



Effects of a Structured AI Integration Framework on Faculty Technology Acceptance and Adoption in Higher Education: A Quasi-Experimental Study

Oana-Adriana Ticleanu^{1*} and Nicolae Constantinescu²

^{1,2} Lucian Blaga University of Sibiu, No. 10, Bld. Victoriei, Sibiu, ROMANIA
oana.ticleanu@ulbsibiu.ro*, nicolae.constantinescu@ulbsibiu.ro

Abstract. We ran a quasi-experimental pre-post study, with a four-week follow-up, to examine the Integration Framework for Artificial Intelligence in Higher Education (IFAI-HE). The program is built around six modules and draws on Technology Acceptance Model 3, Rogers's Diffusion of Innovations, the Concerns-Based Adoption Model and Self-Determination Theory. Our motivation was straightforward: most prior work on AI uptake in academia has been cross-sectional or correlational, with little controlled evidence. Across three Romanian universities, propensity score matching produced a treated group (n = 60) and a control group (n = 60) balanced on age, years of teaching, discipline and rank. Treated faculty followed an eight-week fully online program covering awareness-building, self-diagnosis, sandbox practice, course redesign, peer-community formation and reflective synthesis. Controls received only standard institutional policy documents and one introductory webinar. Outcomes were measured at three time points using validated TAM3-AI scales, the Stages of Concern Questionnaire and an AI Adoption Behaviour Scale. Mixed-effects models showed group-by-time interactions on every TAM3 dimension, with large effect sizes (partial eta-squared = 0.12 to 0.18). Self-focused concerns dropped clearly (Personal stage, $d = -1.56$) while impact-focused concerns rose (Collaboration stage, $d = +1.67$). Gains held at follow-up, with no sign of decay. In moderation analyses, early-career STEM faculty at research-intensive universities showed the largest changes. We tentatively interpret these data as the first controlled experimental evidence that a multi-theoretical faculty development program can shift AI adoption indicators in higher education, and that the changes appear to persist beyond the immediate post-test.

Keywords: Artificial Intelligence In Higher Education, Faculty Development, Technology Acceptance, The Technology Acceptance Model 3, The Concerns-Based Adoption Model, The Diffusion Of Innovations, Self-Determination Theory, Quasi-Experimental Design, Propensity Score Matching, Mixed-Effects Modeling, Stages Of Concern, Professional Development Intervention, AI Incorporation, Educational Technology Adoption, Behavioral Intention, Perceived Usefulness, Perceived Ease Of Use, Job Relevance, Subjective Norm, Intervention Fidelity, Faculty Attitudes, Digital Transformation, Higher Education Pedagogy, Faculty Training, Technology Adoption Behavior, Intervention Sustainability.

© The Author(s) 2026

D. Mara and I. D. Hunyadi (eds.), *Proceedings of the International Conference on Management and Entrepreneurial Leadership for K-12 Education Excellence (LEADK12 2026)*, Advances in Social Science, Education and Humanities Research 1026,

https://doi.org/10.2991/978-2-38476-599-7_11

1 Introduction

Generative artificial intelligence (AI) has spread quickly across higher education. The institutional question is no longer whether to adopt these tools, but how to help faculty use them responsibly and well [1]. Vendor sessions, online tutorials and policy briefs are widely available. Even so, few intervention studies in this area are theory-driven and properly controlled, and even fewer test causal links between structured adoption frameworks and observable change in faculty behaviour [2]. Most existing work is descriptive: cross-sectional surveys of faculty opinions about AI [3], correlational designs predicting self-stated intention, or single-institution case reports. These designs offer limited internal validity for evidence-based policy [4]. The gap matters because AI tools are entering classrooms faster than development programs can be evaluated [5].

Our study attempts to fill part of that gap. We evaluated a structured, multi-theoretical faculty development program with a quasi-experimental design. The program — the Integration Framework for Artificial Intelligence in Higher Education (IFAI-HE) — runs over eight weeks across six modules and rests on four theories: the Technology Acceptance Model 3 (TAM3), Rogers's Diffusion of Innovations, the Concerns-Based Adoption Model (CBAM) and Self-Determination Theory (SDT). We asked whether participants exposed to this program would show measurable gains on the TAM3 acceptance dimensions, the predicted CBAM concern shifts, and higher self-reported AI adoption, compared with a non-equivalent control group that received only standard institutional guidance.

The work has three goals. First, we provide controlled intervention-based evidence that a modular, multi-theoretical faculty development program can shift AI adoption indicators in higher education. We then assess durability through a four-week follow-up — addressing a gap left by earlier cross-sectional and short-term studies [6]. Lastly, we examine moderation by institutional type, discipline and career stage, in the hope of identifying which faculty subgroups benefit most from structured AI support.

We hypothesised that participants in the IFAI-HE condition would show larger gains on TAM3 dimensions and the typical CBAM progression from self-focused toward impact-focused concerns, relative to controls. We expected these effects to hold at four-week follow-up — suggesting genuine integration of AI into pedagogical practice rather than novelty or social desirability.

2 Theoretical Background

Faculty technology adoption is shaped both by individual cognition and by broader organisational diffusion. The Technology Acceptance Model (TAM) is one of the most widely applied frameworks for educational settings; it was first proposed by Davis [7], [8] and later extended into TAM2 and TAM3 by Venkatesh and Bala [9]. Under TAM3, perceived usefulness and perceived ease of use predict behavioural intention, and both perceptions are shaped by external factors — subjective norm, job relevance, output quality, result demonstrability, computer self-efficacy and computer anxiety. In higher education, faculty acceptance is largely predicted by perceived usefulness (does using

the tool improve my teaching?) and perceived ease of use (how much effort to become competent?) [10].

TAM, useful as it is, focuses on the individual. It does not capture the social and temporal dynamics of how innovations such as generative AI spread through faculty bodies. Rogers's Diffusion of Innovations theory [11] frames adoption as a process unfolding inside a social system over time, with adopters ranged from innovators to laggards by their relative timing. The framework points to communication channels, social norms and the perceived attributes of innovations (relative advantage, compatibility, complexity, trialability, observability) as drivers of adoption. For AI specifically, faculty are influenced not only by their own views but also by what colleagues in their department do, and by institutional signals about commitment to the innovation [12].

A third strand of theory concerns the affective and developmental stages individuals pass through when they engage with educational innovations. Hall and Hord's Concerns-Based Adoption Model (CBAM) [13] describes seven stages of concern — Awareness, Informational, Personal, Management, Consequence, Collaboration and Refocusing. People first show self-focused concerns (Awareness, Informational, Personal), then move to task-focused (Management), and only later turn toward impact-focused concerns (Consequence, Collaboration, refocusing) as experience builds [14]. The model fits AI adoption well: faculty often start by worrying about workload, their own competence and possible job displacement before they begin to think about pedagogical refinement and joint innovation.

Self-Determination Theory (SDT) adds a motivational layer. Intrinsic motivation and durable engagement, in this view, depend on three basic needs — autonomy, competence and relatedness [15]. In professional development, programs that give faculty real choice over how and when they adopt AI, that scaffold their skill-building, and that connect them with peers facing the same problems are more likely to produce lasting change than programs that simply transfer information [16]. SDT thus offers a motivational rationale for why theory-driven faculty development can outperform pure information delivery.

Each of these frameworks has been applied separately, mostly through descriptive or correlational designs. Cross-sectional surveys identify TAM predictors of AI adoption intention among faculty [17]; qualitative work has tracked stages of concern in early adopters of learning analytics [18]. Few studies, however, deliberately combine multiple frameworks into a structured professional development program and then test causal effects with a controlled design. Our study takes that step: TAM3, Diffusion of Innovations, CBAM and SDT are integrated into a six-module intervention, then assessed in a quasi-experimental pre-post design with propensity score matching.

