

Does Digital Tool Usage Translate Digital Competence into Workflow Efficiency? Explanatory and Predictive Evidence from AI-Augmented Banking

Meng Chai¹, and Doris Wong Hooi Ten^{1*}

¹Faculty of Artificial Intelligence, Universiti Teknologi Malaysia, Kuala Lumpur, 81310, Malaysia.

*Corresponding Author: Doris Wong Hooi Ten. Email: doriswong1@sina.com

Received: March 27, 2026 Accepted: August 08, 2026

Abstract: Digital competence is frequently presented as a prerequisite for effective work with artificial-intelligence-enabled systems. Yet competence does not show that technology has been incorporated into task execution. This study examines whether task-embedded digital tool usage translates digital competence into perceived technology-enabled task efficiency. A cross-sectional questionnaire yielded 301 complete responses from a banking-sector sample. The proposed capability-enactment model was evaluated with reflective partial least squares structural equation modelling (PLS-SEM) and 5,000 bootstrap resamples. Its predictive implications were tested through repeated nested cross-validation. Ridge regression and random forests used competence-only, usage-only, and combined item-level feature sets. Digital competence was associated with digital tool usage ($\beta = 0.388$) and perceived task efficiency ($\beta = 0.261$). Tool usage retained an association with efficiency after competence was controlled ($\beta = 0.303$). The indirect association was 0.118 (95% bootstrap CI 0.074-0.171), supporting partial statistical mediation. The combined random forest produced the lowest held-out RMSE (0.835) and highest pooled Q²_{predict} (0.240). It also outperformed competence-only and usage-only forests in paired comparisons. Group permutation analysis assigned more predictive information to usage than to competence, although both were useful. The results identify digital tool usage as a proximal behavioural mechanism between capability and perceived workflow return. They also show why explanatory significance should be complemented by out-of-sample validation when designing or evaluating human-AI work systems.

Keywords: Digital competence; System Use; Digital Tool Usage; Perceived Task Efficiency; Human-AI Workflow; PLS-SEM; Predictive Modelling

1. Introduction

Artificial intelligence is entering knowledge work through document processing, information retrieval, decision support, customer interaction, risk assessment, and workflow automation. In most organizational settings, however, AI is not an autonomous production unit. It is one component of a human-AI system. Employees frame tasks, configure tools, interpret output, detect failure, and integrate results into work [1-4]. System effectiveness therefore depends on both computational capability and human capacity. Employees with the same interface can still obtain different workflow returns. They may differ in information evaluation, routine adaptation, troubleshooting, and decisions to accept, revise, or reject an automated recommendation.

Digital competence has consequently become a central construct in research on digital work. Workplace accounts define it more broadly than operational skill: it combines knowledge, evaluative judgement, adaptability, problem solving, and responsible technology use [5-7]. These attributes are especially salient in AI-augmented work because generated outputs can appear fluent while remaining incomplete, biased, or wrong. A competent user may be better equipped to formulate a query, assess source

quality, detect a mismatch between output and task requirements, and preserve professional accountability. It is therefore plausible that digital competence is related to individual performance and efficiency. Empirical work has reported positive associations between competence and job performance, including evidence that intermediate psychological or organizational mechanisms help explain this relationship [14].

Competence, nevertheless, is not equivalent to enactment. A user can possess the knowledge required to operate and evaluate a system without embedding it in daily work. Conversely, frequent use can result from mandatory procedures while remaining shallow, poorly matched to the task, or weakly informed by judgement. This distinction is obscured when studies move directly from competence to performance or operationalize technology adoption through intention, frequency, or access alone. Information-systems research has long argued that system use should be defined in relation to the user, system, task, and function represented by the measure [9,10]. The relevant behavioral question is not simply whether a tool was opened, but whether it became part of the sequence through which work is completed.

The present study addresses this problem by positioning digital tool usage (DTU) as a behavioral enactment layer between digital competence (DC) and perceived technology-enabled task efficiency (PTE). DTU refers to regular, active, multi-step integration of digital or AI-enabled tools into task execution and work-related judgement. PTE captures perceived speed, process smoothness, reduced rework, and reduced unnecessary effort attributable to those tools. The outcome is deliberately narrower than overall job performance. It does not represent objective productivity, work quality, contextual performance, or organizational value. This boundary prevents an improvement in self-reported workflow efficiency from being mistaken for a comprehensive performance gain.

The distinction between capability and enactment also creates a methodological issue. Structural models can establish whether proposed associations are coherent with the observed covariance structure. Statistical significance and in-sample explained variance do not establish predictive performance on new observations [18]. A theoretically sensible mediation model can therefore be useful as explanation while remaining weak as a prediction system. Conversely, an opaque prediction model can reduce held-out error without clarifying the mechanism connecting human capability with system behavior. Information-systems research therefore treats explanation and prediction as complementary rather than interchangeable objectives [18,19].

The study combines both objectives. PLS-SEM estimates the competence-usage-efficiency pathway, while repeated nested cross-validation compares competence-only, usage-only, and combined features with ridge regression and random forest. The design separates measurement, structural association, and held-out generalization.

Three contributions follow. First, the capability-enactment account distinguishes competence from its behavioral conversion into a workflow. Second, PTE is bounded to self-reported technology-enabled efficiency rather than extended to objective or overall performance. Third, the explanation-prediction design evaluates theoretical interpretation and incremental predictive value with the same item pool but different analytical criteria.

