing constructs	to be estimated	to be estimated

Table 5. Invariance ladder (shell), per Section 5.3.

Quantity	Pooled	Innovative	Frontier	No-skill
PE → BI	to be estimated	to be estimated	to be estimated	to be estimated

Quantity	Pooled	Innovative	Frontier	No-skill
EE → BI	to be estimated	to be estimated	to be estimated	to be estimated
SI → BI	to be estimated	to be estimated	to be estimated	to be estimated
FC → BI	to be estimated	to be estimated	to be estimated	to be estimated
CS → BI	to be estimated	to be estimated	to be estimated	to be estimated
TA → BI	to be estimated	to be estimated	to be estimated	to be estimated
SA → BI	to be estimated	to be estimated	to be estimated	to be estimated
RR → BI	to be estimated	to be estimated	to be estimated	to be estimated
EE/PE ratio	to be estimated	to be estimated	to be estimated	to be estimated
FC/PE ratio	to be estimated	to be estimated	to be estimated	to be estimated
Baseline adoption rate	to be estimated	to be estimated	to be estimated	to be estimated

Table 6. Path estimates by segment with the ratio contrasts and baselines required by Section 5.4 (shell). The H6 and H7 interactions and the H8 nested comparison are reported from the pooled model alongside this table.

The reporting commitments are those of the pre-registration (Annex A). All registered tests are reported regardless of outcome. Deviations from the registered plan are listed in a deviations section, with reasons. Achieved sensitivity (the minimum detectable difference at the realized per-group N) is stated whenever the realized sample falls short of the planning targets, in place of any claim to the planned power. The specialty-sample and stated-preference caveats appear in every output, including derived descriptive publications. Anonymized data and analysis code are released with the paper, subject to GDPR constraints.

7. THREATS TO VALIDITY

Common-method bias. All Module A constructs are measured on one instrument from one respondent, which can inflate construct correlations through shared method variance. The procedural remedies are in the instrument: fixed module order keeps advocacy-flavored pooling and sovereignty-trend content after the acceptance measurement; item order is randomized within batteries; referents are stated per section; behavioral anchors (AUX3, UB1 to UB3) break the Likert monotony with event- and fact-based items; an instructed-response attention check (AC1) and a “cannot say” option reduce acquiescence. The statistical remedy is planned: a method-factor analysis reported next to the main estimates (Podsakoff et al. 2003). Method variance that differs by segment would also distort the H9 contrast; the invariance ladder screens part of that, and the behavioral use block gives a same-survey criterion less exposed to it.

Selection into segments. Capability is endogenous in the long run: organizations hire, learn by adopting, and build the very capacities the classifier measures, so segment membership partially reflects past adoption success. The design is cross-sectional and is read that way: it estimates how the adoption calculus differs across current capability positions and claims nothing about what moving an organization across segments would do to its calculus. No causal claim about capability is registered, and the use-behavior block permits descriptive control for adoption history.

Intention-behavior gap. The outcome is stated intention, and the acceptance tradition’s intention-use link is one of its contested points (Bagozzi 2007). The use-behavior items (breadth, depth, policy) anchor intention in current practice and allow a consistency check, and they are a partial remedy only: they are self-reported, cross-sectional, and cannot stand in for observed future adoption. No forecast is made from intention levels.

Expressive and agency contamination. Sovereignty items invite expressive agreement, and the gap between stated sovereignty preference and stated willingness to pay in representative benchmarking (Bitkom 2026) shows how wide expressive responding can

run. Path weights also carry the agency and legitimacy variance of Section 3.4. The design responses are pre-specified: AUX1 and AUX2 enter as robustness covariates for H9; the dyad cross-report detects motive divergence on real decisions; commitment-style and behaviorally anchored items carry the willingness measurement in the companion module; and organizational-economics interpretations are reserved for findings the dyad corroborates.

Specialty sample. The frame reaches the open-source-engaged, and self-selection into association networks correlates with favorable attitudes toward everything this survey measures. Within-sample structural contrasts, the estimand of this study, are robust to that selection so long as it does not differ sharply by capability segment; distributions and levels are not. Accordingly no population share is reported, every output carries the specialty-sample caveat, and comparisons to representative benchmarks are labeled approximate context.