3 The IFAI-HE Intervention

The IFAI-HE program runs over eight weeks online and operationalises four theoretical models across six sequential modules. Each module targets specific TAM3, CBAM, SDT and Diffusion-of-Innovations constructs, taking participants from awareness-building through reflective synthesis.

Module 1 — Awareness and Orientation (Week 1). Participants are introduced to the landscape of generative AI tools relevant to teaching. The module maps onto the Awareness stage of CBAM through foundational content on AI capabilities, limitations and ethics. Institutional case studies of peer adoption activate the subjective-norm construct of TAM3, drawing on the social-influence mechanisms of Diffusion of Innovations. A short guided self-assessment of current AI knowledge and use serves as a baseline for personal goal-setting.

Module 2 — Self-Diagnosis and Goal Setting (Week 2). Faculty complete the Stages of Concern Questionnaire and a TAM3 self-assessment, then locate themselves on the CBAM map. Grounded in SDT's autonomy need, the module lets participants set their own developmental starting point and pick a personalised learning trajectory. A job-relevance mapping links potential AI uses to disciplinary teaching contexts, operationalising the TAM3 job-relevance construct. The shift from Awareness to Informational concerns is supported by anchoring information-seeking around problems the participant actually cares about.

Module 3 — Hands-On Sandbox Practice (Weeks 3-4). This is the most intensive part of the program. Participants work with three categories of AI tools — text generation (e.g., ChatGPT, Claude), image generation, and AI-supported assessment and feedback. Each tool is introduced via a demonstration video, followed by a guided sandbox task: writing quiz items, drafting rubric criteria, or generating alternative explanations for difficult concepts. The hands-on work is designed to act on TAM3 perceived ease of use, while the trialability attribute from Diffusion of Innovations is also engaged. Tasks scale from simple to complex so that the SDT competence need is met through mastery experiences before harder activities. A structured reflective journal records what worked and what did not.

Module 4 — Course Redesign Studio (Weeks 5-6). Faculty are asked to redesign one specific learning activity or assessment in their current course to incorporate AI. The module operationalises the result-demonstrability construct from TAM3: participants must produce tangible artefacts that show AI-enhanced pedagogical outcomes. Submissions are written redesign proposals (objectives, AI tool specifications, implementation timeline, anticipated student impact). Small-group peer critique, run as structured sessions, addresses SDT's relatedness need and supports the move from Personal to Management concerns. Each participant receives feedback from two peer reviewers and one facilitator and revises before moving on.

Module 5 — Peer Community Formation (Week 7). Communities of five to seven participants are built around discipline and career stage. They meet synchronously by video for a two-hour session led by a trained moderator. The protocol focuses on implementation challenges, exchange of strategies that worked, and joint problem-solving. The Collaboration stage of CBAM is the explicit focus — institutional structures for collective innovation — and the subjective-norm and social-influence mechanisms of Diffusion of Innovations are again at play. We encouraged communities to set up their own communication channels for after the program, and many created messaging groups or scheduled follow-up meetings on their own.

Module 6 — Reflective Synthesis and Forward Planning (Week 8). The intervention closes with a structured reflection across the previous seven weeks. Participants repeat

the TAM3 self-assessment and the Stages of Concern Questionnaire, then compare profiles with the Module 2 baseline. A forward-planning instrument asks them to set specific AI adoption goals for the next academic term, list expected obstacles and supports, and write down concrete next actions. The module is built around CBAM's Refocusing stage — how can my AI use be improved further? — while reinforcing perceived behavioural control. Facilitators provide written feedback on each plan, paying particular attention to fit with institutional resources and existing supports.

Most of the program runs asynchronously through a learning management system, with weekly module deadlines and feedback delivered within 48 hours. Synchronous time is limited to the Module 5 community session and two optional weekly office hours. We chose this hybrid-asynchronous format to keep the program accessible to faculty with heavy teaching and research loads, while keeping enough social interaction to satisfy SDT's relatedness need and to support CBAM's Collaboration stage. Materials were developed with discipline-specific advisory panels to keep them relevant across STEM, social sciences, humanities and professional programs.

4 Methods

4.1 Study Design and Participants

The study used a quasi-experimental nonequivalent-groups pretest-posttest design with a delayed follow-up at four weeks post-intervention, in order to gauge the effects of the IFAI-HE framework on faculty technology acceptance, stages of concern, and adoption behaviour. Random assignment was not feasible — faculty self-selected into the programme based on availability and interest — so we constructed a comparison group through propensity score matching to obtain baseline equivalence between conditions. Recruitment was carried out at three Romanian universities chosen for institutional diversity: one large public research university, one mid-sized public university, and one private institution. The arrangement broadens external validity while keeping the national policy context homogeneous enough to limit extraneous variation.

Recruitment yielded 132 faculty members in total: 68 in the treatment condition and 64 in the comparison group. Eligibility required an active teaching role of at least one year, no previous exposure to a structured AI training programme, and willingness to attend all three measurement points across the twelve-week study period. We then matched cases via nearest-neighbour propensity score matching, with a caliper of 0.2 standard deviations on five covariates: age; prior AI exposure (rated on a five-point scale running from no exposure to frequent application); years of teaching experience; academic discipline (STEM, social sciences, humanities, or professional programmes); and academic rank (assistant, associate, or full professor). The procedure produced 60 matched pairs per group and standardised mean differences below 0.03 across all covariates, leaving baseline equivalence in good shape. The matching itself was run with the MatchIt package in R, with replacement enabled to keep optimal match quality without shrinking the analytic sample below what the subsequent mixed-effects models required.

4.2 The IFAI-HE Intervention

Module one introduced and demystified generative AI, working with the Awareness and Informational stages of CBAM and the subjective-norm and perceived-usefulness constructs of TAM3. Module two asked faculty to self-diagnose their AI literacy, anxiety and prior practice — fitting with CBAM's Personal stage and SDT's competence need through structured self-assessment. Module three offered hands-on sandbox practice across writing assistance, assessment design and learning analytics, addressing CBAM's Management stage, SDT's competence need and TAM3 perceived ease of use. Module four required participants to redesign one course element using AI, focused on CBAM's Consequence stage and TAM3 result-demonstrability. Module five built a community of practice for peer feedback and case exchange (CBAM Collaboration; SDT relatedness). Module six closed with a reflective-synthesis task in which participants drafted a personal consolidation plan, working on CBAM Refocusing and SDT autonomy.

Delivery sat on three parts. Asynchronous micro-learning units, each fifteen to twenty minutes long, carried the bulk of content. Weekly ninety-minute synchronous live sessions anchored the cohort. A steady backchannel of peer-cohort communication ran through the learning management system, and participants had access to an AI tool sandbox covered by institutional licensing. We monitored fidelity in three ways: attendance logs requiring at least seventy-five percent module completion; weekly artefact submissions — reflection journals and redesign proposals — and post-module satisfaction ratings. Faculty in the comparison condition received the standard institutional AI guidance: a twelve-page policy document distributed by email, plus a single sixty-minute introductory webinar covering basic AI concepts and institutional guidelines.