2. Theoretical background and research hypotheses

2.1. Digital competence in AI-augmented work

Digital literacy, digital skills, technology competence, and digital competence overlap but are not identical. Digital literacy originally emphasized access to and understanding of information, whereas skill-oriented approaches often concentrated on operational tasks. Workplace research has broadened the concept to include information management, problem solving, critical evaluation, adaptability, ethical responsibility, and technology selection [5-7]. The common element is not mastery of one application. It is the capacity to act effectively and responsibly when tasks, information, and digital systems interact.

This broader interpretation is required for AI-enabled work. An employee may operate a generative interface successfully but still lack the competence to judge whether an output is relevant, sufficiently evidenced, compliant with organizational requirements, or safe to use. Similarly, rapid adoption of a new tool does not guarantee that the user can recover from unexpected output or understand when professional judgement should dominate system advice. Recent evidence links digital literacy to differentiated patterns

of ChatGPT adoption and utilization rather than to a single binary decision to use the system [8]. Experimental work on intelligent decision support likewise shows that user proficiency and system capability jointly shape reliance [17]. Digital competence is therefore best treated as a portfolio of evaluative, adaptive, and problem-solving capabilities.

The study operationalizes DC through seven indicators. They cover digital information search, recommendation reliability, rapid learning, troubleshooting, work adaptation, independent judgement, and integration with system output. This representation is narrower than comprehensive competency frameworks, which may also include communication, content creation, cybersecurity, digital identity, and well-being [7]. The narrower scope is intentional because the model concerns work execution rather than a complete inventory of digital citizenship.

Human capital theory offers a limited but useful starting point. Knowledge and skills can increase productive capacity when they are applied to work [21]. DC can accordingly be viewed as task-relevant human capital for AI-augmented workflows. Yet human capital is a capability stock; returns depend on whether the capability is mobilized in operational behavior. The next sections therefore distinguish the expected direct efficiency return of DC from the behavioral pathway through which competence is enacted.

2.2. Digital competence and perceived technology-enabled task efficiency

Competent users should be better able to reduce avoidable effort when using digital tools. Effective information search can shorten the path to a usable input. Reliability assessment can prevent time being lost to defective recommendations. Troubleshooting can reduce interruptions, and adaptation can help a user reorganize a task around the strengths and limitations of a new system. Professional judgement can also reduce both blind reliance and indiscriminate rejection. These capabilities should be associated with smoother work even when the frequency of tool use is held constant.

Prior studies support a positive relationship between digital competence and work outcomes. Workplace reviews connect competence with efficient performance across digitally mediated tasks [5]. Empirical research also relates competence to job performance through intermediate mechanisms [14]. Enterprise social media research links information-technology competence and work cooperation with employee performance [15]. The present outcome is narrower than job performance, but the same logic applies to technology-supported tasks. Thus:

H1. Digital competence is positively associated with perceived technology-enabled task efficiency.

2.3. Digital competence and task-embedded digital tool usage

System use is not a unitary behavior. It may refer to duration, frequency, breadth of functions, degree of dependence, or the extent to which a system is incorporated into a task. Burton-Jones and Straub argue that a usage measure should represent the system, user, and task implied by the theoretical model [9]. Doll and Torkzadeh similarly describe organizational use through task-related functions rather than through access alone [10]. For AI-enabled work, this distinction is decisive: opening an application is different from using it across information gathering, drafting, verification, decision support, and revision.

DC should be associated with deeper DTU because competent users have a larger repertoire for selecting functions, resolving friction, and adapting the technology to task requirements. They are also better positioned to calibrate reliance. A user who can identify errors and limitations may integrate a tool more confidently because use no longer requires uncritical trust. A less competent user may either avoid the tool or use it in a superficial manner that does not survive task complexity. Digital-transformation research accordingly treats employee competence as part of the process through which technical infrastructure becomes operational change [16]. Therefore:

H2. Digital competence is positively associated with task-embedded digital tool usage.

2.4. Digital tool usage and perceived technology-enabled task efficiency

The information-systems success model distinguishes system use from downstream benefits [11]. Use is therefore not inherently valuable. High frequency can reflect mandatory compliance, poor usability, or repeated correction. Task-technology fit theory further predicts that performance consequences depend on the alignment between task requirements and technological functionality [12]. The present construct does not measure raw frequency alone. Its indicators capture regular and active use, integration across multiple

work steps, importance to task completion, support for work-related judgement, and continued integration as tasks become more demanding.

When use has these task-embedded properties, it should be positively associated with PTE. Repeated incorporation of a tool can reduce switching costs, standardize routine operations, and accelerate information exchange. Multi-step integration can also reduce duplication when outputs move directly into later stages of a workflow. Active use for judgement can shorten search and comparison processes, provided that the user retains responsibility for verification. Measures of information-technology impact similarly distinguish task productivity from other consequences such as innovation and customer satisfaction [13]. The current claim is limited to the perceived efficiency dimension:

H3. Task-embedded digital tool usage is positively associated with perceived technology-enabled task efficiency after accounting for digital competence.

2.5. The mediating role of digital tool usage

H1-H3 imply an indirect pathway. DC provides the capacity to select, evaluate, adapt, and troubleshoot digital resources. DTU indicates whether that capacity becomes part of work execution. PTE represents the immediate perceived return of the enacted workflow. This sequence is described as a **capability-enactment mechanism**. The term names the theoretical process; it is not an additional latent variable.