8. POSITIONING AND CONTRIBUTION

The contribution is to bring an economics-grounded segmentation into technology-acceptance research on open source, and to test whether the acceptance structure is segment-dependent, together with the two 2026 institutional moderators. This keeps UTAUT as the instrument (the constructs as measures) while replacing its representative user with the value-network typology, and it stays within the parsimony the model's critics demand. The typology is cited. It belongs to Jullien, Viseur & Zimmermann (2025) and the literature behind it. Whether AI assistance shifts an adopter's capability is a distinct claim with its own measurement problem, out of scope for this paper. And every coefficient is deferred: this paper's deliverable is a model, an instrument, a registered analysis plan, and the discipline the companion theory imposes on the central test.

The paper travels with companions. The theory companion (acceptance-structure-wp.md) derives the segment-dependence hypothesis from a cost primitive and disciplines its test. The shared fieldwork's second study models demand pooling for digital autonomy; until its data arrive, its theory and empirical design are one paper (collective-provision-wp.md, §7). A third companion (decision-map-wp.md) fixes, before fieldwork, what each estimable quantity would license for policy, advocacy and practice, which is why the present paper carries the design and none of the reading. Descriptives from the survey, cut by the capability segments, feed a trilingual policy barometer for the distributing associations, with the specialty-sample caveat fixed in every output.

9. CONCLUSION

Open-source adoption is decided by a chain of capacities and needs that the economics of open source already sorts into a typology, and by the institutional forces of 2026 that raise its risks and its strategic stakes. A UTAUT model that grounds its segmentation in that typology, keeps its constructs few, and adds regulation and sovereignty as moderators can say where the acceptance structure differs and why. This draft specifies that model to the level a referee can audit: constructs mapped to items, hypotheses to mechanisms, segments to predictions, tests to their identification conditions, and reporting to pre-committed shells. The estimation is the next phase, and until it is done the model puts forward no coefficient.

REFERENCES

- Abrahamson, E. (1991). Managerial fads and fashions: The diffusion and rejection of innovations. *Academy of Management Review*, 16(3), 586–612.
- ACM, AFM, AP, DNB, & RDI. (2026). *Towards digital autonomy: Increased control over digital dependencies*. Joint position paper of five Dutch supervisory authorities.

- Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179–211. [https://doi.org/10.1016/0749-5978\(91\)90020-T](https://doi.org/10.1016/0749-5978(91)90020-T)
- Ajzen, I., & Fishbein, M. (1970). The prediction of behavior from attitudinal and normative variables. *Journal of Experimental Social Psychology*, 6(4), 466–487. [https://doi.org/10.1016/0022-1031\(70\)90057-0](https://doi.org/10.1016/0022-1031(70)90057-0)
- APELL. (2021). *Europe's Open Source Industry's Statement on the Cyber Resilience Act*.
- Bagozzi, R. P. (2007). The Legacy of the Technology Acceptance Model and a Proposal for a Paradigm Shift. *Journal of the Association for Information Systems*, 8(4), 244–254. <https://doi.org/10.17705/1jais.00122>
- Bandura, A. (1986). *Social foundations of thought and action: A social cognitive theory*. Prentice-Hall, Inc. <https://psycnet.apa.org/record/1985-98423-000>
- Bikhchandani, S., Hirshleifer, D., & Welch, I. (1992). A theory of fads, fashion, custom, and cultural change as informational cascades. *Journal of Political Economy*, 100(5), 992–1026.
- Bitkom. (2026). *Cloud Report 2026: Status Quo & Trends in Wirtschaft und Politik*. Bitkom Research. doi:10.64022/2026-cloud-report-2026.
- Blind, K., Böhm, M., Grzegorzewska, P., Katz, A., Muto, S., Pätsch, S., & Schubert, T. (2021). *Study about the impact of open source software and hardware on technological independence, competitiveness and innovation in the EU economy*. European Commission, Directorate-General for Communications Networks, Content and Technology. <https://op.europa.eu/en/publication-detail/-/publication/d38e0892-7eb6-11eb-9ac9-01aa75ed71a1/language-en/format-PDF>
- Blind, K., & Schubert, T. (2023). Estimating the GDP effect of Open Source Software and its complementarities with R&D and patents: Evidence and policy implications. *The Journal of Technology Transfer*. <https://doi.org/10.1007/s10961-023-10033-1>
- Blut, M., Chong, A. Y. L., Tsiga, Z., & Venkatesh, V. (2022). Meta-Analysis of the Unified Theory of Acceptance and Use of Technology (UTAUT): Challenging its Validity and Charting a Research Agenda in the Red Ocean. *Journal of the Association for Information Systems*, 23(1), 13–95. <https://doi.org/10.17705/1JAIS.00719>
- Bonaccorsi, A., & Rossi, C. (2003). Why open source software can succeed. *Research Policy*, 32(7), 1243–1258.
- Chao, C.-M. (2019). Factors Determining the Behavioral Intention to Use Mobile Learning: An Application and Extension of the UTAUT Model. *Frontiers in Psychology*, 10, 1652.
- Chauhan, S., & Jaiswal, M. (2016). Determinants of acceptance of ERP software training in business schools: Empirical investigation using UTAUT model. *The International Journal of Management Education*, 14(3), 248–262. <https://doi.org/10.1016/J.IJME.2016.05.005>
- Cimperman, M., Makovec Brenčič, M., & Trkman, P. (2016). Analyzing older users' home telehealth services acceptance behavior—applying an Extended UTAUT model. *International Journal of Medical Informatics*, 90, 22–31. <https://doi.org/10.1016/j.ijmedinf.2016.03.002>
- CNLL & Inno³. (2024). *Cyber Resilience Act Compliance Guide for Open Source Players*.
- Daffara, C. (2007). *Logiciels libres et open source: Guide pour PME*. OPENTTT EU project (FP6-033982).
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly: Management Information Systems*, 13(3), 319–339. <https://doi.org/10.2307/249008>
- Davis, F. D., Bagozzi, R. P., & Warshaw, P. R. (1992). Extrinsic and Intrinsic Motivation to Use Computers in the Workplace. *Journal of Applied Social Psychology*, 22(14), 1111–1132. <https://doi.org/10.1111/J.1559-1816.1992.TB00945.X>

