

ALGORITHMIC INVESTMENT LITERACY: TEACHING STUDENTS TO EVALUATE, CALIBRATE, AND ETHICALLY USE AI INVESTMENT ADVICE

Fahmi Abdul Rahim^{1*}

Nur Hafidzah Idris²

Amirudin Mohd Nor³

Rohaiza Kamis⁴

Mohd Hafiz Bakar⁵

Faezah Othman⁶

¹ Faculty of Business and Management, Universiti Teknologi MARA (UiTM), Cawangan Melaka, Kampus Bandaraya Melaka, 75300 Melaka, Malaysia; Email: fahmi029@uitm.edu.my

² Faculty of Business and Management, Universiti Teknologi MARA (UiTM), Cawangan Melaka, Kampus Bandaraya Melaka, 75300 Melaka, Malaysia; Email: fidza769@uitm.edu.my

³ Faculty of Business and Management, Universiti Teknologi MARA (UiTM), Cawangan Melaka, Kampus Bandaraya Melaka, 75300 Melaka, Malaysia; Email: amirudinmn@uitm.edu.my

⁴ Faculty of Business and Management, Universiti Teknologi MARA (UiTM), Cawangan Melaka, Kampus Bandaraya Melaka, 75300 Melaka, Malaysia; Email: rohai451@uitm.edu.my

⁵ Faculty of Business and Management, Universiti Teknologi MARA (UiTM), Cawangan Melaka, Kampus Bandaraya Melaka, 75300 Melaka, Malaysia; Email: hafizbakar@uitm.edu.my

⁶ Faculty of Business and Management, Universiti Teknologi MARA (UiTM), Cawangan Melaka, Kampus Bandaraya Melaka, 75300 Melaka, Malaysia; Email: faezahothman@uitm.edu.my

* Corresponding Author

Article history

Received date : 14-7-2026

Revised date : 15-7-2026

Accepted date : 29-8-2026

Published date : 30-8-2026

To cite this document:

Abdul Rahim, F., Idris, N. H., Mohd Nor, A., Kamis, R., Bakar, M. H., & Othman, F. (2026). Algorithmic investment literacy: Teaching students to evaluate, calibrate, and ethically use AI investment advice. *International Journal of Accounting, Finance and Business (IJAFB)*, 11 (67), 657 - 677.

Abstract: *The diffusion of robo-advisors and generative artificial intelligence has changed the central problem of investment education. When advanced language models score near the ceiling on standard financial-literacy tests, the decisive competency is no longer whether learners possess financial knowledge, but whether they can judge the machine-generated advice they receive. This paper defines Algorithmic Investment Literacy (AIL) by a single operation that no neighbouring construct centres: the learner lets a personal-suitability judgement, whether a recommendation fits this investor's own goals, risk tolerance, horizon and constraints, govern reliance on an adviser whose reasoning is opaque, so that suitability can veto a reliable-looking recommendation and license a plainly generated one. The human-automation reliance literature calibrates reliance to a system's verifiable task accuracy rather than to a recommendation's private, unverifiable fit with the person; financial literacy concerns what a person knows and not how they appraise an external adviser; generic AI literacy asks whether an output is sound in itself and not whether it is right for this person. Each supplies part of the operation; AIL joins them, and we argue it predicts reliance-decision quality beyond their additive combination, showing construct by construct the behaviour AIL predicts that each neighbour cannot. Four finance-specific dimensions follow from the operation rather than*

standing beside it: provenance awareness, suitability appraisal, suitability-governed reliance, and fiduciary and responsible use. The paper advances nine propositions linking pedagogical and learner antecedents to AIL and its outcomes, and translates the construct into a learning progression, a worked classroom task with an assessment rubric, and a research agenda, giving educators a teachable, assessable target and researchers a foundation for instrument development and validation.

Keywords: *algorithmic investment literacy; artificial-intelligence literacy; investment education; robo-advisors; self-regulated learning; metacognition; competency framework*

Introduction

Consider a final-year undergraduate who opens a mobile investing application, answers six risk-profiling questions, and receives a fully specified portfolio recommendation generated by an algorithm. She can also ask a general-purpose chatbot whether to accept it, and the chatbot returns a fluent, confident answer within seconds. Nothing in her prior coursework tells her how that advice was produced, whether it fits her actual circumstances, or when she should trust it against her own judgement. This situation is now ordinary. Automated advice has moved from a niche service to a mainstream channel for retail investing, and the person who must act on it is very often a student who is simultaneously entering both formal investment education and independent financial life.

