

Perceived Architectural Fit in Workplace Reliance on Artificial Intelligence: Evidence from a UTAUT-Based Model

Author(s): Eric Strandt¹, Daniel N. Strandt²

1 Indiana Tech, College of Business, Fort Wayne, USA.

2 Independent Researcher.

Corresponding Author: EStrandt01@indianatech.net

Received: July, 2026, Published: Sep, 2026

ARTICLE INFO

Keywords:

AI adoption; Perceived architectural fit; UTAUT; Workplace AI reliance

© 2026 by the Author(s). This open-access article is distributed under a Creative Commons Attribution (CC-BY) 4.0 license, making research freely available to the public and supporting a greater global exchange of knowledge and human experiments.



ABSTRACT

Artificial intelligence (AI) is increasingly integrated into workplace operations through general-purpose systems, including large language models, foundation-model platforms, and AI assistants. The Unified Theory of Acceptance and Use of Technology (UTAUT) explains adoption through performance expectancy, effort expectancy, social influence, and facilitating conditions; however, the four UTAUT-related constructs used here do not directly represent whether an AI system's technical and governance supports fit the accountability demands of its intended use. This study tests perceived architectural fit, defined as the perceived match between those supports and the accountability demands of the intended use, as an added predictor in a UTAUT-based model. A quantitative vignette-based survey with prebalanced 2 × 2 conditions and post-vignette perception measures collected 350 responses from U.S. workers with workplace AI exposure and retained 282 after screening. The four UTAUT-related predictors jointly explained 59.7% of the variance in AI reliance intention (ARI). Adding perceived architectural fit increased explained variance by 10.0 percentage points, and perceived architectural fit remained significant in the full model. Participants receiving the enhanced architectural-support description reported higher perceived architectural fit than participants receiving the limited-support description. Assigned conditions did not produce statistically significant differences in perceived legitimacy or ARI. Secondary analyses found that perceived architectural fit was associated with perceived legitimacy and that perceived legitimacy was associated with ARI. A bootstrap analysis also identified a positive indirect association through perceived legitimacy after controlling for the UTAUT-related predictors. The findings provide quantitative evidence that perceived architectural fit may contribute explanatory value to UTAUT-based research on workplace ARI.

1.0 Introduction

Artificial intelligence (AI) is reshaping workplace productivity by supporting tasks such as summarization, drafting, classification, recommendation, forecasting, and decision support (Bommasani et al., 2021). Enterprise AI assistants and large language models (LLMs) are entering organizations through common platforms, interfaces, and implementation methods, not only as stand-alone applications (Wang et al., 2026). These developments make AI acceptance and AI reliance important issues for organizational technology adoption research (Xia & Chen, 2025).

Technology acceptance research offers a strong starting point for addressing those issues. Venkatesh et al. (2003) explain technology adoption through the Unified Theory of Acceptance and Use of Technology (UTAUT), which identifies performance expectancy, effort expectancy, social influence, and facilitating conditions as central predictors of behavioral intention and use. Performance expectancy assesses whether people believe that using a system will improve task performance. Effort expectancy concerns perceived

ease of use. Social influence concerns whether important others are perceived to support system use. Facilitating conditions concern whether organizational and technical resources are available to support use. In this study, the UTAUT outcome logic is adapted to AI reliance intention (ARI) rather than general behavioral intention.

AI reliance is a more specific adoption question than general acceptance or use. In this study, ARI refers to respondents' self-reported intention to use the AI system, support its deployment, and rely on its outputs as part of the work process described in the vignette. This outcome adapts UTAUT's behavioral-intention logic to the specific question of workplace AI reliance. AI systems are often treated as digital tools, but responsible AI frameworks show that deployment also depends on documentation, oversight, accountability, transparency, robustness, and risk controls (National Institute of Standards and Technology, 2023). For example, two similar systems could differ in whether they support traceability, maintain audit logs, communicate uncertainty, reproduce results, or enable human review. Raji et al. (2020) argue that reviewable artifacts and audit records are central to closing accountability gaps in AI systems. In high-accountability settings, these architectural and governance differences may affect whether relying on an AI system appears reasonable, defensible, reviewable, open to challenge, and legitimate (European Parliament & Council of the European Union, 2024).

This paper argues that research on AI adoption and reliance should consider perceived architectural fit: the perceived match between an AI system's technical and governance architecture and the accountability demands of the use context. Perceived architectural fit reflects whether users believe the system's technical and governance support is adequate to meet the task's accountability demands. The vignette design varied the architectural and governance features described in the scenario, while the architectural-fit scale assessed whether participants perceived those features as fitting the task's accountability needs. This construct, therefore, moves AI adoption research from general acceptance toward context-sensitive reliance judgments. For example, an AI system may be useful and easy to use but still be a poor fit when the task requires traceability and the system lacks an audit log, or when stable records are needed, and the system blurs stored records with generated claims.

This argument depends on a distinction between perceived task benefit and perceived legitimacy of reliance. This study did not include a separate capability-trust scale. Performance expectancy captured perceived task benefit, while perceived legitimacy captured whether reliance on the system appeared appropriate, defensible, reviewable, open to challenge, and governable in the specific context. In this paper, perceived legitimacy is an individual judgment about the legitimacy of relying on the described system. It is not a measure of field-level institutional legitimacy.

This study uses the vignette design in two ways. First, it tests whether ratings differed across the scenario conditions. Second, it examines whether participants who perceived stronger architectural fit also saw the AI system as more legitimate and reported stronger ARI. This distinction keeps the study's claims clear: the paper tests both assigned-condition differences and perception-based relationships, but it does not test actual workplace reliance behavior.

The study's principal contribution is an incremental-prediction test of perceived architectural fit within a UTAUT-based model of workplace ARI. The prebalanced vignette conditions provide supporting evidence that participants distinguished the architectural-support descriptions and reported different levels of perceived architectural fit. Associations involving perceived legitimacy are examined as a secondary explanatory pattern rather than as causal mediation or full construct validation.

2.0 Literature Review and Conceptual Development

2.1 UTAUT and AI Adoption

UTAUT remains a central framework for explaining technology acceptance and use. The model identifies performance expectancy, effort expectancy, social influence, and facilitating conditions as core predictors of behavioral intention and use (Venkatesh et al., 2003). Research on AI acceptance continues to support the applicability of UTAUT while showing where AI raises additional concerns. For example, Xia and Chen (2025) found that task-tool fitness was the strongest positive predictor of behavioral intention to use generative AI in new product development teams.

Consistent with the reviewed literature, this study retains four UTAUT-related predictors as an acceptance benchmark and adds PAF as a separate predictor of ARI. This preserves UTAUT's focus on intention while recognizing that workplace AI use often requires workers to rely on system outputs within work processes.

