

The Influence of AI Tool Acceptance on Achievement Levels in Youth Robotics Competitions: A Case Study of Innovedex 2026

1st Puvasit Aphisinrungrot
King Mongkut's Institute of
Technology Ladkrabang International
Demonstration School (KMIDS)
Bangkok, Thailand
jaonai1301@gmail.com

2nd Sirin Liukasemsarn
Darunsikkhalai School for
Innovative Learning, King Mongkut's
University of Technology Thonburi
Bangkok, Thailand
angiesirin09@gmail.com

3rd Chacharin Lertyosbordin*
School of Information Technology,
King Mongkut's Institute of
Technology Ladkrabang
Bangkok, Thailand
chacharin.le@kmitl.ac.th

Abstract—This study examined whether AI tool acceptance—assessed through the Unified Theory of Acceptance and Use of Technology (UTAUT)—predicts achievement level among youth competitors in the Innovedex 2026 robotics competition. A cross-sectional survey of 297 competitors measured four UTAUT constructs: Performance Expectancy (PE), Effort Expectancy (EE), Social Influence (SI), and Facilitating Conditions (FC). Results showed generally high acceptance, led by Social Influence and Effort Expectancy, while Performance Expectancy was lowest. All four constructs correlated significantly and positively with award level, with Facilitating Conditions and Performance Expectancy showing the strongest associations. Because ceiling effects violated linear regression assumptions, a multinomial logistic regression model was used instead, achieving 68.7% cross-validation accuracy against a 46.8% baseline. Findings suggest that genuine confidence in AI's practical benefit and adequate technical support, rather than mere ease of use, most strongly drive competitive success, offering actionable guidance for competition organizers and directions for future replication research.

Keywords— Artificial intelligence tools, UTAUT, Technology acceptance, Robotics competition

I. INTRODUCTION

Artificial intelligence (AI) tools such as ChatGPT, Gemini, and Claude have come to assume an increasingly prominent role in assisting engineers with the design and programming of robots [1]. Within the specific context of youth robotics competitions, however, no study to date has established with any certainty whether the acceptance and use of these AI tools bears any relationship to competitors' levels of achievement or to the awards they ultimately receive.

To address this gap in the literature, the present study draws on the Innovedex 2026 national robotics and artificial intelligence competition — organized by the present research team under the theme “AI & Robotics System Integration” [2] — as its case study. The study analyzes competitors' acceptance of AI tools through the lens of the Unified Theory of Acceptance and Use of Technology (UTAUT) [3], which is assessed across four constructs: Performance Expectancy (PE), Effort Expectancy (EE), Social Influence (SI), and Facilitating Conditions (FC).

Accordingly, this study pursues three objectives, each formalized as a corresponding research question and summarized in Table I: to establish competitors' baseline level of AI tool acceptance across the four UTAUT constructs; to determine whether this acceptance is associated with the achievement level attained in competition; and to develop

a model capable of predicting competitors' award level from their UTAUT profile.

TABLE I. RESEARCH QUESTIONS AND RESEARCH OBJECTIVES

Research Question	Research Objective
RQ1: What is the level of AI tool acceptance — across Performance Expectancy, Effort Expectancy, Social Influence, and Facilitating Conditions?	OBJ1: Establish competitors' baseline level of AI tool acceptance across the four UTAUT constructs.
RQ2: Is competitors' acceptance of AI tools significantly associated with the award level they achieve?	OBJ2: Determine whether competitors' acceptance of AI tools is associated with their achievement level in the competition.
RQ3: Can a predictive model based on the four UTAUT constructs accurately classify competitors' award level?	OBJ3: Develop a model capable of predicting competitors' award level from their UTAUT profile.

To visualize what this study sets out to establish across all three research questions, Fig. 1 summarizes the conceptual model: the four UTAUT constructs (PE, EE, SI, FC) are measured (RQ1), examined for their association with award level (RQ2), and used to predict it (RQ3).