Investment education has not yet named the competency this student lacks. The field has an extensive vocabulary for what people know about finance and, more recently, for how they use digital financial services, but it does not have a settled construct for the capacity to engage critically with advice produced by an algorithm. The purpose of this paper is to argue that this capacity is distinct, that it is teachable, and that it should be an explicit objective of investment curricula. We name it Algorithmic Investment Literacy (AIL) and define it, provisionally, as the competency that enables a person to understand how artificial-intelligence systems generate investment advice, to critically evaluate that advice, to calibrate reliance on it, and to use it ethically and responsibly in service of sound and autonomous investment decisions.

The distinctive claim is narrow and can be stated at once. What makes AIL more than a synthesis of existing literacies is a single joint operation: the learner lets a personal-suitability judgement, whether a specific recommendation fits this investor's own goals, risk tolerance, horizon, liquidity and constraints, govern reliance on an adviser whose reasoning is opaque. Suitability can veto a recommendation from a demonstrably reliable system and can license one from a plainly generated model. No neighbouring construct centres this operation. The human-automation reliance literature calibrates reliance to a system's verifiable task accuracy rather than to a recommendation's fit with the individual's private and not externally verifiable goals; financial literacy concerns what a person knows and not how they appraise an external adviser; generic artificial-intelligence literacy evaluates whether an output is sound in itself and not whether it is right for this person. Each supplies part of the operation; AIL is the competency that joins them and, we argue, predicts reliance decisions beyond their additive combination. Section 6 shows, construct by construct, the behaviour AIL predicts that each neighbour cannot, and Proposition 9 states the incremental claim directly. Everything else in the construct, including its four dimensions, follows from this operation rather than standing beside it.

The urgency of naming this competency follows from a simple observation. Financial-literacy research has long treated financial knowledge as the scarce resource that education must supply (Goyal & Kumar, 2021). Yet in an early study an advanced generative model scored close to full marks on a standard financial-literacy test (Niszczota & Abbas, 2023). That single result awaits wider replication, and we do not rest the argument on it alone; but the rapid and continuing improvement of these models makes the underlying shift difficult to dismiss. When the machine is itself financially literate, the pedagogically important question moves from whether students know finance to whether students can judge the financial advice a machine gives them. Evidence that automated-advice adoption is driven in part by overconfidence rather than by demonstrated competence sharpens the concern, because it suggests that the learners least able to evaluate algorithmic advice may be among the most willing to follow it (Piehlmaier, 2022; Shanmuganathan, 2020).

The paper is deliberately theoretical rather than empirical, in the tradition of construct-development work that precedes measurement. Construct clarification is a necessary first step before an instrument can be built or a hypothesis tested, and we return in Section 9 to a worked classroom task, an assessment rubric, a concrete item pool, and a validation plan, so that the construct is empirically actionable rather than merely asserted. We make four moves. First, we show the gap by examining the nearest constructs and identifying the assumption each shares that the phenomenon breaks. Second, we define AIL by the joint operation above and derive its four finance-specific dimensions from it. Third, we differentiate AIL from its neighbours not by asserting a difference but by stating, for each, a behaviour AIL predicts that the neighbour cannot, with the observable that separates them. Fourth, we advance a set of propositions and boundary conditions that render the construct testable, and we translate it into a university-level learning progression, an assessment blueprint with a worked task, and a research agenda. Throughout, we hold the focus on university students, the population for whom the stakes of algorithmic advice and the opportunity of formal instruction coincide most closely.

The gap in the adjacent literatures

Six literatures sit close to the phenomenon, and each stops short of it. We take them in turn and, for each, name the assumption that the phenomenon breaks.

Financial literacy

Financial literacy, the most established of the neighbours, treats the individual's stock of financial knowledge as the object of interest and the lever of intervention. Bibliometric syntheses show a field organised around measuring knowledge and linking it to behaviour, largely through the Theory of Planned Behaviour and structural-equation designs (Goyal & Kumar, 2021). The shared assumption is that the binding constraint on good decisions is what the individual knows. When a fluent external adviser supplies the knowledge on demand, that assumption weakens, and the constraint shifts to the individual's ability to appraise the adviser.

Financial capability and financial self-efficacy

Financial capability extends literacy by adding the opportunity to act, locating competence in the relationship between the individual and financial institutions rather than in the individual alone (Sherraden, 2013; Xiao et al., 2022). Financial self-efficacy, in turn, captures a learner's belief in their own ability to manage money and is a motivational-belief construct, rooted in self-efficacy theory (Bandura, 1997), rather than a knowledge or skill construct (Al Rahahleh, 2023). Both are indispensable, and both share an assumption the phenomenon breaks: that the challenge is either access to act or confidence to act. A student can have access to an investing

platform and feel confident using it, yet still be unable to tell whether the algorithm's recommendation is sound. Confidence without the capacity to appraise machine advice is precisely the condition that overconfidence research identifies as risky (Piehlmaier, 2022).