2.2 General-Purpose AI and Architecture-Sensitive Reliance

General-purpose AI systems differ from many earlier workplace technologies because the same broad class of system can be applied across many tasks and accountability settings. Bommasani et al. (2021) describe foundation models as adaptable systems that can support many downstream uses. Wang et al. (2026) show that enterprise-wide LLM implementation is a staged sociotechnical process rather than a single act of deploying technology. Foundation models and enterprise AI assistants can support drafting, retrieval, classification, and decision support (Bommasani et al., 2021). This flexibility is valuable, but it creates a fit issue. A system that is suitable for informal drafting may not provide the same basis for reliance in work involving formal decisions or compliance documents. Organizational LLM adoption can stall when experimentation is not matched with governance, implementation pathways, and assigned responsibility (Wang et al., 2026).

Responsible AI use and risk-management frameworks make this fit issue visible. The National Institute of Standards and Technology (NIST) AI Risk Management Framework emphasizes governance, mapping, measurement, and management of AI risks across the system lifecycle (National Institute of Standards and Technology, 2023). Similarly, the NIST generative AI profile identifies documentation, evaluation, content provenance, risk monitoring, and governance controls as relevant to managing generative AI risks (Autio et al., 2024). Additionally, the European Union AI Act links higher-risk AI systems with obligations for risk management, transparency, human oversight, recordkeeping, and technical documentation (European Parliament & Council of the European Union, 2024). Although these frameworks are not technology adoption models, they identify system features that may matter to adoption, as users may ask whether reliance can be traced, reviewed, challenged, reproduced, and governed.

Established technology acceptance theories address multiple determinants of technology use; however, the UTAUT predictors used in this study do not directly capture whether an AI system's technical and governance support fits the accountability demands of a specific use context. Two AI systems may appear equally capable for a task but differ in whether they preserve source content, maintain audit trails, communicate uncertainty, reproduce results, support human review, or allow local governance configuration. These design and governance differences can shape whether users consider reliance on the system acceptable within the described context.

2.3 Trust, Accountability, and Appropriate Reliance

Trust and reliance are related but conceptually distinct. Trust is an attitude toward a system, whereas ARI in this study represents behavioral intention to use, support deployment of, and rely on the AI system. Trust may influence such reliance intentions, but appropriate reliance also requires that reliance be calibrated to whether use is reasonable in the specific

context. Lee and See (2004) define appropriate reliance as a calibration issue in which users rely on automation when reliance is reasonable and avoid reliance when it is not. Parasuraman and Riley (1997) distinguish use, misuse, disuse, and abuse of automation. Hoff and Bashir (2015) show that trust in automation is shaped by system performance, process, purpose, user experience, work environment, and context. Huynh and Aichner (2025) extend this concern to generative AI by treating fairness, accountability, transparency, and related responsible-AI concerns as antecedents of cognitive trust. In their model, trust also shapes attitudes, perceived performance, and behavioral intention toward generative AI systems. This distinction is relevant in this study because architectural and governance characteristics may inform attitudes toward an AI system while ARI captures respondents' intended reliance on it. This distinction is especially relevant to generative AI because fluent outputs may contain errors or use contextual information inconsistently, making calibrated reliance necessary (Farquhar et al., 2024; Liu et al., 2024).

Research on algorithmic accountability reaches a similar conclusion from a governance standpoint. Internal auditing frameworks emphasize artifacts, review points, and records that make AI development and deployment traceable across the system lifecycle (Raji et al., 2020). Because organizations and users may hesitate to rely on AI when the conditions for review or accountability are missing, trust issues in AI use can arise from psychological, organizational, and architectural sources.

2.4 Perceived Legitimacy of Reliance in AI Use

Legitimacy is well established in institutional and organizational theory. Suchman (1995) defines legitimacy as a generalized perception that actions are desirable, proper, or appropriate within a socially constructed system of norms, values, beliefs, and definitions. Scott (2014) locates legitimacy within institutional systems that include regulative, normative, and cultural-cognitive elements. In this study, that theory provides conceptual background rather than a direct measure of organizational or field-level legitimacy.

Applied to AI reliance, this suggests that users may assess not only whether an AI system can help them complete a task, but also whether reliance on the system can be justified to supervisors, affected people, auditors, or the organization itself. The measured construct is therefore perceived legitimacy of reliance: an individual respondent's judgment that reliance on the described system appears appropriate, defensible, reviewable, open to challenge, and governable. In this study, perceived legitimacy is related to accountability because it reflects whether reliance can be justified and reviewed in light of the task's responsibilities and accountability demands. Perceived legitimacy of reliance is a more context-specific application of legitimacy to AI use and should be read as an individual perception formed in response to the vignette rather than as a measure of institutional legitimacy. The next section details perceived architectural fit, which is proposed as one reason reliance may appear more or less legitimate in a specific AI use context.

2.5 Perceived Architectural Fit

Perceived architectural fit (PAF) refers to respondents' perceived match between the technical and governance supports offered by an AI system and the accountability demands described in the context of use. Conceptually, PAF adapts the logic of task-technology fit, which considers whether system functions match task requirements (Goodhue & Thompson, 1995), by focusing specifically on whether technical and governance supports align with accountability requirements. Its content is also informed by responsible AI guidance, which identifies documentation, provenance, risk monitoring, and governance controls as relevant to responsible AI use (Autio et al., 2024). Together, these sources provide a rationale for the construct's conceptual content, but they do not constitute empirical validation of PAF.

PAF offers a focused way to understand why some AI systems may be accepted for routine work yet questioned in settings that require reviewability, provenance, record separation,

reproducibility, uncertainty communication, and governance support. Table 1 situates PAF relative to constructs selected from the study's UTAUT benchmark, task-technology fit, and perceived legitimacy: performance expectancy and facilitating conditions derive from UTAUT (Venkatesh et al., 2003), task-technology fit from Goodhue and Thompson (1995), and perceived legitimacy from the legitimacy literature discussed above (Suchman, 1995; Scott, 2014). The main questions in Table 1 summarize the central concern of each construct and clarify its distinction from PAF.

Table 1. Conceptual Distinctions Between Perceived Architectural Fit and Related Constructs

Construct	Main question	Difference from perceived architectural fit
Performance expectancy	Will the system improve performance?	Architectural fit concerns accountability and review support, not task benefit alone.
Task-technology fit	Do system functions fit task requirements?	Architectural fit concerns whether review, provenance, records, uncertainty, and governance supports fit accountability requirements.
Facilitating conditions	Are organizational and technical resources available?	Architectural fit concerns the match between the system's support design and the use context.
Perceived legitimacy	Is relying on the system appropriate and defensible?	Architectural fit is a system-context match judgment; legitimacy is a broader normative evaluation of reliance.