The argument resembles mediation research in which digital competence is linked to job performance through a psychological or organizational state [14], but the mechanism is different. Here the intervening layer is observable self-reported behavior: the extent to which tools are incorporated into tasks. The model does not require full mediation. Competence may still be related to efficiency through better information judgement, faster problem diagnosis, and more effective coordination that are not completely represented by DTU. A positive indirect association together with a remaining direct association would therefore be consistent with partial statistical mediation.

H4. Digital tool usage carries a positive indirect association between digital competence and perceived technology-enabled task efficiency.

2.6. From structural explanation to out-of-sample prediction

The structural hypotheses concern associations among latent constructs. They do not determine whether the item responses predict PTE for cases not used to estimate the model. Predictive evaluation requires explicit separation of training and test observations and must include model selection inside the training data [18]. Without this separation, performance estimates are optimistic because the same observations influence both parameter selection and evaluation.

Three feature representations follow from the capability-enactment argument. A competence-only representation captures whether a respondent possesses the human capability needed to use digital tools effectively. A usage-only representation captures the proximal behavioral layer. A combined representation retains both the capability and its enactment. If DTU only repeats information already contained in DC, the combined representation should not improve held-out error. If the layers are complementary, the combined features should perform best. Because the relationship may contain nonlinearities or interactions, random forest is assessed alongside a regularized linear benchmark [24,25].

H5. A combined digital-competence and digital-tool-usage feature representation predicts perceived technology-enabled task efficiency more accurately than either feature group alone.

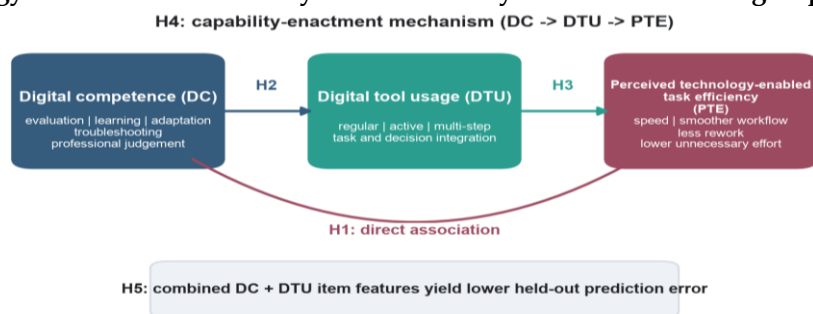


Figure 1. Proposed capability-enactment model. Arrows denote hypothesized associations rather than established causal effects.

3. Method

3.1. Research design and data collection

The study used a cross-sectional questionnaire design. The intended population was banking-sector employees whose work involved digital or AI-enabled tools. The questionnaire was distributed through WeChat from 24 to 27 March 2026. All 301 records in the raw export contained responses to the twenty substantive items, and the cleaned analytical file retained all 301 cases. The sampling design was non-probabilistic; consequently, the sample should be described as a banking-sector convenience sample rather than as a representative sample of the national banking workforce.

The raw export contained response timestamps, completion duration, source channel, and an IP-derived string. IP-derived information was used only to check for repeated submission patterns and was excluded from the analytical data. The analytical file contains only the twenty item responses and three derived construct means. The participating banking population and country were not encoded in the files. They must be completed from the authorized study record before submission: **[banking population and country/region]**.

The export contained no demographic variables. Age, gender, education, organizational tenure, job level, bank type, and functional role cannot be reported or used as controls. This limitation is disclosed rather than repaired through inference from IP location or other indirect identifiers. It narrows the conclusions to the observed construct relationships and prevents demographic subgroup analysis.

Ethical approval and consent information must match the original protocol. The final manuscript should state: **[full ethics committee name, approval or exemption number, and approval date]**. Participants should also be described according to the authorized consent record: **[verified electronic consent wording and confidentiality procedure]**. These fields are retained because an institutional name or approval number cannot be reconstructed from the questionnaire export.

3.2. Measures

The questionnaire used a five-point response scale coded from 1 to 5. The item wording was task oriented and referred to “digital or AI-enabled tools” so that respondents could answer across the heterogeneous applications used in their work. Construct scores used for descriptive summaries were the arithmetic means of the relevant items. The PLS-SEM analysis estimated latent-variable scores from the reflective measurement model rather than treating the arithmetic means as error-free measures.

Digital competence. DC was measured with seven items. They represented digital information search, recommendation reliability, rapid learning, troubleshooting, adaptation, independent judgement, and integration of professional judgement with tool output. The content reflects workplace definitions that combine operational capability with critical and adaptive dimensions [5-7].

Digital tool usage. DTU was measured with seven items. They assessed regular use, multi-step integration, importance to task completion, active task execution, decision support, and continued use under complexity. This operationalization follows the system-use literature by representing task-embedded behavior rather than behavioral intention or access [9,10].

Perceived technology-enabled task efficiency. PTE was measured with six items. They covered faster completion, less routine time, smoother processes, efficient work steps, reduced rework, and lower unnecessary effort. The construct is aligned with the task-productivity dimension of perceived information-technology impact [13]. It does not measure objective throughput, accuracy, service quality, creativity, or overall job performance.

Table 1. Construct operationalization.