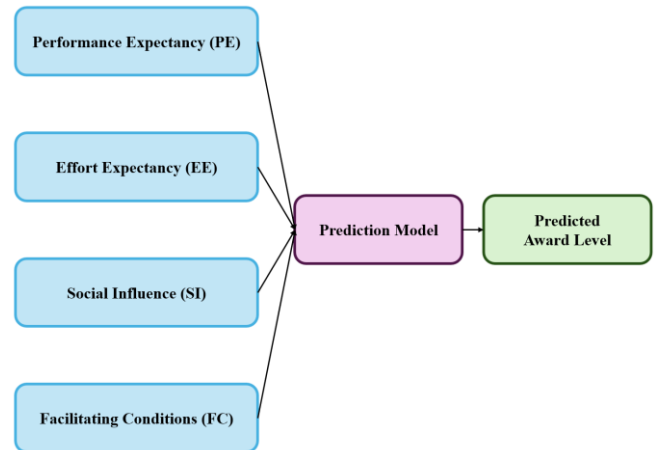


Fig. 1. Conceptual model of the study

II. RELATED THEORETICAL BACKGROUND

The Unified Theory of Acceptance and Use of Technology (UTAUT) was developed by Venkatesh [3] through the synthesis of eight influential technology acceptance models, with the aim of explaining users' technology-related behavior and behavioral intention. UTAUT proposes four core constructs that shape technology adoption decisions:

- 1) **Performance Expectancy (PE)** refers to an individual's belief that a given technology will enhance task performance. In the context of this study, this construct captures competitors' belief that AI tools will strengthen their capacity to design robots, resolve mission-related problems, and increase their chances of winning an award.
- 2) **Effort Expectancy (EE)** refers to the belief that a technology can be understood and used with ease and without undue complexity. In this study, this construct reflects competitors' perception that AI tools can be learned quickly (for example, through prompting to generate code or retrieve information) and applied to robot development conveniently, without requiring advanced expertise.
- 3) **Social Influence (SI)** refers to an individual's perception of encouragement from significant others. In this study, this construct reflects competitors' sense that the people around them — such as team-supervising teachers or peers at the same institution — encourage, recommend, or expect them to use AI tools as an aid in robot development.
- 4) **Facilitating Conditions (FC)** refers to the belief that adequate infrastructure and support exist to enable technology use. In this study, this construct captures the extent to which competitors received support in the form of computing equipment, AI tool accounts, and robotics hardware sufficient to test and implement AI-generated outputs in practice.

Since AI tools began to see widespread adoption, UTAUT and its extension, UTAUT2 [4], have been applied with increasing frequency to investigate learners' AI usage behavior. For instance, a recent study applying the UTAUT model to examine AI tool acceptance among Chinese university students found that Social Influence (SI) was the strongest predictor of behavioral intention to use AI, followed by Performance Expectancy (PE), whereas Facilitating Conditions (FC) exerted no statistically significant effect [5].

That study, however, measured only the “intention” to use AI within a general learning context, leaving unresolved whether this same pattern of relationships would hold — or shift — once users are required to apply AI tools to hands-on tasks involving direct technical and equipment-related constraints, such as robot development. The present study therefore uses the Innovedex 2026 robotics competition, in which competitors must perform integrated work with physical hardware, as a case study through which to examine this question empirically.

III. METHODOLOGY

A. Design and Sample

This study employed a cross-sectional quantitative design. The population comprised competitors in the Innovedex 2026 competition across all four regional qualifying rounds. Sampling was conducted through non-probability, voluntary sampling: all competitors were invited to take part in the study, and those who consented and completed the questionnaire in full constituted the final sample of 297 respondents. Demographically, respondents ranged in age from 14 to 19 years ($M = 16.40$, $SD = 1.69$), and

56.2% were male. With respect to competition outcomes, 46.8% of respondents received a gold award, 23.9% a silver award, 16.8% a bronze award, and 12.5% a participation-only award.

B. Instrument

AI tool acceptance was measured using a 12-item, 5-point Likert-type rating scale (1 = strongly disagree, 5 = strongly agree; full item wording is presented in Table II), comprising three items per construct and adapted from [3]. Internal consistency of all four subscales proved satisfactory, with Cronbach's alpha exceeding 0.70 for every subscale — surpassing the threshold recommended in [6].

TABLE II. AI TOOL ACCEPTANCE SURVEY ITEMS

Item	Question
PE1	The AI tool helps me refine my robot design more quickly.
PE2	The AI tool helps me complete robot programming tasks more quickly.
PE3	Using AI increases my team's chances of winning an award in the competition.
EE1	The AI tool's design is clear, so I know intuitively how to use it.
EE2	I can learn to use the AI tools needed for robot programming within a short time.
EE3	I can produce output from the AI tool without needing to memorize complex procedures.
SI1	My supervising teacher encourages me to use AI in robot development.
SI2	My teammates believe AI is important for the competition.
SI3	Experts recommend using AI in robot development work.
FC1	My school provides sufficient AI tools for my robot development work.
FC2	I have received adequate training or instruction on using AI tools for my needs.
FC3	Technical support is available to help me when I encounter problems using AI tools.