Note: Performance expectancy and facilitating conditions are based on Venkatesh et al. (2003); task-technology fit is based on Goodhue and Thompson (1995); perceived legitimacy draws on Suchman (1995) and Scott (2014) as applied to AI reliance in this study.

2.6 Architectural and Governance Features Informing the Construct

Perceived architectural fit is operationalized as an overall perceived-fit construct. The item pool is informed by seven feature categories: auditability, provenance, record separation, reproducibility, uncertainty communication, governance visibility, and local control. These categories align with the NIST generative AI profile, which identifies documentation, evaluation, content provenance, risk monitoring, and governance controls as relevant to generative AI risk management (Autio et al., 2024).

Table 2. Architectural and Governance Features Informing Perceived Architectural Fit

Feature category	Definition
Auditability	The system leaves a reviewable trail of inputs, outputs, sources, decisions, and relevant process steps.
Provenance	The system identifies the sources, records, or evidence used to generate an output.
Record separation	The system distinguishes official records, user notes, retrieved evidence, generated suggestions, and temporary context.
Reproducibility	The system's process and outputs can be repeated or reconstructed when needed.
Uncertainty communication	The system communicates limits, ambiguity, missing information, and conditions under which outputs should not be relied on.
Governance visibility	Users can identify review paths, escalation procedures, human oversight, policy constraints, and accountability mechanisms.

Feature category	Definition
Local control	The organization can configure relevant settings, data use, review rules, retention practices, or governance requirements.

These features define the conceptual content used to develop the vignette and global perceived-fit items. In this study, accountability is operationalized as the level of consequences and review requirements associated with the task, ranging from routine internal work to work involving compliance obligations, professional duties, formal decisions, or institutional records. They are not proposed or tested as seven empirical subscales in this study. Future research should examine whether perceived architectural fit is best modeled as a reflective construct, a formative construct, or a higher-order construct.

3.0 Research Model and Hypotheses

3.1 Research Model

The research model adds perceived architectural fit to an adapted UTAUT-based intention model. Performance expectancy, effort expectancy, social influence, and facilitating conditions form the acceptance benchmark for ARI. Perceived architectural fit is the focal added predictor, while perceived legitimacy is examined as a secondary explanatory perception.

The primary outcome is ARI, representing the respondent's intention to use, support deployment of, or rely on the AI system as part of the work process described in the vignette. Accountability condition is included as an assigned boundary condition rather than as another UTAUT construct.

The UTAUT-related predictors are expected to predict ARI, and perceived architectural fit is expected to add explanatory value beyond that block. Perceived architectural fit is also expected to be associated with perceived legitimacy, which is expected to be associated with reliance intention. The indirect pathway is treated as a secondary concurrent-association analysis. Figure 1 illustrates this design.

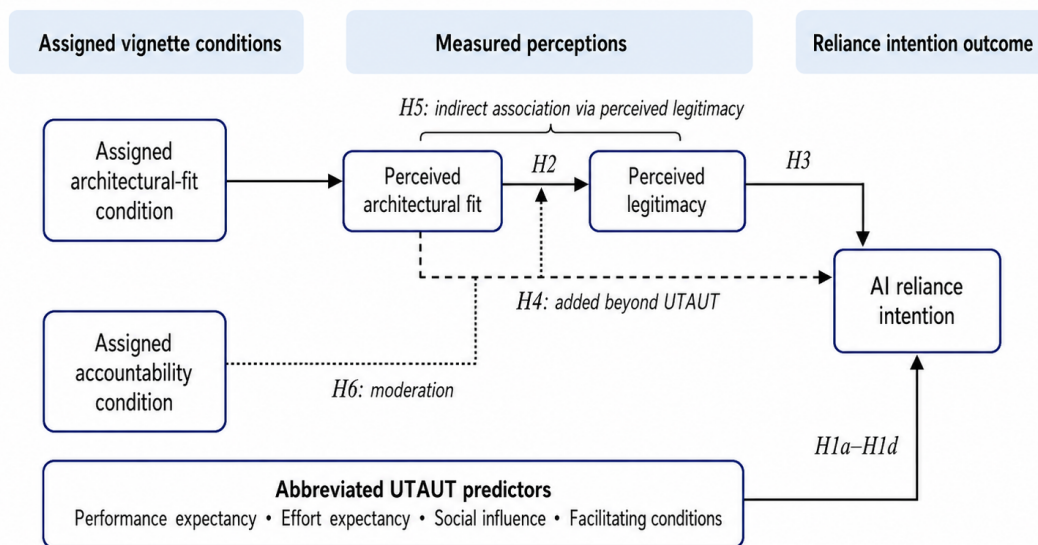


Fig. 1. Research Model and Evidence Boundaries

The hypotheses define these expectations as follows: H4 provides the principal test of whether PAF offers incremental statistical prediction of ARI beyond the four UTAUT-related predictors used in this study. H1a through H1d establish the contribution of the UTAUT-related predictor block. H2, H3, and H5 are secondary explanatory hypotheses, and H6 tests a possible accountability boundary condition.

3.2 Hypotheses

The following hypotheses guided the analysis:

- H1a: Performance expectancy positively predicts ARI toward the AI system.
- H1b: Effort expectancy positively predicts ARI toward the AI system.
- H1c: Social influence positively predicts ARI toward the AI system.
- H1d: Facilitating conditions positively predict ARI toward the AI system.
- H2: Perceived architectural fit is positively associated with perceived legitimacy.
- H3: Perceived legitimacy is positively associated with ARI toward the AI system.
- H4: Perceived architectural fit contributes significant incremental explanatory power to ARI beyond the UTAUT-related predictor block.
- H5: Perceived architectural fit has a positive indirect association with ARI through perceived legitimacy.
- H6: Accountability level moderates the relationship between perceived architectural fit and perceived legitimacy, such that the relationship is stronger in the high-accountability condition than in the low-accountability condition.

4.0 Method

4.1 Research Design

A quantitative vignette-based survey with prebalanced 2×2 conditions and post-vignette perception measures tested the hypotheses. Vignette methods are useful when researchers need to vary features of a described situation while holding other details constant (Aguinis & Bradley, 2014). Accountability was rated as low or high, and bundled architectural support as limited or enhanced. Assigned-condition analyses tested condition differences. Hierarchical regression and indirect-effect analyses tested associations among measured post-vignette perceptions. Each participant viewed one condition and then completed the same survey measures.