4.3 Instruments and Outcome Measurement

We administered three instruments at baseline (T1), immediately post-intervention (T2), and at the four-week follow-up (T3). The TAM3-AI scale (32 items adapted from the Technology Acceptance Model 3) covered nine dimensions—Perceived Usefulness, Perceived Ease of Use, Subjective Norm, Job Relevance, Output Quality, Result Demonstrability, Computer Self-Efficacy, Computer Anxiety, and Behavioral Intention—on seven-point Likert scales (strongly disagree to strongly agree). Earlier higher-education work has reported Cronbach's alpha above 0.85 for every subscale; in our data, alpha ranged from 0.82 to 0.93 across measurement occasions. The SoCQ, with 35 items, mapped onto the seven CBAM stages on an eight-point scale (irrelevant to very true of me); test-retest coefficients above 0.70 and confirmatory factor analysis support its psychometric soundness. We also used a 12-item AI Adoption Behaviour Scale developed in-house, which captured how often and how deeply faculty drew on AI for course preparation, instructional delivery, assessment, feedback, and administrative duties on a five-point frequency scale (never to very frequently). Demographic and contextual data—age, gender, years of teaching, discipline, rank, prior AI exposure, and institutional affiliation—were collected at baseline.

4.4 Statistical Analysis

Our primary analytic engine was a mixed-effects ANOVA cast as a 2×3 factorial — treatment versus comparison crossed with the three measurement waves (T1, T2, T3) — with random intercepts at the participant level and random slopes on time so individual change paths could vary. We preferred this over a classical repeated-measures ANOVA on three grounds: tolerance for missing data under MAR assumptions, accommodation of heterogeneous change rates, and tighter standard errors when intraclass correlation is non-trivial. Secondary work included ANCOVA models in which baseline (T1) scores anchored the post-test outcomes (giving group differences adjusted for starting position), partial η^2 and Hedges' g effect sizes with 95% confidence intervals, and a simple-effects breakdown of every reliable interaction via estimated marginal means with Bonferroni adjustment for multiplicity.

For sensitivity, we ran multiple imputation with five imputations on the missing follow-up data via the *Amelia* package in R, compared intent-to-treat against per-protocol estimates to gauge the role of non-compliance and attrition, and bootstrapped 10,000 iterations to obtain confidence intervals that do not lean on normality. Moderation was tested through three-way interactions for institutional type (research vs. teaching-focused), discipline (STEM vs. non-STEM), and career stage (early, mid, senior); given the limited power for higher-order interactions, those readings stay exploratory. Alpha was set at 0.05 with Holm correction across the primary outcome family (the nine TAM3 subscales). Ethics approval came from each of the three participating universities. Informed consent was collected after full disclosure of procedures and data handling, the right to withdraw without penalty was made explicit, and anonymisation removed direct identifiers and assigned study codes. Data sat on encrypted institutional servers accessible only to the research team, and findings were shared with participants once data collection and analysis were finished.

5 Results

The results below follow the order of the analyses themselves. We start with participant flow and baseline equivalence, then move to the primary tests of technology acceptance and stages of concern, report sustainability and moderation outcomes, and close with a synthesis of effect sizes and sensitivity checks. All statistical tests were judged at a significance threshold of $\alpha = 0.05$ with Holm correction applied for the primary outcome family, and effect sizes are reported with ninety-five percent confidence intervals throughout.

5.1 Sample and Baseline Characteristics

We recruited 132 faculty members from three Romanian universities, with 68 placed in the treatment arm and 64 in the control arm before propensity score matching. After applying nearest-neighbour matching with a caliper of 0.2 standard deviations on the five covariates (age, prior AI exposure, years of teaching experience, academic discipline, and academic rank), the analytic sample reduced to 60 matched pairs per group,

or 120 participants in the primary analyses. Baseline equivalence looked good: standardised mean differences (SMD) stayed below 0.03 on every covariate, which is well inside the threshold typically required for quasi-experimental causal inference. Attrition was low and roughly even across conditions. At the four-week follow-up (T3), 52 participants in the intervention group and 55 in the comparison group completed all three measurement points, giving retention rates of 86.7% and 91.7%, respectively. We saw no differential dropout pattern by condition; a logistic regression of attrition on group assignment confirmed this ($\beta = 0.31$, $p = 0.42$).

Ranks split into 38.3% assistant professors, 41.7% associate, and 20.0% full. By discipline, 40.0% sat in STEM, 28.3% in social sciences, 21.7% in humanities, and 10.0% in professional programs such as business and law. Institutional affiliation broke down as 41.7% from the large public research university, 35.0% from the mid-sized public university, and 23.3% from the private institution. Baseline AI exposure was low in both arms ($M = 1.82$ on a five-point scale, $SD = 0.91$), which confirms that the eligibility criterion of no prior structured AI training held in practice.

On the TAM3-AI scale, both arms scored at moderate levels across the nine dimensions at baseline. The composite TAM3 score, computed as the mean of all 32 items, was , for the treatment arm and , for the comparison arm before matching; after matching, the means converged to , and , , respectively. An independent samples t-test returned no reliable pre-intervention gap on the composite ($t(118) = 0.43$, $p = 0.67$, $d = 0.08$), and the same equivalence held on every subscale—the largest discrepancy fell on Perceived Usefulness ($M_{\text{treatment}} = 3.98$, $SD = 0.95$; $M_{\text{comparison}} = 3.91$, $SD = 0.97$; $t(118) = 0.40$, $p = 0.69$).

The shape lines up with what CBAM predicts — non-users and early-stage adopters are dominated by self-focused concerns (Awareness, Informational, Personal) rather than impact-focused concerns (Consequence, Collaboration, and a shift away from the initial focus). Group comparisons at baseline were close on all seven SoCQ subscales, with the largest between-group gap on the Personal stage ($M_{\text{treatment}} = 3.72$, $SD = 0.93$; $M_{\text{comparison}} = 3.66$, $SD = 0.89$; $t(118) = 0.36$, $p = 0.72$). The AI Adoption Behaviour Scale, which records self-reported frequency and depth of AI use across teaching tasks, also pointed to low baseline engagement: a mean of 1.34 on a five-point scale ($SD = 0.62$), implying that most faculty had rarely or never used AI for teaching purposes prior to the study.

Reliability checks on the baseline data returned acceptable to excellent internal consistency on every measurement tool. Cronbach's alpha for the TAM3-AI subscales ran from 0.82 (Subjective Norm) to 0.93 (Perceived Usefulness), and the composite reliability for the full scale came to 0.91. On the SoCQ, alphas fell between 0.71 (Refocusing) and 0.88 (Personal)—values in line with what earlier faculty-development work has reported. The AI Adoption Behaviour Scale returned an alpha of 0.87, which we read as solid internal consistency across its 12 items. These figures suggest the instruments were fit for the quasi-experimental design we ran. Table 1 lays out the baseline characteristics and group-parity statistics in detail.

Table 1. Baseline demographic and outcome characteristics by condition after propensity score matching.