C. Analysis

To address Research Question 1, descriptive statistics were computed and Cronbach's alpha reliability was tested for each UTAUT construct. To address Research Question 2, a full Spearman's rank correlation analysis was conducted among all four constructs and the four award levels.

To address Research Question 3, the analysis initially proceeded via Ordinary Least Squares (OLS) regression; the data, however, significantly violated the assumptions of linearity [7] and homoscedasticity [8] ($p < .001$), a consequence of a ceiling effect whereby 43.8% of the sample achieved the maximum possible score of 100. Because data transformation techniques could not remedy this limitation, the OLS approach was abandoned. The analysis was instead reformulated as a multinomial logistic regression [9] modeling the four award levels against the four predictor constructs. Diagnostic checks confirmed that variance inflation factors (VIF) for all predictors remained below 1.3, indicating no meaningful risk of multicollinearity.

Model performance was evaluated using 5-fold stratified cross-validation, with results reported in terms of accuracy, macro-F1, and quadratic-weighted Cohen's kappa, benchmarked against the majority-class baseline of 46.8%. McFadden's pseudo- R^2 [10] was additionally reported to assess overall model fit. All analyses were conducted in Python using the scikit-learn [11] and statsmodels [12] libraries.

IV. RESULTS

A. Level of AI Tool Acceptance among Competitors (RQ 1)

Table III presents descriptive statistics for the four UTAUT constructs. All construct means exceeded 3.00 (the scale midpoint), indicating a relatively high level of AI tool acceptance among competitors prior to the start of the competition. Social Influence (SI) recorded the highest mean ($M = 4.45$, $SD = 0.54$), followed by Effort Expectancy (EE) ($M = 4.29$, $SD = 0.54$), suggesting that competitors perceived strong encouragement from those around them and regarded the AI tools as easy to use.

Facilitating Conditions (FC) fell within a moderate range ($M = 3.93$, $SD = 0.72$), while Performance Expectancy (PE) recorded the lowest mean ($M = 3.55$, $SD = 0.78$), indicating that, prior to the competition, competitors held comparatively limited confidence that AI tools would concretely elevate the quality of their work. All four subscales demonstrated good internal consistency.

TABLE III. DESCRIPTIVE STATISTICS AND RELIABILITY OF UTAUT CONSTRUCTS (N = 297)

Construct	Mean	SD	Cronbach's α
PE	3.55	0.78	.895
EE	4.29	0.54	.812
SI	4.45	0.54	.831
FC	3.93	0.72	.894

B. Relationship between Technology Acceptance and Award Level (RQ2)

Table IV presents the full Spearman's rank correlation matrix among the four UTAUT constructs and award level. All four constructs were significantly and positively correlated with award level (PE: $\rho = .42$; EE: $\rho = .19$; SI: $\rho = .39$; FC: $\rho = .54$; all $p < .01$), with Facilitating Conditions (FC) exhibiting the strongest association and Effort Expectancy (EE) the weakest.

Furthermore, inter-predictor correlations were weak to moderate in magnitude ($|\rho| \leq .32$), consistent with the low VIF values reported in Section III-C, indicating no significant multicollinearity among the predictors and thereby confirming their suitability for subsequent predictive modeling.

TABLE IV. SPEARMAN CORRELATION MATRIX AMONG UTAUT CONSTRUCTS AND AWARD LEVEL (N = 297)

	PE	EE	SI	FC
EE	-.26***			
SI	-.12*	.32***		
FC	-.05	-.04	.22***	
Award Level	.42***	.19**	.39***	.54***

Note: * $p < .05$; ** $p < .01$; *** $p < .001$

C. Predictive Equation for Award Level (RQ3)

Table V presents the coefficient weights of the multinomial logistic regression model, estimated using raw mean scores for the PE, EE, SI, and FC constructs (each on a 1–5 scale), with the “participation” award level specified as the reference group.