4.2 Participants and Sampling Strategy

The target population was U.S. workers reporting workplace AI exposure. Participants were recruited through Amazon Mechanical Turk (MTurk) in May 2026. The MTurk batch opened on May 1, 2026, and closed on May 2, 2026, after reaching 350 completed responses. Although the posting allowed a maximum two-week collection window, the sampling goal was met before the full window elapsed. Data collection therefore closed early. The study was administered entirely through MTurk, and no external survey platform was used.

The MTurk task was restricted to U.S.-based workers through MTurk posting settings. Eligible participants were required to report that they were currently employed, self-employed, or employed or self-employed within the previous 12 months and that they had workplace familiarity with AI or AI-enabled digital tools. The workplace AI item was intentionally broad and did not verify experience with a specific generative AI product, task type, or recent frequency beyond the background item reported below.

The MTurk batch contained 350 survey tasks. Each task row contained one preassigned vignette condition, and each task allowed one completed assignment. The vignette condition was assigned via a prebalanced MTurk batch file rather than by an external survey randomizer. Workers saw the scenario text rather than analytic condition labels, and completed responses were matched to assigned conditions through MTurk input variables. Duplicate-response screening based on MTurk Worker ID was used to ensure that all cases in the analytic dataset were unique respondents. The batch was balanced as closely as possible across the four cells, with 87 or 88 task rows per condition and equal marginal counts for accountability level and architectural-support condition.

The minimum usable analytic target was 256 respondents, or at least 64 usable respondents per assigned condition. This threshold provided a practical cell-level floor for detecting medium-condition effects while preserving sufficient observations for the intended regression-based tests (Cohen, 1988). The target was better suited to main-effect and regression tests than to small interaction effects, which is important for interpreting H6. The 350-response fielding target provided a buffer above that minimum after eligibility, attention-check, duplicate-response, incomplete-response, and completion-time exclusions.

Demographic quotas were not used. Because this study examined relationships among model variables rather than population prevalence estimates, a nonprobability MTurk sample was acceptable for this initial empirical test. MTurk has been used successfully in behavioral and social science research when appropriate screening and data-quality procedures (Thomas & Clifford, 2017). Sampling limitations are discussed in the Limitations and Future Research section.

4.3 Procedure

Workers participated in the study directly through the MTurk platform. The introductory screen described the study purpose in general terms, stated that participation was voluntary, and asked participants to provide electronic consent before continuing. Participants then answered screening questions about their employment status and familiarity with AI or AI-enabled digital tools in the workplace. Eligible participants received one of the four vignette conditions through the prebalanced MTurk input file. Each task row contained the assigned condition variables and the corresponding vignette text. Workers were not shown analytic condition labels; they saw only the participant-facing scenario text and survey items.

A single background item near the end of the survey asked participants how often they used or interacted with AI-enabled tools in their current or most recent work. This item was used to describe the sample and review consistency in workplace AI exposure. The survey also included one yes/no attention-check item placed approximately midway through the survey. Workers who did not answer this item correctly were removed from the analytic sample in accordance with the data-quality rules described below.

Participants were automatically compensated through MTurk for a task described as taking approximately 3 minutes. The observed median completion time across all submitted responses before screening was 154 seconds, or approximately 2.6 minutes.

4.4 Vignette Development

The vignettes were written to be similar in length, tone, and task setting. Evaluative labels such as trustworthy, legitimate, safe, risky, accountable, and reliable were avoided because those judgments were measured as outcomes rather than supplied directly in the scenario. Table 3 provides the participant-facing components needed to reconstruct all four conditions.

The accountability manipulation varied the task's consequences and review requirements. The low-accountability descriptions depicted routine internal planning tasks, including drafting meeting agendas, summarizing brainstorming notes, and preparing project-

planning drafts. The high-accountability descriptions depicted AI-assisted work that could affect compliance obligations, professional duties, formal decisions, or institutional records.

The architectural-support manipulation varied the system features that supported responsible reliance. The enhanced-support versions described source links, reviewable process records, separation of official records from generated suggestions, uncertainty notices, reproducible output processes, human review pathways, and local governance options. The limited-support versions described a professionally implemented AI assistant designed for flexible drafting and fast text generation, but not formal review trails, source-level process records, record separation, or configurable governance options. The architectural-support condition should therefore be understood as a bundled description rather than as a test of any single feature.

The limited-support versions were worded neutrally so they did not describe a low-quality or malfunctioning system, but the enhanced-support versions necessarily provided more positive and specific support details. Both architectural-support versions described an AI assistant capable of generating useful text quickly.

Table 3. Participant-Facing Vignette Components

Component	Wording
Instruction	Please read the scenario carefully. After reading it, answer the questions based only on the organization, task, and AI assistant described in the scenario.
Low accountability block	A medium-sized organization is considering an AI assistant to help employees prepare routine internal planning materials. Employees would use the assistant to draft meeting agendas, summarize brainstorming notes, and create first drafts of project-planning documents. The materials would be reviewed informally by the employees who create them and could be edited before being shared with a team.
High accountability block	A medium-sized organization is considering an AI assistant to help employees prepare materials for decisions that may affect compliance obligations, professional duties, and institutional records. Employees would use the assistant to summarize case information, draft recommendations, and prepare documents that supervisors may use when making formal decisions. The materials could become part of an internal record and may later need to be reviewed.
Limited architectural-support block	The AI assistant has been professionally implemented for fast drafting and flexible interaction. It can produce useful text quickly and respond to follow-up prompts. The system helps users generate and revise text, but it is not designed to create a formal review trail for each output. It does not provide source-level process records for later reconstruction. Its workspace does not clearly separate official records, user notes, and generated suggestions. Local managers have limited ability to configure review steps, retention settings, or governance rules for the tool.
Enhanced architectural-support block, low-accountability condition	The AI assistant has been professionally implemented with process support features. It can produce useful text quickly and respond to follow-up prompts. It provides links to the sources or prior notes used in each answer. It keeps a reviewable record of prompts, outputs, and source materials. It separates official records, user notes, and generated suggestions. It includes notices when information is incomplete or uncertain. It can reproduce the same output process when the same materials and settings are used. It also supports defined review steps that local managers can configure along with retention settings and governance rules for the tool.

Component	Wording
Enhanced architectural-support block, high-accountability condition	The AI assistant has been professionally implemented with process support features. It can produce useful text quickly and respond to follow-up prompts. It provides links to the sources or prior records used in each answer. It keeps a reviewable record of prompts, outputs, and source materials. It separates official records, user notes, and generated suggestions. It includes notices when information is incomplete or uncertain. It can reproduce the same output process when the same materials and settings are used. It also supports defined review steps that local managers can configure along with retention settings and governance rules for the tool.