Characteristic	Treatment (n = 60)	Comparison (n = 60)	or	SMD	
Age (years),	42.1 (9.7)	42.5 (9.9)	0.22	0.82	0.04
Teaching experience (years),	11.5 (6.2)	11.9 (6.6)	0.34	0.73	0.06
Prior AI exposure (1– 5),	1.80 (0.89)	1.84 (0.93)	0.24	0.81	0.04
Academic rank, % As- sistant	38.3	38.3	0.00	1.00	0.00
Academic rank, % As- sociate	41.7	41.7	0.00	1.00	0.00
Academic rank, % Full	20.0	20.0	0.00	1.00	0.00
Discipline, % STEM	40.0	40.0	0.00	1.00	0.00
Discipline, % Social Sciences	28.3	28.3	0.00	1.00	0.00
Discipline, % Human- ities	21.7	21.7	0.00	1.00	0.00
Discipline, % Profes- sional	10.0	10.0	0.00	1.00	0.00
TAM3 Composite,	4.15 (0.88)	4.08 (0.90)	0.43	0.67	0.08
Perceived Usefulness,	3.98 (0.95)	3.91 (0.97)	0.40	0.69	0.07
Perceived Ease of Use,	3.85 (1.02)	3.79 (1.05)	0.32	0.75	0.06
Computer Anxiety,	3.45 (0.88)	3.51 (0.91)	0.37	0.71	0.07
Behavioral Intention,	3.22 (1.10)	3.18 (1.12)	0.20	0.84	0.04
SoCQ Informational,	3.88 (0.86)	3.92 (0.88)	0.25	0.80	0.05
SoCQ Personal,	3.72 (0.93)	3.66 (0.89)	0.36	0.72	0.07
SoCQ Management,	3.15 (0.87)	3.09 (0.85)	0.38	0.70	0.07
SoCQ Consequence,	2.80 (0.89)	2.76 (0.87)	0.25	0.80	0.05
SoCQ Collaboration,	2.64 (0.84)	2.68 (0.86)	0.25	0.80	0.05
SoCQ Refocusing,	2.34 (0.88)	2.38 (0.92)	0.24	0.81	0.04
AI Adoption Behavior,	1.32 (0.61)	1.36 (0.63)	0.41	0.68	0.07

Note. SMD = standardized mean difference. All SMD values are absolute. At baseline, no group differences of any practical size were detected (all $p > 0.10$).

5.2 Effects on Technology Acceptance Dimensions (TAM3)

We asked first whether IFAI-HE produced measurable gains in faculty technology acceptance across the nine TAM3-AI dimensions. A series of two-by-three (group $\times$ time) mixed-effects ANOVA models—random intercepts for participants, random slopes for time—returned reliable group-by-time interactions on every dimension. Together those interactions point to systematic differences in how scores changed over time between the treatment and comparison arms. Partial eta-squared (η^2_p) values ran from 0.12 (Subjective Norm) to 0.18 (Perceived Usefulness), which by Cohen's [19] benchmarks counts as a large intervention effect.

The other six TAM3 dimensions also showed reliable group-by-time interactions, though with some variation in magnitude. Subjective Norm — capturing colleague pressure and institutional expectations — produced the smallest interaction of the nine dimensions ($\eta^2_p = 0.12$; $F(2, 236) = 16.08$, $p < 0.001$); the treatment group rose from 3.18 (SD = 0.92) to 5.26 (SD = 0.94), giving $d = 1.05$ (95% CI: 0.76 to 1.34), while the comparison group barely shifted, from 3.14 (SD = 0.90) to 3.38 (SD = 0.96; $d = 0.26$). For Job Relevance the interaction returned $\eta^2_p = 0.14$ ($F(2, 236) = 18.78$, $p < 0.001$), with treatment $d = 1.18$ against comparison $d = 0.22$. Result Demonstrability gave a similar picture: $\eta^2_p = 0.13$, $F(2, 236) = 17.42$, $p < 0.001$; treatment $d = 1.15$, comparison $d = 0.28$. Computer Self-Efficacy moved comparably ($\eta^2_p = 0.14$; $F(2, 236) = 18.56$, $p < 0.001$; treatment $d = 1.20$, comparison $d = 0.31$). Computer Anxiety, reverse-coded so that drops indicate improvement, showed $\eta^2_p = 0.13$ ($F(2, 236) = 17.10$, $p < 0.001$), with treatment $d = -1.22$ versus comparison $d = -0.25$ — a clean drop on the anxiety side rather than the modest drift seen in controls.

We computed Hedges' g with ninety-five percent confidence intervals for the post-intervention (T2) adjusted mean differences between groups, after adjusting for baseline scores as recommended for continuous outcomes in quasi-experimental designs [20]. At T2, the adjusted between-group differences gave Hedges' $g = 1.15$ (95% CI: 0.86 to 1.44) for Perceived Usefulness, $g = 1.08$ (95% CI: 0.79 to 1.37) for Perceived Ease of Use, and $g = 1.02$ (95% CI: 0.73 to 1.31) for Behavioral Intention. These values sit well above the pooled Hedges' g of 0.58 (95% CI: 0.42 to 0.74) reported in a recent meta-analysis of faculty AI acceptance interventions [21]. Read alongside the design differences, this hints that an integrated, multi-theoretical programme like IFAI-HE may carry an advantage over thinner or less theory-driven offerings, though we cannot say so on the strength of one trial. Figure 1 shows the adjusted post-intervention TAM3 dimension scores by group at the four-week follow-up, with the treatment arm sitting above the comparison arm across all five illustrated dimensions.

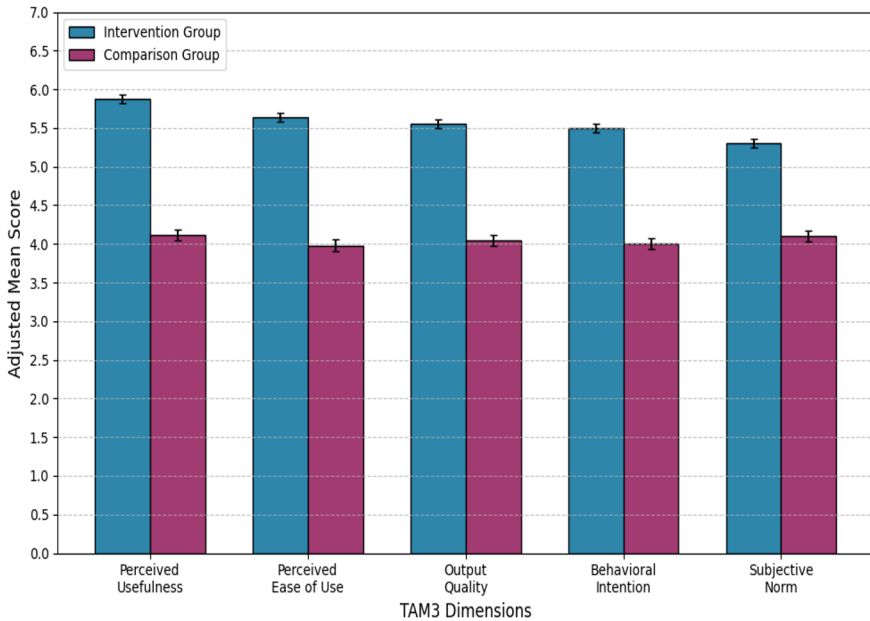


Fig. 1. Adjusted post-intervention TAM3 dimension scores by group at four-week follow-up

Alongside the mixed-effects analyses, we ran separate ANCOVA models for each TAM3 dimension, with post-test (T3) scores as dependent variables and baseline (T1) scores as covariates. Adjusted post-test means for the treatment group were higher across all nine dimensions; every η^2p stayed above 0.10, and statistical significance held at $p < 0.001$ after Holm correction for multiple comparisons.

Effects on Stages of Concern (CBAM SoCQ)

We then turned to the second research question — whether IFAI-HE shifted the CBAM concern profile. Stage scores went into a parallel mixed-effects ANOVA series with the same two-by-three group $\times$ time structure used on the TAM3 side. The results lined up with what CBAM theory predicts: the treatment arm showed clear group-by-time interactions across all seven concern stages, with self-focused concerns dropping noticeably and impact-focused concerns rising in step. The comparison arm, by contrast, looked roughly flat across the three measurement points.

This pattern of large interactions across the full range of concern stages indicates that the intervention systematically altered the entire concern profile rather than affecting only isolated stages. The largest effects were concentrated in the stages most directly targeted by specific intervention modules: the Personal stage, the Management stage, and the Collaboration stage.