Construct	Items	Representative Content	Scale	Conceptual Basis
Digital competence (DC)	7	Search and evaluate information; learn and troubleshoot tools; adapt routines; integrate	1–5	Workplace digital competence [5–7]

Digital tool usage (DTU)	7	professional judgement Regular and active use; multi-step integration; decision support; continued use under complexity Faster completion; smoother workflow; less repetition, rework, and unnecessary effort	1–5	Task-embedded system use [9,10]
Perceived technology-enabled task efficiency (PTE)	6		1–5	Perceived task productivity [13]

3.3. Response-quality assessment

All 301 twenty-item response vectors were unique, complete, and non-constant. The median completion time was 71 seconds (IQR 64-82); 34 records were under 60 seconds, but none was below half the sample median. Three long identical-response runs were retained because they were not complete straight lines, duplicates, or jointly extreme in duration. Three IP-derived strings appeared twice, but their response vectors differed; shared network identifiers were not treated as proof of duplicate participation. No record was removed on a single noisy quality indicator.

3.4. Explanatory analysis

The explanatory analysis used reflective PLS-SEM for prediction-oriented latent constructs measured with multiple indicators. It did not require multivariate normality. All three blocks were estimated in Mode A. The measurement model used loadings, Cronbach’s alpha, composite reliability (rhoC), average variance extracted (AVE), and the heterotrait-monotrait ratio (HTMT). Thresholds were approximately 0.70 for loadings and reliability, 0.50 for AVE, and 0.85 for HTMT [20,22].

The structural model contained paths from DC to PTE, DC to DTU, and DTU to PTE. Endogenous explained variance was summarized with R2. The indirect association was the product of the DC-to-DTU and DTU-to-PTE coefficients. Uncertainty used 5,000 nonparametric bootstrap resamples and percentile intervals [23]. An effect was supported when its 95% interval excluded zero. Temporal ordering was not observed, and all constructs were self-reported once. The result is therefore an **indirect association** or **statistical mediation**, not a causal mediated effect.

3.5. Predictive analysis

The predictive analysis used item-level features rather than PLS construct scores. The competence-only representation contained DC1-DC7, the usage-only representation contained DTU1-DTU7, and the combined representation contained all fourteen items. The prediction target was the mean of PTE1-PTE6. No outcome item was included among the predictors.

Ridge regression served as a regularized linear benchmark. Random forest allowed nonlinear relations and interactions without prescribing their form [24]. Repeated nested cross-validation used five outer folds and three repeats, producing fifteen held-out evaluations. Three-fold inner cross-validation selected hyperparameters within each outer training partition. Preprocessing and selection were confined to training data. The same outer splits supported paired comparison of every feature representation and algorithm.

Root mean squared error (RMSE) was the primary loss measure; mean absolute error (MAE) and fold-level R2 were secondary metrics. Pooled Q2predict compared squared prediction error with a naive benchmark that assigned each test case the mean PTE score of the corresponding training set [27]. Values above zero indicate improvement over the benchmark. Paired Wilcoxon signed-rank tests compared the combined random forest with competence-only and usage-only forests across the fifteen matched outer evaluations.

Permutation importance was calculated only in held-out data. Item-level permutation estimates can be diluted when reflective indicators are correlated, because an unpermuted item can substitute for a permuted item. The analysis therefore also jointly permuted all DC items or all DTU items. The resulting increase in held-out RMSE represents feature-group information within the fitted predictive model; it is not an estimate of causal importance.

All analyses were conducted in Python with fixed random seeds. The reproducible pipeline writes the measurement, structural, bootstrap, cross-validation, model-comparison, and permutation-importance outputs used in the tables and figures.

4. Results

4.1. Descriptive results and measurement model

The construct means were 3.302 for DC (SD = 0.932), 3.334 for DTU (SD = 0.932), and 3.426 for PTE (SD = 0.954). These values indicate that the sample covered the full substantive range rather than clustering uniformly at the upper endpoint.

The reflective measurement model showed adequate indicator reliability and convergence (Table 2). DC loadings ranged from 0.784 to 0.810, DTU loadings from 0.746 to 0.832, and PTE loadings from 0.783 to 0.854. Cronbach's alpha ranged from 0.898 to 0.907, composite reliability ranged from 0.921 to 0.926, and AVE ranged from 0.639 to 0.662. HTMT values were 0.425 for DC-DTU, 0.419 for DC-PTE, and 0.444 for DTU-PTE. All were well below 0.85, supporting empirical separation of capability, enactment, and perceived efficiency [22].

Table 2. Reflective measurement quality and discriminant validity.

Construct	Mean (SD)	Loading range	Cronbach's alpha	rhoC	AVE	Maximum HTMT
DC	3.302 (0.932)	0.784–0.810	0.906	0.925	0.639	0.425
DTU	3.334 (0.932)	0.746–0.832	0.907	0.926	0.643	0.444
PTE	3.426 (0.954)	0.783–0.854	0.898	0.921	0.662	0.444

Pairwise HTMT: DC-DTU = 0.425; DC-PTE = 0.419; DTU-PTE = 0.444.

4.2. Structural model and hypothesis tests

The structural results are reported in Table 3 and Figure 2. DC was positively associated with PTE (beta = 0.261, 95% CI 0.160–0.361), supporting H1. DC was also positively associated with DTU (beta = 0.388, 95% CI 0.296–0.486), supporting H2. After DC was included in the model, DTU retained a positive association with PTE (beta = 0.303, 95% CI 0.204–0.403), supporting H3. The two predictors had a score correlation of 0.388, corresponding to a structural VIF of approximately 1.18, which does not indicate problematic predictor collinearity.

The model explained 15.0% of the variance in DTU and 22.2% of the variance in PTE. These R² values are substantively meaningful for a parsimonious behavioral model but leave most individual variation unexplained. They should not be interpreted as a complete account of workflow efficiency.