Note. Each participant saw the instruction, one accountability block, and one architectural-support block. The enhanced-support wording differed only in the reference to prior notes for low-accountability work and prior records for high-accountability work.

4.5 Measures

Each focal construct was measured with five-point Likert-style items, with responses ranging from 1 (Strongly Disagree) to 5 (Strongly Agree). To reduce participant burden while retaining multi-item measurement, each focal construct was measured with three items. The survey included measures of performance expectancy, effort expectancy, social influence, facilitating conditions, perceived architectural fit, perceived legitimacy, ARI, and manipulation checks.

The four UTAUT-related measures were brief, scenario-specific adaptations used to provide a theoretically grounded benchmark for acceptance. They were not intended to reproduce the complete original UTAUT instrument. Performance expectancy measured perceived usefulness for the described task; effort expectancy measured perceived ease of use; social influence measured conditional support from important work contacts; and facilitating conditions measured perceived organizational resources, guidance, and assistance.

Perceived architectural fit measured whether the AI system's design was compatible with the accountability and review demands described in the work context. The scale assessed overall perceived fit rather than separate feature-category subscales. Item development was guided by auditability, provenance, record separation, reproducibility, uncertainty communication, governance visibility, and local control. The study did not independently measure objective context-level requirements.

Perceived legitimacy measured whether reliance on the AI system appeared appropriate, defensible, reviewable, open to challenge, and consistent with the responsibilities of the task. One perceived-legitimacy item asked whether reliance could be reviewed or challenged, which creates possible item-content overlap with architectural fit. Because that item overlapped with architectural-support content, a reduced two-item perceived-legitimacy score excluding that item was also computed for sensitivity analyses.

ARI was treated as a composite measure. Its three items captured intention to use the AI system, support for deployment, and willingness to rely on the system as part of the work process. Because these items reflected related but distinct judgments, item-level sensitivity models were also estimated. Scale scores were created by averaging items within each construct and were treated as approximately continuous variables for regression-based analysis.

Table 4. Construct and Manipulation-Check Measures

Construct or check	Items
Performance expectancy (PE)	1. This AI system would help me complete the task more effectively. 2. This AI system would improve the quality of the work produced. 3. This AI system would be useful for the work described in the scenario.
Effort expectancy (EE)	1. Learning to use this AI system would be easy for me. 2. I would find this AI system clear and understandable to use. 3. Using this AI system would not require unnecessary effort.
Social influence (SI)	1. I would be more likely to use this AI system if coworkers supported its use. 2. I would be more likely to use this AI system if supervisors encouraged its use. 3. People important to my work would likely support use of this AI system.
Facilitating conditions (FC)	1. The organization would likely have the resources needed to support use of this AI system. 2. I would likely receive enough guidance to use this AI system appropriately. 3. I would likely have access to help if I had problems using this AI system.
Perceived architectural fit (PAF)	1. The system design appears appropriate for the level of accountability in this scenario. 2. The way the system handles records, review, and uncertainty fits the kind of work described. 3. The system provides the right kind of support for responsible reliance in this setting.
Perceived legitimacy (PL)	1. The organization could justify allowing this system for the task described. 2. The system would support a defensible work process. 3. The system would allow reliance to be reviewed or challenged if needed.
AI reliance intention (ARI)	1. If I worked in this organization, I would intend to use this AI system for the task described. 2. I would support deploying this AI system for the task described. 3. I would be willing to rely on this AI system as part of the work process.
Accountability manipulation check	1. The task described in the scenario could have meaningful consequences for the organization or affected people. 2. Errors in this task could create significant problems.
Architectural-support recognition check	1. The AI system described in the scenario provides a reviewable process. 2. The AI system described in the scenario makes sources or records traceable.
Employment eligibility screening	Which best describes your current employment status? Options: currently employed; self-employed; employed or self-employed within the previous 12 months; not employed or self-employed within the previous 12 months. The first three responses were eligible.
Workplace AI familiarity screening	In your current or most recent work, have you used or been exposed to AI or AI-enabled digital tools? Options: Yes or No. Yes was required for eligibility.
Workplace AI-use frequency	How often have you used or interacted with AI-enabled tools in your current or most recent work? Options: Never, Rarely, Sometimes, Often, Very often.

Note. Items used a five-point response scale from 1 (strongly disagree) to 5 (strongly agree). Construct composites and manipulation-check scores were item means.

4.6 Data Quality and Analysis

Data was screened before analysis. Survey responses were excluded if participants failed eligibility screening, did not complete the survey, failed the attention check, submitted duplicate responses, gave inconsistent workplace AI exposure information, or completed the survey implausibly quickly. Duplicate responses were identified using MTurk Worker IDs before the IDs were removed from the analytic dataset.

Completion-time screening used a median-relative approach. Prior research has shown that rapid completion can be associated with satisficing because respondents in web surveys may have incentives to finish quickly while minimizing effort (Malhotra, 2008). Malhotra (2008) also cautions that completion time is not a perfect standalone measure of attention. For that reason, completion time was used only as one part of a broader screening process. Responses were excluded for speed only when completion time fell below one-third of the median completion time for submitted responses. In this study, that rule produced a cutoff of 85.4 seconds.

Hypothesis tests were conducted using JASP, an open-source statistical package used in psychological and social science research (JASP Team, 2025). Analyses began with descriptive statistics, Cronbach's alpha, correlations, condition counts, manipulation checks, and 2×2 analyses of variance. H1a through H1d were tested by regressing ARI on the four UTAUT-related predictors. H4 tested whether PAF provided incremental statistical prediction of ARI beyond those predictors by entering PAF after the initial predictor block. H2, H3, and H5 were secondary explanatory analyses involving perceived legitimacy; H5 used 5,000 bootstrap samples and a conservative model controlling for the UTAUT-related predictors (Hayes, 2022). H6 tested the interaction between measured perceived architectural fit and assigned accountability. Heteroskedasticity-consistent type 3 (HC3) standard errors and influence diagnostics were reviewed for focal models. Because the post-vignette perceptions were measured concurrently, H2, H3, and H5 are interpreted as associations rather than causal paths (Podsakoff et al., 2003).

5.0 Results

5.1 Screening and Sample

The MTurk batch produced 350 unique responses. Sequential exclusions removed 6 attention-check failures, 1 respondent without current or recent employment, 41 respondents without workplace AI familiarity, and 20 responses below the speed cutoff. The final sample included 282 respondents. Cell sizes were 71 for low accountability/limited support, 75 for low accountability/enhanced support, 66 for high accountability/limited support, and 70 for high accountability/enhanced support. All cells exceeded the planned minimum.