Digital financial literacy

Digital financial literacy (DFL) adds the knowledge, skills, and attitudes needed to access and safely use digital financial services (Koskelainen et al., 2023; Lyons & Kass-Hanna, 2021). International measurement now codifies DFL around awareness of digital services, practical know-how, and self-protection against digital fraud (OECD/INFE, 2024). DFL is essential, but its object is the safe operation of a platform, not the adjudication of the advice the platform's algorithm produces. A learner can be adept at using an application and protecting a password while remaining unable to judge whether the portfolio the application recommends is suitable.

Generic artificial-intelligence literacy

Generic artificial-intelligence literacy supplies the evaluative and ethical vocabulary that the finance constructs lack. Reviews converge on a small set of competencies: knowing and understanding AI, using and applying it, evaluating and creating with it, and engaging with it ethically (Long & Magerko, 2020; Ng et al., 2021), with later frameworks adding confidence and self-reflection (Chiu et al., 2024). This literature is, however, deliberately domain-neutral. It does not address the specific stakes, opacity, and suitability logic of investment decisions, where an error is not a wrong answer on a worksheet but a mis-allocated portfolio.

Robo-advisor adoption

The behavioural literature on robo-advisor adoption studies whether individuals intend to use automated advice, typically through the Technology Acceptance Model or the Unified Theory of Acceptance and Use of Technology, with financial literacy entering as a predictor or moderator (Aristei & Gallo, 2025; Pattnaik & Hota, 2025). Adoption research asks whether people use the tool. It does not ask whether they use it well, and using it well is a learning outcome rather than an intention. The digital-fluency and literacy nexus in this setting is itself flagged as under-theorised (Pattnaik & Hota, 2025).

Human-automation trust and reliance

A sixth literature, from human factors and behavioural decision research, is the closest of all and must be confronted directly. Work on trust in automation shows that good outcomes depend on appropriate reliance, that is, trust calibrated to a system's actual reliability, with both over-trust and under-trust harmful (Lee & See, 2004; Parasuraman & Riley, 1997). Related work documents algorithm aversion, the tendency to abandon an algorithm after seeing it err even when it outperforms humans (Dietvorst et al., 2015; Burton et al., 2020), and its opposite, automation complacency, with algorithmic vigilance as the desired mid-point (Zerilli et al., 2022). This literature is the true origin of the calibrated-reliance idea, and intellectual honesty requires saying so. What it does not do is treat calibration as a competency to be taught. It describes how operators and consumers rely, and it designs systems and interfaces to nudge reliance, but it does not ask how a learner acquires the capacity to calibrate, how that capacity develops across a curriculum, or how it combines with the financial-suitability judgement specific to investing. The finance evidence confirms the phenomenon is real in this domain: disclosing that advice is AI-generated shifts investor reactions (Downen et al., 2024), and aversion intensifies as the stakes of a decision rise (Filiz et al., 2023), which is precisely the condition under which the competency proposed here is most needed.

The empty space these literatures leave is a domain-specific, teachable competency for engaging critically with algorithmic investment advice. AIL is proposed to occupy exactly that space.

Partial resources from the adjacent literatures

Naming a new construct does not mean discarding its neighbours. Each supplies something to AIL and withholds something else, and stating both credits the lineage while earning the construct its independence.

Financial literacy supplies the substantive knowledge base against which advice is judged. A learner cannot recognise an unsuitable recommendation without some grasp of diversification, risk, and cost. What financial literacy does not supply is any stance toward an external, machine adviser; it locates competence inside the individual's own knowledge rather than in the relation between the individual and a system that advises.

Financial capability supplies the insight that competence is relational, residing partly in the fit between a person and the institutions and tools available to them (Sherraden, 2013). AIL adopts this relational stance and applies it to a new counterpart, the algorithmic adviser. Financial self-efficacy supplies the reminder that belief and confidence shape action (Al Rahahleh, 2023); AIL treats calibrated confidence as part of the competency, while insisting that confidence must be earned by the capacity to appraise, not assumed.

Digital financial literacy supplies operational fluency and risk awareness for transacting safely in digital environments (OECD/INFE, 2024). What it does not supply is the evaluative judgement to weigh the quality of advice, as distinct from the safety of a transaction. Generic AI literacy supplies the evaluate-and-ethics architecture, in particular the ideas that AI outputs must be appraised and that their use carries ethical weight (Ng et al., 2021; Allen & Kendeou, 2024). What it does not supply is contextualisation to a domain in which advice is personalised, consequential, and often opaque.

The human-automation reliance literature supplies the core mechanism, a validated account of how reliance becomes miscalibrated and what makes calibration accurate (Küper & Krämer, 2025; Lee & See, 2004). What it does not supply is the educational move. It studies operators in situ rather than learners in formation, and it seeks to engineer appropriate reliance through system and interface design rather than to cultivate it as a durable, transferable capability in the person. AIL takes the mechanism and asks the educational question the source literature does not: how is suitability-governed reliance learned, developed, and assessed.