The indirect association from DC through DTU to PTE was 0.118, and its 95% bootstrap interval excluded zero (0.074–0.171). H4 was therefore supported. The direct path from DC to PTE remained positive after DTU was introduced, while the total association was 0.379 (95% CI 0.287–0.473). This pattern is consistent with partial statistical mediation: DTU carries part, but not all, of the competence-efficiency relationship.

Table 3. Structural effects and hypothesis decisions.

Hypothesis/effect	Estimate	Bootstrap SE	95% percentile CI	Decision
H1: DC → PTE	0.261	0.051	0.160–0.361	Supported

H2: DC → DTU	0.388	0.049	0.296–0.486	Supported
H3: DTU → PTE	0.303	0.051	0.204–0.403	Supported
H4: DC → DTU → PTE	0.118	0.025	0.074–0.171	Supported
Total DC–PTE association	0.379	0.047	0.287–0.473	–

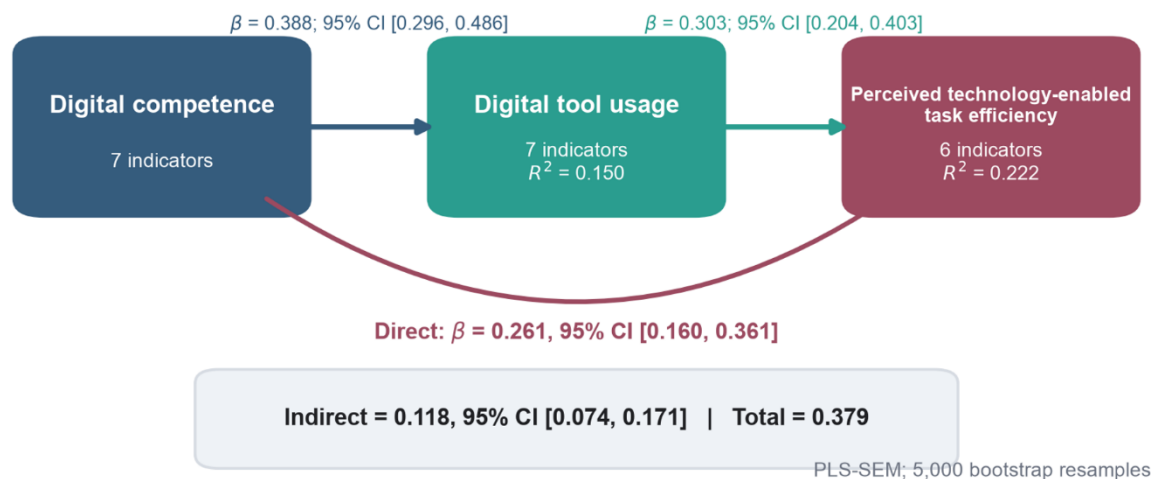


Figure 2. Estimated structural model. Coefficients are standardized PLS-SEM estimates; intervals are percentile bootstrap intervals from 5,000 resamples.

4.3. Out-of-sample prediction

All fitted models achieved positive pooled Q²predict values, indicating lower pooled squared error than the training-mean benchmark (Table 4). Performance nevertheless differed by algorithm and feature representation. For competence-only features, the random forest achieved mean RMSE = 0.863 and pooled Q²predict = 0.186. The usage-only random forest achieved RMSE = 0.849 and Q²predict = 0.213. The combined random forest performed best, with RMSE = 0.835, MAE = 0.709, mean fold R² = 0.208, and pooled Q²predict = 0.240.

Ridge regression showed the same ordering across feature sets. The combined ridge produced RMSE = 0.855 and Q²predict = 0.203. Corresponding values were 0.894 and 0.130 for competence only, and 0.886 and 0.143 for usage only. The random forest showed a modest advantage over ridge in the combined representation. Nonlinearities or interactions may therefore contribute without dominating the signal.

Table 4. Repeated nested cross-validated predictive performance

Feature set	Model	RMSE, mean (SD)	MAE, mean (SD)	Mean fold R ²	Pooled Q ² predict
DC only	Ridge	0.894 (0.038)	0.762 (0.032)	0.093	0.130
DC only	Random forest	0.863 (0.053)	0.737 (0.044)	0.152	0.186
DTU only	Ridge	0.886 (0.054)	0.751 (0.044)	0.108	0.143
DTU only	Random forest	0.849 (0.049)	0.719 (0.034)	0.181	0.213
DC + DTU	Ridge	0.855 (0.047)	0.720 (0.036)	0.170	0.203
DC + DTU	Random forest	0.835 (0.045)	0.709 (0.033)	0.208	0.240

Paired comparisons used the same fifteen outer test partitions. The combined random forest reduced mean RMSE by 0.029 relative to the competence-only forest (Wilcoxon $p = 0.005$). It reduced RMSE by 0.015 relative to the usage-only forest ($p = 0.018$). H5 was therefore supported. Joint permutation of DTU increased held-out RMSE by 0.107 (SD = 0.034). Permuting DC increased RMSE by 0.050 (SD = 0.019). DTU was the stronger proximal group, but competence remained informative because the combined representation outperformed usage alone.