In the analytic sample, 3.9% reported workplace AI use rarely, 25.2% sometimes, 55.0% often, and 16.0% very often. Age, gender, education, occupation, and industry were not collected. The sample should therefore be treated as a screened nonprobability sample of U.S.-restricted workers with self-reported workplace AI exposure, not as a population estimate.

5.2 Descriptive Statistics and Measurement Checks

Table 5 reports means, standard deviations, internal consistency, and correlations. Mean ratings were near the agreement range, which may have constrained condition effects. Reliability was acceptable for perceived architectural fit ($\alpha = .743$) and perceived legitimacy ($\alpha = .704$), and near the conventional threshold for ARI ($\alpha = .682$). Reliability for the UTAUT scales was lower ($\alpha = .514$ to $.666$); the UTAUT block is therefore interpreted as a scenario-specific benchmark whose individual coefficients may be attenuated by measurement error. The architectural-support recognition check had good reliability ($\alpha = .838$), whereas the accountability check had low two-item consistency ($\alpha = .359$).

Table 5. Descriptive Statistics, Internal Consistency, and Correlations

Variable	M	SD	α	1	2	3	4	5	6	7
1. PE	4.042	.470	.514	1.000						
2. EE	4.055	.424	.614	.515	1.000					
3. SI	4.085	.421	.583	.513	.607	1.000				
4. FC	4.059	.472	.666	.449	.568	.582	1.000			
5. PAF	3.940	.679	.743	.341	.408	.414	.693	1.000		
6. PL	4.011	.617	.704	.393	.448	.491	.696	.802	1.000	
7. ARI	4.056	.520	.682	.455	.525	.582	.746	.750	.778	1.000

Note. $N = 282$. PE = performance expectancy; EE = effort expectancy; SI = social influence; FC = facilitating conditions; PAF = perceived architectural fit; PL = perceived legitimacy; ARI = AI reliance intention. All reported correlations were positive and statistically significant, $p < .001$.

Perceived architectural fit correlated strongly with perceived legitimacy ($r = .802$) and ARI ($r = .750$); perceived legitimacy correlated strongly with ARI ($r = .778$). These associations fit the proposed model but also raise common-method and construct-overlap concerns. The six-item component check produced one dominant unrotated component with an eigenvalue of 3.504, explaining 58.4% of item variance. The second component had an eigenvalue of .786 and explained 13.1%. Only the first eigenvalue exceeded 1.0, indicating that PAF and perceived legitimacy were empirically close in this preliminary analysis. Additional research using formal discriminant-validity tests is needed to determine whether they are distinct constructs.

5.3 Manipulation Checks and Assigned-Condition Effects

The accountability check was higher in the high-accountability condition ($M = 4.188$, $SD = 0.467$, $n = 136$) than in the low-accountability condition ($M = 4.055$, $SD = 0.522$, $n = 146$), Welch's $t(279.5) = 2.252$, $p = .025$, $d = 0.268$, 95% CI [0.033, 0.502]. This difference was small. The architectural-support recognition check was higher in the enhanced-support condition ($M = 4.148$, $SD = 0.755$, $n = 145$) than in the limited-support condition ($M = 3.836$, $SD = 1.054$, $n = 137$), Welch's $t(245.4) = 2.848$, $p = .005$, $d = 0.341$, 95% CI [0.105, 0.576].

Participants in the enhanced architectural-support condition reported higher perceived architectural fit than those in the limited-support condition, $F(1, 278) = 5.503$, $p = .020$, partial $\eta^2 = .019$; perceived legitimacy and ARI did not differ significantly. The accountability conditions differed on the manipulation check, $F(1, 278) = 5.001$, $p = .026$, partial $\eta^2 = .018$, but not on the focal outcomes. Tables 6 and 7 report the complete results. The assigned-condition evidence therefore supports recognition of the architectural-support distinction, whereas the principal incremental-prediction finding comes from the hierarchical measured-perception model.

Table 6. Four-Cell Means and Standard Deviations

Variable	HH (n = 70)	HL (n = 66)	LH (n = 75)	LL (n = 71)
Perceived legitimacy	4.132 (.509)	4.015 (.727)	4.013 (.604)	3.887 (.607)
ARI	4.209 (.421)	4.014 (.631)	4.018 (.511)	3.985 (.483)
Accountability check	4.171 (.450)	4.205 (.488)	4.040 (.568)	4.070 (.473)
Architectural-support recognition	4.243 (.588)	3.947 (1.053)	4.060 (.877)	3.732 (1.052)
Perceived architectural fit	4.128 (.437)	3.888 (.787)	3.942 (.654)	3.802 (.760)

Note. HH = high accountability/enhanced architectural support; HL = high accountability/limited architectural support; LH = low accountability/enhanced architectural support; LL = low accountability/limited architectural support. Values are M (SD).

Table 7. Assigned-Condition Effects

Outcome	Accountability main effect F, p, partial η^2	Architectural- support main effect F, p, partial η^2	Interaction F, p, partial η^2
Accountability check	5.001, .026, .018	.286, .593, .001	.001, .982, < .001
Architectural-support recognition	3.341, .069, .012	8.293, .004, .029	.021, .884, < .001
Perceived architectural fit	2.959, .087, .011	5.503, .020, .019	.389, .534, .001
Perceived legitimacy	2.855, .092, .010	2.782, .097, .010	.003, .956, < .001
ARI	3.355, .068, .012	3.255, .072, .012	1.750, .187, .006

5.4 UTAUT-Related Baseline and Incremental-Prediction Test

The UTAUT-related block was significant at the model level, $F(4, 277) = 102.70$, $p < .001$, $R^2 = .597$, adjusted $R^2 = .591$. Social influence and facilitating conditions had significant unique coefficients, whereas performance expectancy and effort expectancy did not in the simultaneous model.

H4 was the principal test of whether PAF provides incremental statistical prediction of ARI beyond the study's four UTAUT-related predictors. Adding perceived architectural fit increased explained variance from .597 to .698, $\Delta R^2 = .100$, $\Delta F(1, 276) = 91.60$, $p < .001$. Perceived architectural fit remained significant in the full model, $B = 0.337$, $SE = 0.035$, $\beta = .440$, 95% CI [0.268, 0.406], $p < .001$. Table 8 reports the complete Model 1 and Model 2 coefficients.

The H2, H3, and H5 results were consistent with the proposed explanatory pattern. In the UTAUT-adjusted H2 model, perceived architectural fit was associated with perceived legitimacy, $B = 0.554$, $SE = 0.042$, $\beta = .609$, $p < .001$. Perceived legitimacy was associated with ARI, $B = 0.656$, $SE = 0.032$, $\beta = .778$, $p < .001$. The conservative indirect association was $B = 0.145$, $SE = 0.028$, with a 95% bootstrap CI of [0.090, 0.201]. Because perceived architectural fit, perceived legitimacy, and reliance intention were measured concurrently using the same response method, these analyses establish statistical associations rather than temporal ordering or causal mediation. Because PAF and perceived legitimacy were not entered simultaneously in a reported ARI regression model, the present study does not establish their relative predictive strength.