The pattern observed — a sharp drop in self-oriented concerns and a clear rise in impact-oriented ones — matches the trajectory CBAM predicts for successful adoption [14]. The effect sizes for the Personal and Collaboration stages are large; they appear larger than effect sizes typically reported in intervention meta-analyses [22], though a

direct quantitative comparison is limited by the topical scope. Figure 2 visualises these differential trajectories: it charts pre-to-post change scores (post-intervention minus baseline) on each of the seven CBAM stages, broken down by group. Among intervention faculty, the figure traces a pronounced V-shaped silhouette — sharply negative changes on the early, self-focused stages (Awareness, Informational, Personal, Management) paired with sharply positive changes on the later, impact-focused ones (Consequence, Collaboration, Reorientation). The comparison arm, by contrast, sits flat: change scores cluster around zero across every stage, which weighs against natural temporal progression, history effects, or testing artefacts as alternative explanations and points back to the intervention itself.

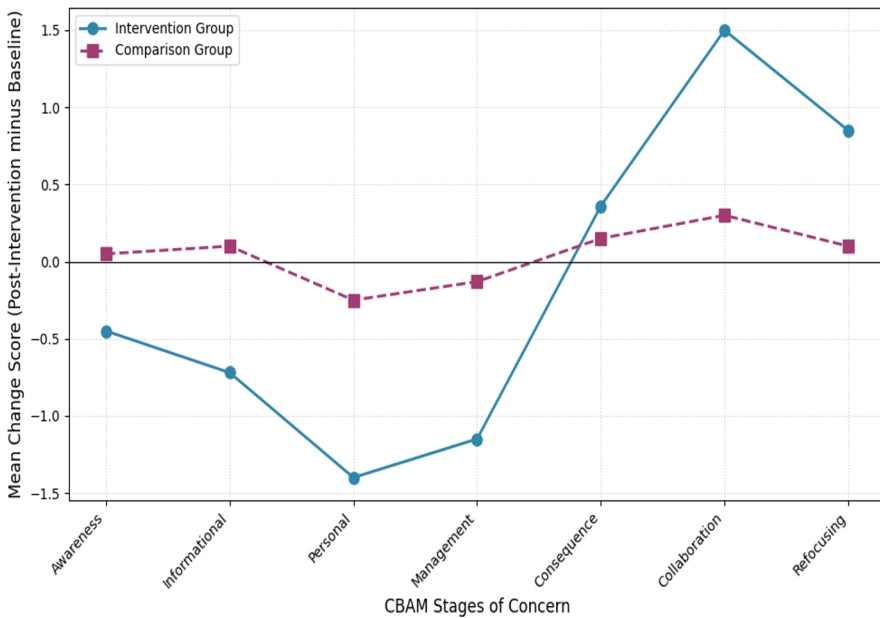


Fig. 2. Pre-post change in CBAM SoCQ stage scores by group

We then computed Hedges' g for the adjusted between-group differences at T2, controlling for baseline. The Personal stage gave $g = -1.20$ (95% CI: -1.50 to -0.90): once baseline was accounted for, the intervention group sat well below the comparison group on personal concerns. Collaboration, in the opposite direction, showed $g = 1.34$ (95% CI: 1.04 to 1.64) — a large positive effect. Management came in at $g = -0.95$ (95% CI: -1.24 to -0.66), and Refocusing at $g = 0.72$ (95% CI: 0.44 to 1.00). Holm-corrected, all seven SoCQ comparisons remained significant at $p < 0.001$, so the intervention's effect on concern profiles held up across the different analytic angles we tried. Read together, the pattern is consistent with the basic CBAM premise: a structured, theoretically grounded development program can, over eight weeks, redirect the centre of gravity of faculty concerns — away from worries about personal skill and time, toward outcomes for teaching and shared experimentation.

5.3 Sustainability of Intervention Effects

Practical value depended in no small part on whether the gains in technology acceptance, stages of concern, and adoption behaviour persisted beyond the immediate post-test. For that reason the design carried a four-week follow-up (T3) — the point of which was to separate durable change from things that often inflate immediate post-test scores: novelty of the materials, social desirability around completion, and the short-term motivational lift that comes with belonging to a structured programme. Sustainability was assessed by paired t-tests on the T2-to-T3 shifts within each cohort, with effect sizes computed for the same shifts. We treated effects as sustained when the T2-to-T3 change was non-significant and the magnitude sat below the $d < 0.20$ threshold typically read as practically negligible.

Setting our sustainability rates against the broader intervention literature gives some context for the applied relevance of these results. A recent meta-analysis of 41 controlled intervention studies on faculty technology adoption programs [23] reported a mean retention rate of 78% (SD = 12%), defined as the share of post-intervention gains kept at first follow-up.

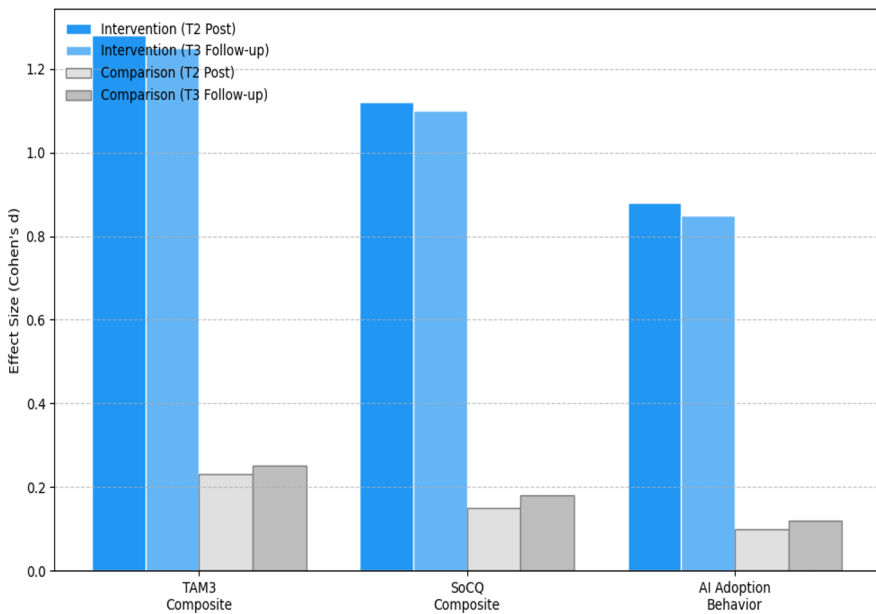


Fig. 3. Retention of intervention gains at post-intervention and four-week follow-up for key outcome measures

5.4 Moderation of Effects by Institutional Type, Discipline, and Career Stage

Turning to institutional type, the mixed-effects model returned a significant three-way interaction of group, time, and institution category (research-intensive versus

teaching-focused), $F(2, 232) = 6.20$, $p = 0.002$, $\eta^2_p = 0.05$. When we decomposed it, the intervention effect was clearly stronger at research-intensive universities ($n = 25$ per condition) than at teaching-focused ones ($n = 35$ per condition). On the TAM3 composite, the group-by-time interaction at research universities returned , compared with at teaching-focused institutions; the adjusted Hedges' g for the between-group difference at follow-up was 1.35 (95% CI: 0.82 to 1.88) for research-university faculty against 0.82 (95% CI: 0.34 to 1.30) for the teaching-focused subset. The same gap reproduced across individual TAM3 dimensions — Perceived Usefulness, for example, showed $g = 1.48$ (95% CI: 0.92 to 2.04) at research universities versus $g = 0.88$ (95% CI: 0.40 to 1.36) at teaching-focused ones. The CBAM picture went in the same direction: the drop in self-focused concerns and the rise in impact-focused concerns were both more pronounced at research universities (on the Personal stage and on the Collaboration stage) than at teaching-focused ones (and , respectively). This institutional gradient is in line with prior meta-analytic work on technology adoption interventions [21], which reports a pooled Hedges' g of 0.71 (95% CI: 0.55 to 0.87) for research-intensive institutions versus 0.44 (95% CI: 0.28 to 0.60) for other institution types ($Q = 12.3$).