TABLE V. MULTINOMIAL LOGISTIC REGRESSION COEFFICIENTS (RELATIVE TO PARTICIPATION LEVEL, N = 297)

Predictor	Bronze	Silver	Gold
Intercept	-26.11	-80.76	-109.97
PE	2.43	6.64	8.36
EE	1.42	5.59	7.10
SI	0.98	3.41	5.44
FC	2.65	6.14	8.05

From these coefficients, the following log-odds (logit) equations can be derived to predict the probability of each of the three award levels relative to the participation level:

$$L_{Bronze} = -26.11 + 2.43 \cdot PE + 1.42 \cdot EE + 0.98 \cdot SI + 2.65 \cdot FC$$

$$L_{Silver} = -80.76 + 6.64 \cdot PE + 5.59 \cdot EE + 3.41 \cdot SI + 6.14 \cdot FC$$

$$L_{Gold} = -109.97 + 8.36 \cdot PE + 7.10 \cdot EE + 5.44 \cdot SI + 8.05 \cdot FC$$

Under this scheme, a respondent is classified into whichever award level corresponds to the highest L value among the three equations, or into the “participation” level if none of the three L values exceeds zero (with $L(\text{Participation})$ defined as 0). Notably, every coefficient across all three equations carries a positive sign, meaning that an increase in the raw score of any UTAUT construct can never lower the predicted award level — a pattern consistent with the positive associations reported in Table IV.

Moreover, Performance Expectancy (PE) and Facilitating Conditions (FC) consistently carried the two largest coefficient weights across all award levels, mirroring the strongest correlation magnitudes observed in Table IV, notwithstanding that PE recorded the lowest self-reported mean score in Table III.

D. Confusion Matrix and Model Evaluation

The logistic model classified award levels with a cross-validation accuracy of 68.7%, significantly exceeding the majority-class baseline of 46.8% (macro-F1 = .611; quadratic-weighted Cohen's kappa = .826; McFadden's pseudo- $R^2 = .472$; $p < .001$).

Table VI presents the confusion matrix derived from all 297 cases under the cross-validation procedure. The model correctly classified 91% of gold-award recipients, and no gold-award case was ever misclassified as participation- or bronze-level. Most misclassification errors involved only a single adjacent-class shift — for example, predicting “silver” where the true outcome was “gold.” Overall, 97.3% of all cases were classified either exactly correctly or with an error of no more than one adjacent award level.

TABLE VI. CONFUSION MATRIX (5-FOLD CROSS-VALIDATION, N = 297)

Actual \ Predicted	Participation	Bronze	Silver	Gold
	Participation	Bronze	Silver	Gold
Participation	21 (57%)	12 (32%)	4 (11%)	0 (0%)
Bronze	10 (20%)	29 (58%)	9 (18%)	2 (4%)
Silver	2 (3%)	10 (14%)	27 (38%)	32 (45%)
Gold	0 (0%)	0 (0%)	12 (9%)	127 (91%)

V. DISCUSSION

The findings indicate that Facilitating Conditions (FC) and Performance Expectancy (PE) exhibited the strongest associations with award level, even though PE recorded the lowest self-reported mean score and FC fell within a moderate range (cf. Table III). By contrast, although competitors reported high levels of perceived social support (SI) and ease of use (EE), both constructs proved less strongly associated with competitive success (cf. Table IV). This pattern suggests that genuine confidence that AI tools can tangibly elevate the quality of one's work, together with adequate equipment and technical support, are more decisive determinants of competitive success than a mere sense that the technology is easy to use or conformity to social expectations.

From the perspective of competition organizers, these findings point to two concrete practical implications. First, organizers should cultivate clear, outcome-oriented expectations by presenting concrete use cases or hosting workshops that demonstrate how AI can help resolve robot control programming challenges or improve mission scores. Second, organizers should invest in infrastructure that is genuinely usable in practice — for instance, ensuring stable, high-speed internet connectivity at competition venues to allow reliable access to cloud-based AI services, alongside sufficient practice space and on-site technical staff to support the testing of AI-generated code. Prioritizing these two areas is likely to have a greater impact on performance outcomes than training sessions that focus solely on basic AI usage skills (e.g., introductory prompt writing), given that the youth competitors already display a high degree of familiarity with the ease-of-use aspects of these technologies.

In addition, the ceiling effect observed during preliminary testing of the linear (OLS) model — whereby 43.8% of competitors achieved the maximum score of 100 — constitutes a statistically noteworthy finding in its own right. This phenomenon suggests that competition outcomes bounded by a fixed scoring range and subject to judges' discretionary evaluation are better analyzed as ordinal categorical data than as an unbounded continuous variable; the detection of this limitation itself underscores the importance of rigorously verifying statistical assumptions prior to model selection. It should also be noted that this study relied on a single multinomial logistic regression specification without benchmarking against alternative classification algorithms; the reported accuracy therefore reflects the performance of a statistically well-justified model rather than an attempt to maximize the highest accuracy achievable from this dataset.