H6 was not supported. The interaction between perceived architectural fit and accountability condition did not reach significance, $B = -0.116$, $SE = 0.066$, $p = .078$, 95% CI [-0.246, 0.013]. Table 8 presents the principal UTAUT-based incremental-prediction model.

Table 8. Hierarchical Regression Predicting ARI

Model	Predictor	B	SE	β	95% CI	p
1	PE	0.088	0.052	0.080	[-0.014, 0.190]	0.091
1	EE	0.058	0.064	0.047	[-0.069, 0.185]	0.367
1	SI	0.215	0.065	0.174	[0.086, 0.344]	0.001
1	FC	0.642	0.055	0.582	[0.533, 0.751]	< .001
2	PE	0.072	0.045	0.065	[-0.017, 0.161]	0.114

Model	Predictor	B	SE	β	95% CI	p
2	EE	0.055	0.056	0.045	[-0.055, 0.165]	0.329
2	SI	0.214	0.057	0.173	[0.102, 0.326]	< .001
2	FC	0.315	0.059	0.286	[0.199, 0.431]	< .001
2	PAF	0.337	0.035	0.440	[0.268, 0.406]	< .001

Note. $N = 282$. PE = performance expectancy; EE = effort expectancy; SI = social influence; FC = facilitating conditions; PAF = perceived architectural fit. Model 1: $R^2 = .597$, adjusted $R^2 = .591$. Model 2: $R^2 = .698$, adjusted $R^2 = .692$, $\Delta R^2 = .100$, $\Delta F(1, 276) = 91.60$, $p < .001$. HC3 standard errors produced the same substantive conclusion for perceived architectural fit ($SE = 0.053$, $p < .001$).

5.5 Sensitivity and Diagnostic Results

Reduced-legitimacy sensitivity analyses produced the same general explanatory pattern as shown in Table 9. Perceived architectural fit also remained significant in separate UTAUT-adjusted models for intention to use ($B = 0.271$, $p < .001$), support for deployment ($B = 0.396$, $p < .001$), and willingness to rely ($B = 0.341$, $p < .001$). Variance inflation factors in the focal noninteraction models ranged from 1.52 to 2.60. HC3 standard errors did not change the substantive conclusion for perceived architectural fit.

Table 9. Reduced-Legitimacy Sensitivity Results

Analysis	Focal estimate	95% CI	p
Adjusted PAF → reduced PL	$B = 0.381$, $\beta = .435$	[0.286, 0.477]	< .001
Reduced PL → ARI	$B = 0.650$, $\beta = .743$	[0.581, 0.719]	< .001
Indirect PAF → reduced PL → ARI	$B = 0.226$	[0.170, 0.281]	< .001

Note. PAF = perceived architectural fit; PL = perceived legitimacy; ARI = AI reliance intention. The reduced two-item legitimacy measure had $\alpha = .572$ and was used only as a sensitivity check.

6.0 Discussion

The principal finding is that perceived architectural fit added significant explanatory value to a UTAUT-based model of workplace ARI. The UTAUT-related block explained 59.7% of the variance, and adding perceived architectural fit increased explained variance to 69.8%. The incremental result remained significant with HC3 robust standard errors, and perceived architectural fit was significant in separate models for the intention-to-use, deployment-support, and willingness-to-rely items. The findings provide quantitative evidence that perceived architectural fit may contribute explanatory value to UTAUT-based research on workplace ARI. The assigned-condition results support participants' recognition of the architectural-support descriptions, not direct causal effects on reliance intention.

6.1 Theoretical Implications

The results support the exploratory addition of PAF to the study's abbreviated UTAUT-based benchmark, not a theoretical or operational extension of UTAUT. The four UTAUT-related predictors jointly explained substantial variance in ARI. Social influence and facilitating conditions made significant unique contributions, whereas performance expectancy and effort expectancy did not contribute significant unique variance in the simultaneous model. These individual coefficients should be interpreted in light of shared predictor variance, high mean ratings, and the lower reliability of the brief scenario-specific measures.

Perceived architectural fit adds a context-sensitive question to this benchmark: whether an AI system's technical and governance supports are suitable for the accountability demands of its intended use. Its addition increased explained variance by ten percentage points after the UTAUT-related predictors were entered. The principal incremental-prediction conclusion

does not depend on perceived legitimacy because legitimacy was not included in the H4 hierarchical model.

6.2 Perceived Legitimacy as a Secondary Explanatory Pattern

Perceived legitimacy offers one possible statistical explanation for the relationship between perceived architectural fit and reliance intention. The evidence does not establish that legitimacy is the mechanism through which architectural fit is related to reliance. The two perceptions were strongly correlated; one legitimacy item overlapped with reviewability; and the component check showed substantial empirical proximity. The reduced-item sensitivity analysis indicates that one overlapping item did not account for the entire pattern, but the study does not claim discriminant validity. Later studies should compare perceived architectural fit with task-technology fit, trust, risk, control, transparency, explainability, and facilitating conditions using formal validity tests (Henseler et al., 2015).

6.3 Practical Implications

Organizations evaluating workplace AI should assess whether prospective users perceive the system's review, record, uncertainty, and governance supports as suitable for the accountability demands of the intended task. The UTAUT-related results also support attention to organizational guidance, resources, and social support.

The study does not show that organizations must implement any particular feature. The vignette bundled source trails, reviewable records, record separation, uncertainty communication, human review, and local governance configuration, so it cannot determine which feature was associated with the observed fit ratings. Participants in the enhanced-support condition reported higher PAF, but the assigned architectural-support conditions did not produce statistically significant differences in perceived legitimacy or ARI. Thus, the experimental results support the first link in the proposed pattern but not direct effects of the condition on the two downstream outcomes.

7.0 Conclusion and Future Research

The UTAUT-related block explained a substantial amount of variance in workplace ARI. Perceived architectural fit added ten percentage points of explained variance beyond that block and remained a significant predictor in the full model.

The enhanced architectural-support condition was associated with higher perceived architectural fit, but perceived legitimacy and ARI did not differ significantly between conditions. Associations involving perceived legitimacy were consistent with the proposed secondary pattern, but they do not establish causal mediation or a validated mechanism.

These findings justify continued study of perceived architectural fit as a context-sensitive predictor in workplace AI adoption research, while fuller measurement validation and behavioral testing remain tasks for later research.