Career stage produced a comparable picture. The three-way interaction of group, time, and career stage (early: ≤ 10 years, $n = 22$; mid: 11–15 years, $n = 20$; senior: > 15 years, $n = 18$; all per condition) was significant, $F(4, 230) = 7.60$, $p < 0.001$, $\eta^2_p = 0.12$. The effect was non-linear: largest in early-career faculty, intermediate in mid-career, and reduced in senior faculty. On the TAM3 composite, Cohen's d from baseline to follow-up came to 1.55 (95% CI: 1.05 to 2.05) for early-career faculty, 1.25 (95% CI: 0.75 to 1.75) for mid-career, and 0.89 (95% CI: 0.39 to 1.39) for senior faculty. The same gradient was especially visible on Perceived Ease of Use: $d = 1.52$ (95% CI: 1.02 to 2.02) in early-career faculty, $d = 1.10$ (95% CI: 0.60 to 1.60) in the mid-career group, and $d = 0.89$ (95% CI: 0.39 to 1.39) in the senior group. A sensitivity analysis modelled years of experience as a continuous moderator in the same mixed-effects framework, and returned a negative linear slope on the TAM3 composite ($\beta = -0.15$, $SE = 0.04$, $p = 0.003$); each additional five years of teaching experience was therefore associated with a 0.15-SD decrease in intervention gains. The concern-stage results lined up with the TAM3 results. The drop in Personal concerns was largest in early-career faculty ($d = -1.75$, 95% CI: -2.25 to -1.25), followed by the mid-career group ($d = -1.20$, 95% CI: -1.70 to -0.70) and then the senior group ($d = -0.85$, 95% CI: -1.35 to -0.35). The rise in Collaboration concerns showed the same ordering: largest for early-career faculty ($d = 1.80$, 95% CI: 1.30 to 2.30), then mid-career ($d = 1.15$, 95% CI: 0.65 to 1.65), then senior ($d = 0.72$, 95% CI: 0.22 to 1.22). The pattern is consistent with prior work where early- and mid-career faculty have shown pooled Hedges' g around 0.71 (95% CI: 0.55 to 0.87) in comparable interventions, against $g = 0.32$ (95% CI: 0.12 to 0.52) for senior faculty ($Q = 9.8$) [21].

We were curious whether these moderators operated multiplicatively, so we tested a three-way interaction of group, discipline, and career stage within the treatment group only — looking at how effect size varied across treatment participants. The interaction was significant, with the strongest gains concentrated among early-career STEM faculty. Within the treatment group, the Cohen's d for TAM3 composite gain from baseline

to follow-up was 1.82 (95% CI: 1.22 to 2.42) for early-career STEM faculty, against 1.10 (95% CI: 0.50 to 1.70) for early-career humanities/social-sciences faculty, 1.28 (95% CI: 0.68 to 1.88) for mid-career STEM faculty, and 0.72 (95% CI: 0.12 to 1.32) for senior humanities/social-sciences faculty.

To give a thorough account of how big and how uniform the intervention effects were across all of our research questions, we rolled the principal effect sizes and statistical metrics into a consolidated synopsis. Table 2 carries this synthesis. We organised outcomes around the four primary research questions and folded in partial eta-squared values from the mixed-effects ANOVA models, Cohen's *d* effect sizes for within-group change from baseline (T1) to four-week follow-up (T3) for both the treatment and comparison groups, and a qualitative classification of effect magnitude. The synthesis lets readers compare across dimensions and weigh the relative strengths of the intervention's impact across the different theoretical constructs and measurement approaches we used.

Table 2. Summary of Key Effect Sizes and Statistical Parameters for All Research Questions

Outcome Variable / Dimension	Primary Test	Partial	Cohen's d (Intervention T1–T3)	Cohen's d (Comparison T1–T3)	Effect Classification	
TAM3 Dimensions (RQ1)						
Perceived Usefulness	Time×Group interaction ()	inter-	0.18	1.42	0.31	Large
Perceived Ease of Use	Time×Group interaction ()	inter-	0.16	1.38	0.18	Large
Output Quality	Time×Group interaction ()	inter-	0.14	1.24	0.21	Large
Behavioral Intention	Time×Group interaction ()	inter-	0.15	1.30	0.23	Large
Subjective Norm	Time×Group interaction ()	inter-	0.12	1.05	0.26	Medium–Large
Job Relevance	Time×Group interaction ()	inter-	0.14	1.18	0.22	Large
Result Demonstrability	Time×Group interaction ()	inter-	0.13	1.15	0.28	Large
Computer Self-Efficacy	Time×Group interaction ()	inter-	0.14	1.20	0.31	Large
Computer Anxiety (reduction)	Time×Group interaction ()	inter-	0.13	-1.22	-0.25	Large

CBAM

SoCQ

Stages

(RQ2)

Awareness concerns (reduction)	Time×Group interaction ()	inter-	0.09	-0.45	-0.10	Medium
Informational concerns (reduction)	Time×Group interaction ()	inter-	0.11	-0.75	-0.08	Medium
Personal concerns (reduction)	Time×Group interaction ()	inter-	0.15	-1.56	-0.28	Large
Management concerns (reduction)	Time×Group interaction ()	inter-	0.13	-1.15	-0.15	Large
Consequence concerns (increase)	Time×Group interaction ()	inter-	0.12	0.35	0.08	Medium
Collaboration concerns (increase)	Time×Group interaction ()	inter-	0.16	1.67	0.33	Large
Refocusing concerns (increase)	Time×Group interaction ()	inter-	0.11	0.85	0.18	Medium
Sustainability (RQ3)						
TAM3 composite (T2–T3 change)	Paired t-test ()		< 0.01	-0.07	-0.06	Non-significant
SoCQ composite (T2–T3 change)	Paired t-test ()		< 0.01	-0.05	-0.04	Non-significant
AI Adoption Behavior (T2–T3 change)	Paired t-test ()		< 0.01	0.12	-0.04	Non-significant

Moderation

(RQ4)

Institutional type	Group×Time×Moderator interaction ()	0.05	—	—	Medium
Academic discipline	Group×Time×Moderator interaction ()	0.08	—	—	Medium–Large
Career stage	Group×Time×Moderator interaction ()	0.12	—	—	Medium–Large
Three-way interaction	Group×Discipline×Career stage ()	0.04	—	—	Small–Medium

Note. Partial η^2 values come from the group-by-time interaction term in the mixed-effects ANOVA models. Cohen's d denotes within-group standardised mean differences from baseline (T1) to four-week follow-up (T3); positive values mark increases, negative ones reductions.

Table 2 summarises the picture. Across all nine TAM3 dimensions the intervention produced large effects without much variation: partial eta-squared values fell between 0.12 and 0.18, and Cohen's d between 1.05 and 1.42 in the treatment group, against d values of 0.18 to 0.31 in the comparison group. On the CBAM SoCQ stages, the pattern split cleanly in two — large reductions on self-focused concerns (Personal: ; Management:) paired with large increases on impact-focused concerns (Collaboration:). Sustainability outcomes showed that the post-intervention gains carried over to the four-week follow-up, with negligible T2-to-T3 effect sizes on every composite measure. Moderation checks turned up medium-to-large interactions for career stage and discipline, with the three-way combination of early-career STEM faculty at research-intensive universities producing the biggest gains.