- DiMaggio, P. J., & Powell, W. W. (1983). The iron cage revisited: Institutional isomorphism and collective rationality in organizational fields. *American Sociological Review*, 48(2), 147–160.
- Dwivedi, Y. K., Rana, N. P., Chen, H., & Williams, M. D. (2011). A meta-analysis of the unified theory of acceptance and use of technology (UTAUT). *IFIP Advances in Information and Communication Technology*, 366, 155–170. https://doi.org/10.1007/978-3-642-24148-2_10
- Dwivedi, Y. K., Rana, N. P., Jeyaraj, A., Clement, M., & Williams, M. D. (2019). Re-examining the Unified Theory of Acceptance and Use of Technology (UTAUT): Towards a Revised Theoretical Model. *Information Systems Frontiers*, 21(3), 719–734. <https://doi.org/10.1007/s10796-017-9774-y>
- Eckhardt, A., Laumer, S., & Weitzel, T. (2009). Who influences whom? Analyzing workplace referents' social influence on IT adoption and non-adoption. *Journal of Information Technology*, 24(1), 11–24. <https://doi.org/10.1057/jit.2008.31>
- 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.
- European Commission. (2021a). *2030 Digital Compass: the European way for the Digital Decade*. COM(2021) 118 final, 9 March 2021.
- European Commission. (2021b). *Commission Decision of 8.12.2021 on the open source licensing and reuse of Commission software*. C(2021) 8759 final, OJ C 495I, 9 December 2021.
- European Commission, Directorate-General for Communications Networks, Content and Technology. (2026). *Study on the EU's critical digital capacities deployment beyond 2027: final report*. Prepared by Technopolis Group; K. Izsak, P. Simmonds, M. Ploeg, S. Viscido, A. Lotito, C. Moreno & R. Alasghar (eds.). Luxembourg: Publications Office of the European Union. ISBN 978-92-68-40079-1, doi:10.2759/0673677.
- Fitzgerald, B. (2006). The transformation of open source software. *MIS Quarterly*, 30(3), 587–598.
- Fritzsche, J., Wyrich, M., Bogner, J., & Wagner, S. (2021). Résumé-driven development: A definition and empirical characterization. In *Proceedings of the IEEE/ACM 43rd International Conference on Software Engineering: Software Engineering in Practice (ICSE-SEIP)* (pp. 19–28).
- Gupta, B., Dasgupta, S., & Gupta, A. (2008). Adoption of ICT in a government organization in a developing country: An empirical study. *The Journal of Strategic Information Systems*, 17(2), 140–154. <https://doi.org/10.1016/J.JSIS.2007.12.004>
- Hoepman, J.-H., & Jacobs, B. (2007). Increased Security Through Open Source. *Communications of the ACM*, 50(1), 79–83.
- Holmström, B. (1999). Managerial incentive problems: A dynamic perspective. *Review of Economic Studies*, 66(1), 169–182.
- Jullien, N., & Roudaut, K. (2012). Can open source projects succeed when the producers are not users? Lessons from the data processing field. *Management international = International management = Gestión internacional*, 16, 113–127.
- 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. (2009). Firms' contribution to open-source software and the dominant user's skill. *European Management Review*, 6, 130–139.
- Jullien, N., & Zimmermann, J.-B. (2011). Floss firms, users and communities: A viable match? *Journal of Innovation Economics*, 1(7), 31–53.
- Keong, M. L., Ramayah, T., Kurnia, S., & Chiun, L. M. (2012). Explaining intention to use an enterprise resource planning (ERP) system: An extension of the UTAUT model.