Some of the research limitations include the measurement of perceived architectural fit, perceived legitimacy, and reliance intention concurrently from the same respondents using similarly valenced items. Common-method variance, semantic overlap, and general system favorability may therefore inflate the secondary associations, and the indirect analysis does not establish temporal ordering or a causal mechanism.

The four UTAUT-related measures were brief scenario-specific adaptations, and several had lower-than-desired internal consistency. Perceived architectural fit was measured with three global items and should be treated as a preliminary construct. The study does not establish full-scale validation, seven empirical subscales, or discriminant validity from perceived legitimacy and other related constructs.

The vignette bundled several architectural and governance features and assessed intention rather than behavior. The enhanced-support description was more specific and reassuring than the limited-support description, and the study cannot identify which feature was

associated with the perceived-fit difference. The accountability manipulation produced only a small between-condition difference, and its two-item check had low internal consistency ($\alpha = .359$); H6 should therefore be interpreted with caution. Future research should use feature-specific manipulations and behavioral outcomes such as verification, challenge, escalation, and calibrated reliance.

The sample was a nonprobability MTurk sample restricted by platform settings to the United States and screened through broad self-reported AI exposure. Demographic and occupational information was not collected, which may limit generalizability.

Ethics Statement

The Indiana Institute of Technology Institutional Review Board (IRB) approved the minimal-risk study. Participants provided electronic informed consent. MTurk Worker IDs were used only for compensation, duplicate prevention, and data-quality screening and were removed before analysis.

Data Availability

De-identified data may be provided upon reasonable request, provided it is consistent with participant privacy, IRB requirements, and MTurk policies.

Conflict of Interest

The authors declare no financial or nonfinancial conflicts of interest related to this research.

Generative AI Use Disclosure

OpenAI's ChatGPT was used to provide editorial feedback on structure, clarity, terminology, presentation, and the boundaries of claims. The authors independently evaluated and implemented all revisions. It was not used to collect data, perform statistical analyses, or generate findings. The authors remain responsible for the manuscript's content and accuracy.

References

- [1] Aguinis, H., & Bradley, K. J. (2014). Best practice recommendations for designing and implementing experimental vignette methodology studies. *Organizational Research Methods*, 17(4), 351–371. <https://doi.org/10.1177/1094428114547952>
- [2] Autio, C., Schwartz, R., Dunietz, J., Jain, S., Stanley, M., Tabassi, E., Hall, P., & Roberts, K. (2024). Artificial intelligence risk management framework: Generative artificial intelligence profile (NIST AI 600-1). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.AI.600-1>
- [3] Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., Buch, S., Card, D., Castellon, R., Chatterji, N., Chen, A., Creel, K., Davis, J. Q., Demszky, D., ... Liang, P. (2021). On the opportunities and risks of foundation models. *arXiv*. <https://arxiv.org/abs/2108.07258>
- [4] Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Routledge. <https://doi.org/10.4324/9780203771587>
- [5] European Parliament & Council of the European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). *Official Journal of the European Union*, L 2024/1689. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- [6] Farquhar, S., Kossen, J., Kuhn, L., & Gal, Y. (2024). Detecting hallucinations in large language models using semantic entropy. *Nature*, 630, 625–630. <https://doi.org/10.1038/s41586-024-07421-0>
- [7] Goodhue, D. L., & Thompson, R. L. (1995). Task-technology fit and individual performance. *MIS Quarterly*, 19(2), 213–236. <https://doi.org/10.2307/249689>

- [8] Hayes, A. F. (2022). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (3rd ed.). Guilford Press.
- [9] Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43, 115–135. <https://doi.org/10.1007/s11747-014-0403-8>
- [10] Hoff, K. A., & Bashir, M. (2015). Trust in automation: Integrating empirical evidence on factors that influence trust. *Human Factors*, 57(3), 407–434. <https://doi.org/10.1177/0018720814547570>
- [11] Huynh, M.-T., & Aichner, T. (2025). In generative artificial intelligence we trust: Unpacking determinants and outcomes for cognitive trust. *AI & Society*, 40, 5849–5869. <https://doi.org/10.1007/s00146-025-02378-8>
- [12] JASP Team. (2025). JASP (Version 0.19.3) [Computer software]. <https://jasp-stats.org>
- [13] Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392
- [14] Liu, N. F., Lin, K., Hewitt, J., Paranjape, A., Bevilacqua, M., Petroni, F., & Liang, P. (2024). Lost in the middle: How language models use long contexts. *Transactions of the Association for Computational Linguistics*, 12, 157–173. https://doi.org/10.1162/tacl_a_00638
- [15] Malhotra, N. (2008). Completion time and response order effects in web surveys. *Public Opinion Quarterly*, 72(5), 914–934. <https://doi.org/10.1093/poq/nfn050>
- [16] National Institute of Standards and Technology. (2023). Artificial intelligence risk management framework (AI RMF 1.0) (NIST AI 100-1). <https://doi.org/10.6028/NIST.AI.100-1>
- [17] Parasuraman, R., & Riley, V. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39(2), 230–253. <https://doi.org/10.1518/001872097778543886>
- [18] 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. <https://doi.org/10.1037/0021-9010.88.5.879>
- [19] Raji, I. D., Smart, A., White, R. N., Mitchell, M., Gebru, T., Hutchinson, B., Smith-Loud, J., Theron, D., & Barnes, P. (2020). Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency*, 33–44. <https://doi.org/10.1145/3351095.3372873>
- [20] Scott, W. R. (2014). *Institutions and organizations: Ideas, interests, and identities* (4th ed.). Sage.
- [21] Suchman, M. C. (1995). Managing legitimacy: Strategic and institutional approaches. *Academy of Management Review*, 20(3), 571–610. <https://doi.org/10.2307/258788>
- [22] Thomas, K. A., & Clifford, S. (2017). Validity and Mechanical Turk: An assessment of exclusion methods and interactive experiments. *Computers in Human Behavior*, 77, 184–197. <https://doi.org/10.1016/j.chb.2017.08.038>
- [23] 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>
- [24] Wang, X., Zhong, W., Huang, K., & Liang, B. (2026). High interest but low adoption: Navigating organizations' journey towards generative artificial intelligence implementation. *International Journal of Information Management*, 87, Article 103009. <https://doi.org/10.1016/j.ijinfomgt.2025.103009>
- [25] Xia, Y., & Chen, Y. (2025). Driving factors of generative AI adoption in new product development teams from a UTAUT perspective. *International Journal of Human–Computer Interaction*, 41(10), 6067–6088. <https://doi.org/10.1080/10447318.2024.2375686>