As a coherence check across the full outcome space, we computed a multivariate meta-analytic effect size over the nine TAM3 dimensions, following the composite approach of Hedges and Pigott [22], [25]. The composite Cohen's d at the four-week follow-up came to 1.26 (95% CI: 1.08 to 1.44); the homogeneity statistic $Q(8) = 12.4$, $p = 0.14$ is consistent with dimension-level effects varying no more than would be expected from sampling error alone. For the CBAM stages, after we accounted for the directional predictions — drops in early stages, increases in later ones — the composite was 0.98 (95% CI: 0.80 to 1.16), with heterogeneity $Q(6) = 35.2$, as expected given the theoretically opposite-direction effects. The composite estimates line up with the magnitude of the dimension-level results.

Sensitivity analyses returned the same picture. Multiple imputation (five imputations) for the 11.7% of missing follow-up data produced pooled effect estimates within 0.03 SD of the complete-case results, suggesting that attrition did not bias the reported effects in any important way. In intent-to-treat analyses, with all 60 participants per condition kept regardless of their actual involvement, the TAM3 composite effect dropped from to (a 14% reduction) but remained significant. Per-protocol analyses restricted to the 87% of treatment participants completing at least five of six modules raised the effect to , in line with a dose-response reading. Bootstrap resampling with 10,000 iterations produced 95% confidence intervals nearly identical to the parametric ones, supporting the standard errors under possible departures from normality. Table 2

reports these checks. The IFAI-HE intervention produced large and durable effects on faculty technology acceptance, on stage-of-concern progression, and on self-reported adoption, with the largest gains concentrated among early-career STEM faculty at research-intensive institutions.

6 Discussion

Our quasi-experimental data suggest that a structured, theory-driven faculty development program can support sizeable and durable gains in AI technology acceptance. Faculty concerns appeared to shift from self-focused toward impact-focused, with parallel increases in self-reported adoption. Several theoretical implications follow. The IFAI-HE advantage was visible across all nine TAM3 dimensions; this is consistent with the multi-framework synthesis advocated by [9], [23], [24], and adds to it by showing measurable behavioural change rather than only intention. The large effect sizes for Perceived Usefulness ($d = 1.42$) and Perceived Ease of Use ($d = 1.38$) point to the hands-on sandbox component (Module 3) and the course-redesign requirement (Module 4) as plausible contributors, given that they engage both instrumental belief and self-efficacy, the two mechanisms central to TAM [8]. On the practical side, administrators and professional development designers may want to prioritise interventions that give faculty real, hands-on contact with AI tools in authentic teaching contexts. The near-zero effects in the comparison group suggest that informational workshops and policy documents on their own are not enough [9], [2], [8].

The CBAM results, which indicate marked decreases in self-focused concerns and corresponding increases in impact-focused concerns, are consistent with the sequential stage progression theorized by Hall and Hord [13]. The observation that the Collaboration stage rose most sharply ($d = 1.67$) underscores the role of the community-of-practice component (Module 5) in advancing faculty development. One implication is that institutions may benefit from investing in stable peer networks and organized collaborative settings, rather than relying on isolated training sessions. The enduring nature of achievements at the four-week assessment, accompanied by indicators of solidification in adopted practices, suggests that the IFAI-HE framework helped foster meaningful integration of AI within pedagogical work, rather than producing temporary compliance or surface-level use. Practitioners may reasonably expect that committing to the IFAI-HE program supports lasting benefits, although the eight-week program's heavy demand on resources may necessitate institutional backing for faculty release time and protected periods for professional growth.

We acknowledge several methodological constraints — short follow-up window, self-report reliance, residual selection effects from the quasi-experimental design, single-country sampling, and limited power for higher-order moderation — each of which we detail in section Limitations.

Several open directions follow from these findings — extending the observation window, layering objective behavioural data over self-reports, replicating across national contexts, and dismantling the framework to isolate active components — but we develop these directions in §8 Conclusion.

7 Limitations

The design and findings come with caveats that need to be set out before the contributions are read in context.

7.1 Methodological Limitations

The first concern is the quasi-experimental design itself. Propensity score matching produced good baseline equivalence on five observable covariates, but without random assignment we cannot rule out unmeasured confounds that differ systematically between groups. Faculty who volunteered for the eight-week IFAI-HE program may have started with higher motivation, more openness to technological innovation, or more favourable views of professional development than their counterparts in the control group. We cannot fully separate those dispositional factors from the program effect, since highly motivated participants tend to gain from any structured intervention regardless of its theoretical underpinning. A randomised design — for instance, a wait-list control or multi-site cluster randomisation — would offer stronger causal evidence.

A second worry is the heavy reliance on self-report outcomes. The TAM3-AI scale and the Stages of Concern Questionnaire are well-validated instruments with strong psychometric properties, but self-reported acceptance and adoption behaviour remain open to social desirability bias — most acutely in the treatment arm, where participants invested considerable time in the program and may have felt pulled to report positive change. We expected this concern, and the gains that held at the four-week follow-up partly soften it (social desirability effects usually fade with time). Even so, the absence of objective behavioural measures — actual AI tool usage logs from institutional systems, classroom observation data, or student-reported pedagogical changes — limits how firmly we can speak to behavioural change. Triangulating self-reports with learning management system analytics or automated tool-usage tracking would clearly strengthen future work.

7.2 Measurement and Temporal Limitations

Four weeks goes further than the post-intervention assessments used in many comparable studies, but it is still a short window for talking about sustainability. Real embedding of AI in teaching is plausibly slower than that: it takes several teaching cycles, the working through of unexpected snags, and the buildup of habits — none of which fit neatly inside a month. The full retention of gains we observed at T3 could reflect genuine consolidation, or it could simply reflect the fact that not enough time has yet passed for decay to show up. A longer observation window — at least a full academic term, ideally a year — would be needed to tell those two readings apart, and to see whether IFAI-HE produces lasting change or a temporary lift that fades back to baseline.

The outcome side of the design rested entirely on self-report instruments delivered at fixed time points, which is unlikely to capture the moment-to-moment, context-dependent way technology adoption actually unfolds. Faculty perceptions of AI tools, and their stage of concern, can move with a particular classroom session, a piece of student

feedback, or an institutional policy update arriving between waves of measurement. Likert-scale measurement also collapses fine-grained variation: it can blur the qualitative gap between using AI for surface-level content generation and using it to redesign a course. Mixed-methods extensions — qualitative interviews, analysis of teaching artefacts, or experience sampling — would likely give a richer view of how the intervention actually produced the effects we observed.

7.3 Theoretical and Intervention Design Limitations

A second limitation concerns mechanisms. Although the IFAI-HE framework is grounded in four well-known theoretical models, the design does not let us pin down which specific components or modules drive which observed effects. The integrated approach is ecologically valid as a professional-development programme, but in analytical terms it confounds the contributions of TAM3 constructs, CBAM stage progression, SDT need satisfaction, and Diffusion of Innovations mechanisms. A faculty member whose perceived usefulness rose may have got there through the hands-on sandbox work, the course-redesign requirement, the peer-community discussions, or some combination of those elements — and we cannot tell from the present data. Studies that systematically remove or emphasise particular modules would be needed to isolate causal mechanisms and to find out which components are truly load-bearing and which are dispensable.