Table 5. Paired comparisons and feature-group importance

Analysis	Estimate	Variability/test	Interpretation
Combined RF vs DC-only RF	RMSE reduction 0.029	Wilcoxon $p = 0.005$	Combined features improved prediction
Combined RF vs DTU-only RF	RMSE reduction 0.015	Wilcoxon $p = 0.018$	Competence retained incremental signal
Joint permutation: DC	RMSE increase 0.050	SD = 0.019 across 15 folds	Capability group was informative
Joint permutation: DTU	RMSE increase 0.107	SD = 0.034 across 15 folds	Usage was the stronger proximal group

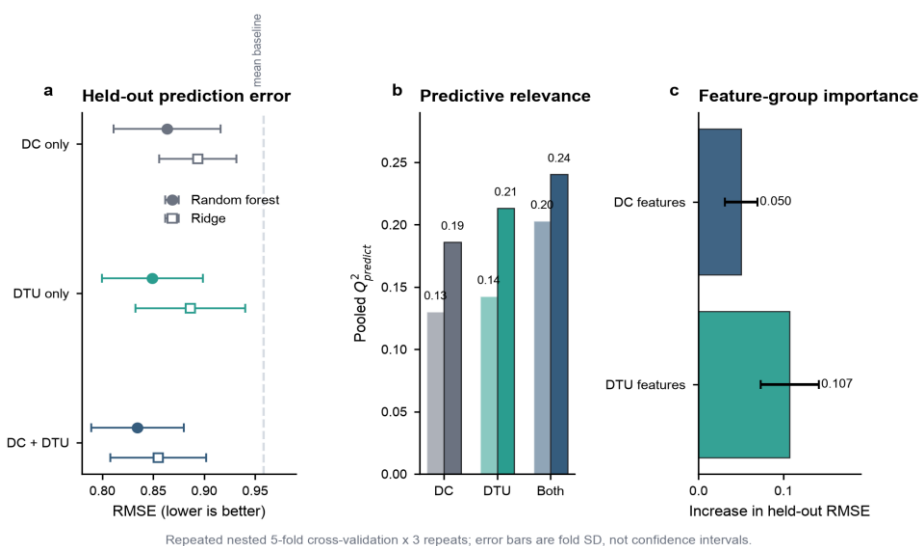


Figure 3. Out-of-sample predictive results. (a) RMSE across feature sets and algorithms; error bars show fold standard deviations. (b) Pooled $Q^2_{predict}$. (c) Joint feature-group permutation importance in held-out folds.

5. Discussion

This study asked whether digital tool usage translates digital competence into perceived workflow efficiency and whether the same construct representation carries information that generalizes beyond the estimation sample. The results support both parts of the argument. DC was positively associated with PTE, DC was positively associated with DTU, and DTU retained a positive association with PTE after DC was controlled. DTU carried a non-zero indirect association, while the direct DC-PTE path remained positive. In prediction, combined competence and usage features produced lower held-out error than either group alone. The structural and predictive evidence therefore converge on a capability-enactment interpretation, although neither analysis establishes causality or objective productivity.

H1 and H2 locate DC at both the outcome and behavioral levels. Evaluative and adaptive capability can reduce wasted effort directly, while also helping a user turn an interface into a repeatable work procedure [5-8,14,16,17]. Availability alone does not produce this integration. Competence is especially relevant when a task requires iterative evaluation, correction, or a decision to reject automation.

H3 and H4 identify DTU as the proximal layer without implying that “more use is better.” The indicators represent active, regular, multi-step, task-relevant integration rather than raw frequency [9,10]. The indirect association shows how capability can become relevant to a workflow. The remaining direct

path indicates that search strategy, judgement, and error recognition can reduce effort without increasing reported integration. The evidence therefore favors a layered capability-enactment mechanism rather than a replacement model in which use renders competence irrelevant.

The predictive findings sharpen this interpretation. Usage-only features predicted PTE better than competence-only features. Joint permutation of DTU also created approximately twice the held-out RMSE increase produced by permuting DC. This pattern is expected if usage is closer to the measured outcome. Yet the combined representation remained superior to usage alone. Competence therefore contained incremental information not recoverable from current usage responses. Structural mediation identifies a meaningful pathway, while cross-validated prediction shows that both layers carry reproducible information.

Absolute prediction remained moderate. A pooled Q2predict of 0.240 indicates improvement over the naive benchmark, not readiness for operational deployment. Most variation in PTE remains unmodelled. The result is useful because it constrains the engineering claim. Deployment would require task context, system quality, latency, error rates, experience, workload, telemetry, sequence logs, and objective outcomes. Survey items alone cannot support high-stakes employee scoring.

5.1. Theoretical implications

The first theoretical contribution is the separation of capability from enactment. Digital-competence research often treats skill acquisition as if its organizational return follows automatically. System-use research, by contrast, can focus on behavior without adequately representing the competence needed to make that behavior effective. The present model joins these streams while preserving their distinction. DC represents what an employee can evaluate, learn, adapt, and troubleshoot; DTU represents what is incorporated into task execution. Low HTMT values and different predictive importance provide empirical support for keeping the constructs separate.

The second contribution is the capability-enactment mechanism. Mediation studies have shown that psychological empowerment can partly connect digital competence with job performance [14]. The present model moves to a behavioral mechanism that is closer to human-computer interaction: competence becomes consequential when it is expressed through task-embedded system use. Psychological and behavioral mediators are not mutually exclusive. Future models could test whether competence encourages both empowered work states and deeper tool integration, or whether organizational constraints determine which pathway dominates.