Business Strategy Series, 13(4), 173–180. <https://doi.org/10.1108/17515631211246249/FULL/XML>

- Kogut, B. M., & Metiu, A. (2001). Open source software development and distributed innovation. *Oxford Review of Economic Policy*, 17(2), 248–264.
- Kuwata, Y., Takeda, K., & Miura, H. (2014). A study on maturity model of open source software community to estimate the quality of products. *Procedia Computer Science*, 35, 1232–1241.
- Lerner, J., & Tirole, J. (2002). Some simple economics of open source. *The Journal of Industrial Economics*, 50(2), 197–234.
- Li, J. (2020). Blockchain technology adoption: Examining the Fundamental Drivers. In *Proceedings of the 2nd International Conference on Management Science and Industrial Engineering* (pp. 253–260). ACM. <https://doi.org/10.1145/3396743.3396750>
- Li, X., Moreschini, S., Zhang, Z., & Taibi, D. (2022). Exploring factors and metrics to select open source software components for integration: An empirical study. *Journal of Systems and Software*, 188, 111255.
- Li, X., Zhang, Y., Osborne, C., Zhou, M., Jin, Z., & Liu, H. (2024). Systematic literature review of commercial participation in open source software. *ACM Transactions on Software Engineering and Methodology*. <https://doi.org/10.1145/3690632>
- Lundell, B., & Gamalielsson, J. (2011). Adopting OSS in business. In *Proceedings of Open Source Systems: Grounding Research* (pp. 22–33). Springer, IFIP Advances in Information and Communication Technology.
- Menon, D., & Shilpa, K. (2023). “Chatting with ChatGPT”: Analyzing the factors influencing users’ intention to Use the Open AI’s ChatGPT using the UTAUT model. *Heliyon*, 9(11), e20962. <https://doi.org/10.1016/j.heliyon.2023.e20962>
- Moon, N. W., & Baker, P. M. A. (2009). *Adoption and Use of Open Source Software: Preliminary Literature Review*. Center for Advanced Communications Policy, Georgia Institute of Technology.
- Morgan, L., Feller, J., & Finnegan, P. (2013). Exploring value networks: Theorising the creation and capture of value with open source software. *European Journal of Information Systems*, 22, 569–588.
- Morgan, L., & Finnegan, P. (2014). Beyond free software: An exploration of the business value of strategic open source. *The Journal of Strategic Information Systems*, 23(3), 226–238.
- Mütterlein, J., Kunz, R. E., & Baier, D. (2019). Effects of lead-usership on the acceptance of media innovations: A mobile augmented reality case. *Technological Forecasting and Social Change*, 145, 113–124. <https://doi.org/10.1016/J.TECHFORE.2019.04.019>
- Nagle, F. (2017). *Open Source Software and Firm Productivity*. Harvard Business School Working Paper, No. 17-089.
- Podsakoff, P. M., MacKenzie, S. B., Lee, J.-Y., & Podsakoff, N. P. (2003). Common method biases in behavioral research: A critical review of the literature and recommended remedies. *Journal of Applied Psychology*, 88(5), 879–903.
- Rogers, E. (1995). *The Diffusion of Innovation*. The Free Press.
- Shaikh, M., & Cornford, T. (2011). *Total cost of ownership of open source software: A report for the UK Cabinet Office supported by OpenForum Europe*. Technical report, UK Cabinet Office.
- Steenkamp, J.-B. E. M., & Baumgartner, H. (1998). Assessing measurement invariance in cross-national consumer research. *Journal of Consumer Research*, 25(1), 78–90.
- Sykes, T. A., Venkatesh, V., & Gosain, S. (2009). Model of acceptance with peer support: A social network perspective to understand employees’ system use. *MIS Quarterly*, 33(2), 371–393. <https://doi.org/10.2307/20650296>