VI. CONCLUSION

Across all four regions, Innovedex 2026 competitors displayed a relatively high level of AI tool acceptance prior to the competition, particularly with respect to perceived Social Influence (SI) and Effort Expectancy (EE), while confidence that the tools would genuinely elevate performance (PE) was the lowest-rated construct. Nevertheless, all four UTAUT constructs were significantly and positively associated with award level, led by Facilitating Conditions (FC) and Performance Expectancy (PE).

For outcome prediction, the multinomial logistic regression model — adopted in place of the linear model, which failed its underlying assumption tests — classified award levels effectively, achieving a cross-validation accuracy of 68.7%, markedly exceeding the 46.8% baseline.

Finally, this study is subject to scope-related limitations, as data were drawn from a single competition event over a one-year period and relied on a self-selected sample of voluntary participants. Accordingly, the magnitude of the relationships identified here — which exceeds what is typically reported in the broader technology acceptance literature — warrants replication with independent samples to confirm the reliability of these findings and to extend their generalizability to a wider population.

ACKNOWLEDGMENT

The researchers wish to thank Dr. Vitsanu Nittayathammakul for graciously serving as advisor and providing strategic guidance on data collection for this study. His generous assistance and feedback throughout the organization of the Innovedex 2026 competition were instrumental to the successful and complete fulfillment of this study's objectives.

REFERENCES

- [1] J. Liang, W. Huang, F. Xia, P. Xu, K. Hausman, B. Ichter, and A. Zeng, "Code as policies: Language model programs for embodied control," in *Proc. 2023 IEEE Int. Conf. Robot. Autom. (ICRA)*, London, UK, May 2023, pp. 9493–9500, doi: 10.1109/ICRA48891.2023.10160591.
- [2] Innovedex, "Innovedex 2026 robotics and artificial intelligence competition: Project document," Innovedex Robotics Competition Thailand, 2026. [Online]. Available: <https://innovedex.com/2026>
- [3] V. Venkatesh, M. G. Morris, G. B. Davis, and F. D. Davis, "User acceptance of information technology: Toward a unified view," *MIS Q.*, vol. 27, no. 3, pp. 425–478, Sep. 2003, doi: 10.2307/30036540.
- [4] V. Venkatesh, J. Y. L. Thong, and X. Xu, "Consumer acceptance and use of information technology: Extending the unified theory of acceptance and use of technology," *MIS Q.*, vol. 36, no. 1, pp. 157–178, Mar. 2012, doi: 10.2307/41410412.
- [5] Q. Ke, Y. Gong, and C. Ke, "Bridging AI literacy and UTAUT constructs: Structural equation modeling of AI adoption among Chinese university students," *Humanit. Soc. Sci. Commun.*, vol. 12, art. 1452, 2025, doi: 10.1057/s41599-025-05775-y.
- [6] K. S. Taber, "The use of Cronbach's alpha when developing and reporting research instruments in science education," *Res. Sci. Educ.*, vol. 48, no. 6, pp. 1273–1296, 2018, doi: 10.1007/s11165-016-9602-2.
- [7] J. B. Ramsey, "Tests for specification errors in classical linear least-squares regression analysis," *J. R. Stat. Soc. B*, vol. 31, no. 2, pp. 350–371, 1969, doi: 10.1111/j.2517-6161.1969.tb00796.x.
- [8] T. S. Breusch and A. R. Pagan, "A simple test for heteroscedasticity and random coefficient variation," *Econometrica*, vol. 47, no. 5, pp. 1287–1294, Sep. 1979, doi: 10.2307/1911963.
- [9] D. W. Hosmer, S. Lemeshow, and R. X. Sturdivant, *Applied Logistic Regression*, 3rd ed. Hoboken, NJ, USA: Wiley, 2013, doi: 10.1002/9781118548387.
- [10] D. McFadden, "Conditional logit analysis of qualitative choice behavior," in *Frontiers in Econometrics*, P. Zarembka, Ed. New York, NY, USA: Academic Press, 1974, pp. 105–142.
- [11] F. Pedregosa et al., "Scikit-learn: Machine learning in Python," *J. Mach. Learn. Res.*, vol. 12, pp. 2825–2830, 2011, doi: 10.5555/1953048.2078195.
- [12] S. Seabold and J. Perktold, "Statsmodels: Econometric and statistical modeling with Python," in *Proc. 9th Python Sci. Conf.*, 2010, pp. 92–96, doi: 10.25080/MAJORA-92BF1922-011.