The third contribution is outcome precision. PTE is one bounded component of individual performance, not a proxy for the complete construct. Digital tools can accelerate a task while reducing accuracy, increasing verification burden, or shifting hidden work to another person. Information-system impact research has long separated task productivity from innovation, control, and customer consequences [13]. Preserving this boundary makes the positive result more credible because it avoids claiming organizational value that the measurement design cannot observe.

The fourth contribution is methodological. PLS-SEM evaluates measurement and the proposed indirect association; nested cross-validation estimates generalization error; and feature-set ablation tests the theory in predictive form. This alignment is relevant to human-AI research, where behavioral models often lack held-out validation and predictive models often lack a defensible theory of representation [18,19].

5.2. Engineering and managerial implications

The model suggests that digital transformation should not be evaluated through software deployment or training completion alone. Organizations need evidence that competence is translated into appropriate workflow behavior. A useful evaluation stack would therefore include three levels: capability assessment, task-level usage telemetry, and independently measured work outcomes. Capability measures can identify whether users understand information quality, system limits, and recovery procedures. Usage telemetry can show where a tool enters the workflow, which functions are used, and where repeated correction occurs. Outcome measurement can test whether speed gains coexist with accuracy, compliance, and service quality.

Training design should likewise move from feature demonstration to scenario-based enactment. Employees need opportunities to practice tool selection, output verification, exception handling,

escalation, and integration with professional judgement. The direct competence path indicates that these capabilities have value beyond increasing usage. A training program should not therefore be judged only by adoption rate. It should assess whether users can recognize when not to use a system and whether they can preserve accountability when a recommendation is uncertain.

The stronger predictive contribution of DTU implies that system designers should observe the workflow rather than assume that skill inventories fully represent human-AI performance. Interfaces can support enactment through provenance cues, uncertainty communication, editable intermediate outputs, recovery paths, and frictionless transfer between task stages. At the same time, the moderate prediction warns against individual-level surveillance or automated personnel decisions. The present model can inform workflow diagnostics and research design, but it is not validated for employee ranking, promotion, or risk scoring.

5.3. Limitations and future research

Several limitations define the boundary of the evidence. First, the design is cross-sectional. The order DC → DTU → PTE is theoretically plausible but not temporally observed. Reverse or reciprocal relationships are possible: employees who experience efficiency gains may use tools more frequently and may subsequently develop competence. Longitudinal panel designs, staged training studies, or field experiments are needed to test within-person change and temporal mediation.

Second, all constructs and the outcome were self-reported in one questionnaire. Common method variance, consistency motives, and positive self-assessment may inflate associations [26]. The low HTMT values show that respondents differentiated the constructs, but they do not eliminate common-source bias. Future studies should combine competence tests, system logs, supervisor or peer ratings, and objective process measures such as cycle time, correction rate, throughput, and error severity.

Third, the outcome concerns perceived technology-enabled efficiency. No conclusion can be drawn about work quality, customer satisfaction, regulatory compliance, creativity, contextual performance, or objective productivity. A faster workflow can still produce an inferior outcome. Future research should treat efficiency and quality as joint outcomes and explicitly model the verification cost created by AI-generated content.

Fourth, recruitment occurred through WeChat and the export contains no demographic or organizational variables. The authorized country, banking population, and institutional context must be restored before submission. The convenience sample does not support population-wide prevalence claims; replication should use multi-bank, multi-region, role-stratified samples.

Fifth, the item scales require further content-validity and measurement-invariance testing. The meaning of DTU may vary across roles, tools, and tasks, and its behavioral dimensions may be formative in some settings. Alternative measurement specifications should therefore be compared.

Finally, nested cross-validation reduces optimistic evaluation but does not establish transportability. Future work should use an external test set, calibrate predictions, quantify uncertainty, and test whether workflow telemetry improves on the DC-DTU survey representation.

6. Conclusion

Digital competence was associated with perceived technology-enabled task efficiency both directly and indirectly through task-embedded digital tool usage. Usage was the more proximal explanatory and predictive layer, but competence retained incremental information. The combined representation consequently produced the strongest held-out prediction, although its absolute performance remained moderate. These findings support a capability-enactment account of AI-augmented work: human capability matters, but organizational return depends partly on whether that capability becomes appropriate workflow behavior. The implication is not to maximize tool use. It is to design, train, and evaluate human-AI systems across capability, enactment, and independently measured outcomes.