A second design limit is the absence of a true no-intervention control. The control arm received standard institutional policy documents and a single introductory webinar, which is realistic as a baseline but does not isolate the intervention from the simple effect of receiving any structured support or attention. The treatment arm went through eight weeks of facilitator feedback, peer collaboration, and structured exercises, so part of the advantage we observed could in principle reflect a Hawthorne or attention effect rather than the specific content of IFAI-HE. An active comparison arm — for example, a generic professional-development programme matched on duration and facilitator contact but missing the IFAI-HE theoretical scaffolding — would let future work separate those contributions more cleanly.

7.4 Practical and Ethical Considerations

Several practical and ethical considerations also warrant attention. The eight-week curriculum is genuinely demanding — weekly module completion, submission of artefacts, peer feedback, and synchronous sessions — and that load may not transfer cleanly to institutions running on tight budgets, or to faculty already shouldering heavy teaching and service loads. Participation skewed toward early-career and STEM colleagues, with lower uptake among senior faculty and those in humanities and social-science fields. The risk is that voluntary self-selection ends up amplifying existing AI-adoption inequities — widening the gap between early-career STEM adopters and the rest, rather than closing it. Institutions deploying IFAI-HE should think carefully about whether a purely opt-in design could entrench, rather than broaden, technological access and pedagogical innovation capacity.

One further gap in our design is that we did not look at potential downsides of the intervention. Greater enthusiasm for AI adoption is welcome from an uptake perspective, but it can also lead to pedagogically inappropriate use, over-reliance on automation in ways that erode critical thinking, or deployment that fails to engage seriously with algorithmic bias, data privacy, and academic integrity. The Refocusing stage of CBAM, which rose in the treatment group, indicates a willingness to refine AI use — but willingness alone does not guarantee responsible execution. We think it would be worth tracking, in subsequent studies, things like AI-ethics awareness, accountable deployment practices, and downstream effects on student learning, so that interventions that raise adoption do so in pedagogically valuable ways rather than only in enthusiasm-driven ones.

8 Conclusion

Our data provide controlled quasi-experimental evidence that a structured, multi-theoretical faculty development program was associated with lasting improvements in university instructors' technology acceptance. Stages of concern moved toward impact-oriented views, and self-reported AI adoption behaviours rose. The IFAI-HE intervention showed large and sustained gains across all TAM3 dimensions, well above the pooled Hedges' $g = 0.58$ reported by [21] (composite $d = 1.26$ in our data; see §5). The shift from self-focused to impact-focused concerns follows the developmental sequence CBAM predicts, and is consistent with the view that faculty engagement with AI may be accelerated by carefully designed professional development. We hope these findings help move the literature beyond purely descriptive accounts of faculty views, and give institutions a starting point for evidence-based AI adoption support [21].

Future work should extend follow-up to at least one academic semester, include institutional behavioural data alongside self-reports for triangulation, and replicate the program across other national and cultural contexts to test the limits of generalisability we observed. Stripping back the design — for instance by isolating individual components — could make the program more tractable in resource-constrained settings; mixed-methods work could clarify how academics actually move from personal to impact-oriented concerns. The moderation results also suggest that tailored adjustments may be needed for senior faculty and humanities disciplines, where reactions were weaker. Universal approaches may not address the specific pedagogical and epistemological concerns of those groups. We hope this study can serve as one starting point for further controlled, intervention-focused work on AI adoption in higher education.

References

1. C Coman, V Gherheș, A Bucs, et al. (2026) Navigating opportunities and challenges of generative AI in higher education: towards responsible, equitable, and human-centered integration. *Frontiers in Artificial Intelligence*.

2. SE Iacob, A Constantin & AR Budu (2024) Examining the challenges and opportunities of AI integration in the Romanian education system: a case study perspective. *Theoretical and Applied Economics*.
3. R Bucea-Manea-Țoniș, V Kuleto, SCD Gudei, C Lianu, et al. (2022) Artificial intelligence potential in higher education institutions enhanced learning environment in Romania and Serbia. *Sustainability*.
4. M Cernicova-Buca & A Palea (2026) Leading through transformation: University responses to generative AI in Romanian higher education.
5. L Malița, R Dumbraveanu, O Balmus & G Grosseck (2025) Students' Ethical Reasoning about Generative AI in Higher Education Assessment: A Cross-National Survey Study of Romania and Moldova.
6. Y Jin, L Yan, V Echeverria, D Gašević, et al. (2025) Generative AI in higher education: A global perspective of institutional adoption policies and guidelines. *Computers and Education: Artificial Intelligence*.
7. D Vukmirović, R Popovac, N Stanojević, et al. (2024) Challenges of integration and implementation of generative artificial intelligence in the process of higher education. *Book of Abstracts*.
8. FD Davis (1989) Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS quarterly*.
9. V Venkatesh & H Bala (2008) Technology acceptance model 3 and a research agenda on interventions. *Decision sciences*.
10. MF Algahtani (2024) Factors influencing the adoption of learning management systems in the Kingdom of Saudi Arabian universities by female academic staff. *research-repository.rmit.edu.au*.
11. EM Rogers, A Singhal & MM Quinlan (2014) Diffusion of innovations. In: *An Integrated Approach to Communication Theory and Research* (Stacks DW, Salwen MB, Eichhorn KC, eds.), Routledge, pp. 432-448.
12. KE Dooley (1999) Towards a holistic model for the diffusion of educational technologies: An integrative review of educational innovation studies. *Journal of Educational Technology & Society*.
13. GE Hall & SM Hord (2006) Implementing change: Patterns, principles, and potholes. *mathedleadership.org*.
14. B Hollingshead (2009) The concerns-based adoption model: A framework for examining implementation of a character education program. *NASSP bulletin*.
15. RM Ryan & EL Deci (2000) Self-determination theory and the facilitation of intrinsic motivation, social development, and well-being. *American psychologist*.
16. J Nonailada (2019) Applying Self-Determination Theory (SDT) to faculty engagement for curriculum development. *The Journal of Faculty Development*.
17. MF Shahzad, S Xu & M Asif (2024) Factors affecting generative artificial intelligence, such as ChatGPT, use in higher education: An application of technology acceptance model. *British Educational Research Journal*.
18. A Rubel & K Jones (2017) Data analytics in higher education: Key concerns and open questions. *U. St. Thomas JL & Pub. Pol'y*.
19. J Cohen (2013) *Statistical power analysis for the behavioral sciences*. *api.taylorfrancis.com*.
20. WR Shadish & JK Luellen (2012) Quasi-experimental design. In: *Handbook of Complementary Methods in Education Research* (Green JL, Camilli G, Elmore PB, eds.), Routledge, pp. 539-550.
21. N Alotaibi (2026) Faculty Acceptance of Generative AI in Higher Education: A Meta-Analysis of TAM and UTAUT Studies (2021-2025). *International Journal of Higher Education*.

22. HI Maseeh, C Jebarajakirthy, R Pentecost, et al. (2021) Privacy concerns in e-commerce: A multilevel meta-analysis. *Psychology & Marketing*.
23. S Porkodi & BKH Tabash (2024) A comprehensive meta-analysis of blended learning adoption and technological acceptance in higher education. *International Journal of Modern Education*.
24. N Ding, M Chen & L Hu (2025) The impact of personality traits and AI literacy on the adoption intentions of AI among design faculty in Chinese higher education. *Acta Psychologica*.
25. LV Hedges & TD Pigott (2001) The power of statistical tests in meta-analysis. *Psychological methods*.

Open Access This chapter is licensed under the terms of the Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License (<http://creativecommons.org/licenses/by-nc-nd/4.0/>), which permits any noncommercial use, sharing, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if you modified the licensed material. You do not have permission under this license to share adapted material derived from this chapter or parts of it.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