- Taylor, S., & Todd, P. A. (1995). Understanding Information Technology Usage: A Test of Competing Models. *Information Systems Research*, 6(2), 144–176. <https://doi.org/10.1287/ISRE.6.2.144>
- Thompson, R. L., Higgins, C. A., & Howell, J. M. (1991). Personal computing: Toward a conceptual model of utilization. *MIS Quarterly: Management Information Systems*, 15(1), 125–142. <https://doi.org/10.2307/249443>
- Van Raaij, E. M., & Schepers, J. J. L. (2008). The acceptance and use of a virtual learning environment in China. *Computers & Education*, 50(3), 838–852. <https://doi.org/10.1016/j.compedu.2006.09.001>
- Venkatesh, V., & Bala, H. (2008). Technology Acceptance Model 3 and a Research Agenda on Interventions. *Decision Sciences*, 39(2), 273–315. <https://doi.org/10.1111/J.1540-5915.2008.00192.X>
- 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>
- Venkatesh, V., Thong, J., & Xu, X. (2012). Consumer Acceptance and Use of Information Technology: Extending the Unified Theory of Acceptance and Use of Technology. *Management Information Systems Quarterly*, 36(1), 157–178. <https://aisel.aisnet.org/misq/vol36/iss1/13>
- Viseur, R., & Jullien, N. (2023). Communesplone, an original open-source model of resource pooling in the public sector. *IEEE Software*, 40, 46–54.
- von Hippel, E. (2001). Innovation by user communities: Learning from open-source software. *MIT Sloan Management Review*, 42(4), 82–86.
- von Hippel, E., & von Krogh, G. (2003). Open source software and the “private-collective” innovation model: Issues for organization science. *Organization Science*, 14(2), 209–223.
- Wang, Y.-S., Wu, M.-C., & Wang, H.-Y. (2009). Investigating the determinants and age and gender differences in the acceptance of mobile learning. *British Journal of Educational Technology*, 40(1), 92–118. <https://doi.org/10.1111/j.1467-8535.2007.00809.x>
- West, J. (2003). How open is open enough? Melding proprietary and open source platform strategies. *Research Policy*, 32(7), 1259–1285.
- Yaseen, M. G., A. Abd, S. A., & Adeb, I. (2019). Critical Factors Affecting the Adoption of Open Source Software in Public Organizations. *Iraqi Journal for Computer Science and Mathematics*, 1(1).

ANNEX A: WORKING NOTES (TEMPORARY ANNEX)

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

- `../notes/stratified-adoption/instrument.md`: Module A, the acceptance survey, and the multi-group analysis plan.
- `../notes/stratified-adoption/prereg-paper1.md`: the pre-registration draft, to be frozen after the pilot and before any full-sample response.
- `../notes/common/questionnaire-en.md`: the English master questionnaire with item-level wording, provenance, and budgets.
- `../notes/common/poll-design.md`: the one-poll, three-module survey architecture.
- `../notes/common/power-analysis.md`: computed power and quota guidance.
- `../notes/common/model-reconciliation.md`: the full hypothesis cut list (the source drafts’ 23 hypotheses onto the pruned set).
- `../notes/common/behavioral-agency.md`: the agency and fashion taxonomy and the auxiliary items AUX1 to AUX3.
- `../notes/common/theory-foundation.md`: the shared theoretical foundation of the project.

- `../notes/common/research-program.md`: the wider research program.

ANNEX B: REPRODUCTION: SCRIPTS, DATA, AND FIGURES

- `../verify/power_analysis.py` (seed 20260716, 2,000 replications per cell, alpha .05): the H9-contrast power simulation of Section 5.5; writes `../verify/results/power_analysis.json` and generates the figure `figures/power-curves.png`.
- `../verify/acceptance_structure.py` (seed 20260716): the companion theory paper's verification suite, source of the misordering shares quoted in Sections 4 and 5.4; results in `../verify/results/acceptance_structure.json`.
- `../verify/run_all.py`: reproduces every computed number and figure of the project in one command.
- The design artifacts (questionnaire and pre-registration drafts) are listed in Annex A.