References

1. Jarrahi, M. H. Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making. *Business Horizons* 61, 577-586 (2018). <https://doi.org/10.1016/j.bushor.2018.03.007>
2. Langer, M. & Landers, R. N. The future of artificial intelligence at work: A review on effects of decision automation and augmentation on workers targeted by algorithms and third-party observers. *Computers in Human Behavior* 123, 106878 (2021). <https://doi.org/10.1016/j.chb.2021.106878>
3. Raisch, S. & Krakowski, S. Artificial intelligence and management: The automation-augmentation paradox. *Academy of Management Review* 46, 192-210 (2021). <https://doi.org/10.5465/amr.2018.0072>
4. Dellermann, D., Ebel, P., Sollner, M. & Leimeister, J. M. Hybrid intelligence. *Business & Information Systems Engineering* 61, 637-643 (2019). <https://doi.org/10.1007/s12599-019-00595-2>
5. Oberlander, M., Beinicke, A. & Bipp, T. Digital competencies: A review of the literature and applications in the workplace. *Computers & Education* 146, 103752 (2020). <https://doi.org/10.1016/j.compedu.2019.103752>
6. van Laar, E., van Deursen, A. J. A. M., van Dijk, J. A. G. M. & de Haan, J. The relation between 21st-century skills and digital skills: A systematic literature review. *Computers in Human Behavior* 72, 577-588 (2017). <https://doi.org/10.1016/j.chb.2017.03.010>
7. Audrin, B., Audrin, C. & Salamin, X. Digital skills at work: Conceptual development and empirical validation of a measurement scale. *Technological Forecasting and Social Change* 202, 123279 (2024). <https://doi.org/10.1016/j.techfore.2024.123279>
8. Moravec, V., Hynek, N., Gavurova, B. & Rigelsky, M. Who uses it and for what purpose? The role of digital literacy in ChatGPT adoption and utilisation. *Journal of Innovation & Knowledge* 9, 100602 (2024). <https://doi.org/10.1016/j.jik.2024.100602>
9. Burton-Jones, A. & Straub, D. W. Reconceptualizing system usage: An approach and empirical test. *Information Systems Research* 17, 228-246 (2006). <https://doi.org/10.1287/isre.1060.0096>
10. Doll, W. J. & Torkzadeh, G. Developing a multidimensional measure of system-use in an organizational context. *Information & Management* 33, 171-185 (1998). [https://doi.org/10.1016/S0378-7206\(98\)00028-7](https://doi.org/10.1016/S0378-7206(98)00028-7)
11. DeLone, W. H. & McLean, E. R. The DeLone and McLean model of information systems success: A ten-year update. *Journal of Management Information Systems* 19, 9-30 (2003). <https://doi.org/10.1080/07421222.2003.11045748>
12. Goodhue, D. L. & Thompson, R. L. Task-technology fit and individual performance. *MIS Quarterly* 19, 213-236 (1995). <https://doi.org/10.2307/249689>
13. Torkzadeh, G. & Doll, W. J. The development of a tool for measuring the perceived impact of information technology on work. *Omega* 27, 327-339 (1999). [https://doi.org/10.1016/S0305-0483\(98\)00049-8](https://doi.org/10.1016/S0305-0483(98)00049-8)
14. Ochoa Pacheco, P. & Coello-Montecel, D. Does psychological empowerment mediate the relationship between digital competencies and job performance? *Computers in Human Behavior* 140, 107575 (2023). <https://doi.org/10.1016/j.chb.2022.107575>
15. Pitafi, A. H., Kanwal, S., Ali, A., Khan, A. N. & Ameen, M. W. Moderating roles of IT competency and work cooperation on employee work performance in an ESM environment. *Technology in Society* 55, 199-208 (2018). <https://doi.org/10.1016/j.techsoc.2018.08.002>
16. Blanka, C., Krumay, B. & Rueckel, D. The interplay of digital transformation and employee competency: A design science approach. *Technological Forecasting and Social Change* 178, 121575 (2022). <https://doi.org/10.1016/j.techfore.2022.121575>
17. Peng, L. *et al.* Human-AI collaboration: Unraveling the effects of user proficiency and AI agent capability in intelligent decision support systems. *International Journal of Industrial Ergonomics* 103, 103629 (2024). <https://doi.org/10.1016/j.ergon.2024.103629>
18. Shmueli, G. & Koppius, O. R. Predictive analytics in information systems research. *MIS Quarterly* 35, 553-572 (2011). <http://hdl.handle.net/1765/25837>
19. Richter, N. F. & Tudoran, A. A. Elevating theoretical insight and predictive accuracy in business research: Combining PLS-SEM and selected machine learning algorithms. *Journal of Business Research* 173, 114453 (2024). <https://doi.org/10.1016/j.jbusres.2023.114453>
20. Hair, J. F., Hult, G. T. M., Ringle, C. M., Sarstedt, M., Danks, N. P. & Ray, S. *Partial Least Squares Structural Equation Modeling (PLS-SEM) Using R: A Workbook*. (Springer, 2021). <https://doi.org/10.1007/978-3-030-80519-7>
21. Becker, G. S. *Human Capital: A Theoretical and Empirical Analysis, with Special Reference to Education*, 3rd edn. (University of Chicago Press, 1993). <https://press.uchicago.edu/ucp/books/book/chicago/H/bo3684031>

22. Henseler, J., Ringle, C. M. & Sarstedt, M. A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science* 43, 115-135 (2015). <https://doi.org/10.1007/s11747-014-0403-8>
23. Preacher, K. J. & Hayes, A. F. Asymptotic and resampling strategies for assessing and comparing indirect effects in multiple mediator models. *Behavior Research Methods* 40, 879-891 (2008). <https://doi.org/10.3758/BRM.40.3.879>
24. Breiman, L. Random forests. *Machine Learning* 45, 5-32 (2001). <https://doi.org/10.1023/A:1010933404324>
25. Pedregosa, F. *et al.* Scikit-learn: Machine learning in Python. *Journal of Machine Learning Research* 12, 2825-2830 (2011). <https://jmlr.org/papers/v12/pedregosa11a.html>
26. Podsakoff, P. M., MacKenzie, S. B., Lee, J.-Y. & Podsakoff, N. P. Common method biases in behavioral research: A critical review of the literature and recommended remedies. *Journal of Applied Psychology* 88, 879-903 (2003). <https://doi.org/10.1037/0021-9010.88.5.879>
27. Shmueli, G. *et al.* Predictive model assessment in PLS-SEM: Guidelines for using PLSpredict. *European Journal of Marketing* 53, 2322-2347 (2019). <https://doi.org/10.1108/EJM-02-2019-0189>