Self-regulated learning and metacognition supply the internal architecture that turns reliance from a reflex into a regulated act, and this is the literature that carries the educational move the automation account lacks. The metacognitive tradition separates a monitoring level from a control level and treats skilled performance as the accurate feeding of the first into the second, so that what a learner knows about their own cognition governs what they do (Flavell, 1979; Nelson & Narens, 1990). Self-regulated learning models place that loop inside a cycle of forethought, performance, and reflection that is not fixed but develops with instruction and feedback (Zimmerman, 2002). What this literature does not supply is the object of regulation that concerns us: an opaque external adviser in a high-stakes financial domain. AIL takes the monitoring-and-control account and points it at that adviser, asking how a learner develops accurate monitoring of machine advice and sound control over reliance on it.

Experiential and simulation-based finance pedagogy supplies the vehicle through which contextualised judgement can be developed, since trading laboratories build decision competence through concrete experience and reflection (Bakoush, 2022; Marriott et al., 2015; Yuan, 2026). On the instructor side, the technological-pedagogical-content-knowledge (TPACK) tradition supplies the account of what teachers must know to integrate a technology into a specific content domain (Koehler et al., 2013; Schmidt et al., 2009), now being extended to the integration of artificial intelligence itself (Ning et al., 2024). What these supply is the means of teaching; what they do not supply is a defined learner competency to target and assess. AIL draws on all of these and belongs to none.

Defining the construct

Algorithmic Investment Literacy is the domain-specific competency that enables an individual to understand how artificial-intelligence systems, including robo-advisors and generative AI, produce investment advice; to critically evaluate that advice; to calibrate reliance on it appropriately; and to use it ethically and responsibly, in service of sound and autonomous investment decision-making.

Two elements of this definition are constitutive, meaning they must hold for the construct to apply, and two are capabilities the construct develops. The constitutive conditions are, first, that advice originates from an artificial-intelligence system whose reasoning is at least partly opaque to the user, and second, that the decision at stake is consequential, so that reliance carries real cost. Where advice is transparent and trivial, AIL is not engaged. The capabilities are evaluation and calibration: the learner appraises the advice and then regulates how far to rely on it. The phrase autonomous decision-making marks the intended end state, in which the learner remains the decision-maker rather than a passive recipient of automation.

The definition also keeps three things distinct that finance-education research often conflates: knowing, feeling capable, and doing. AIL is not knowledge alone, since a knowledgeable learner may still defer uncritically to a confident machine; it is not confidence alone, since confidence uncalibrated to the quality of advice is the overconfidence that drives poor reliance (Piehlmaier, 2022); and it is not behaviour alone, since a learner may follow good advice by luck without the competency to tell good advice from bad. AIL is the capability that connects appraisal to action through suitability-governed reliance.

Dimensions of the construct

AIL comprises four interrelated dimensions, but they are not four independent competencies placed side by side. They are what the single operation defined above, letting suitability govern reliance on an opaque adviser, requires at each step: the learner must understand why machine investment advice can be wrong, appraise whether a given recommendation suits her, control reliance accordingly, and act responsibly on the result. We name the dimensions by their finance-specific function rather than by the generic knowing, using, evaluating and ethics labels of domain-neutral AI-literacy instruments, because the domain re-specifies each one; the parallel is a lineage, not an identity.

Provenance awareness

The learner can explain, in non-technical terms, how a robo-advisor or a generative model produces investment advice: the data it draws on, the assumptions embedded in its models, and its inherent limits, such as training cut-offs and the non-stationarity of financial markets. This is the cognitive foundation on which evaluation rests. Without a working mental model of how

the advice was generated, a learner cannot form reasoned expectations about where it is likely to be reliable or unreliable.

Suitability appraisal

The learner can appraise whether a specific recommendation is sound and, above all, whether it suits this investor. This subsumes detecting bias in the underlying model, recognising fabricated or unsupported claims from a generative system, and noticing conflicts of interest, but it treats these as inputs to the governing question the domain forces: does this recommendation fit this person's goals, risk tolerance, horizon and constraints? Generic quality appraisal, the evaluative competency at the centre of domain-neutral AI-literacy frameworks (Ng et al., 2021), can judge whether an output is good in itself but cannot answer that question, which is why suitability appraisal is domain-specific rather than a relabelling of generic evaluation.

Suitability-governed reliance

The learner can self-regulate the degree of reliance on algorithmic advice, deciding when to accept it, when to adjust it, and when to override it, and can justify that decision. This is a metacognitive operation in the strict sense. In the monitoring-and-control architecture of metacognition, a monitoring process reads the state of the object level, here the learner's confidence in a recommendation and in their own competing judgement, and feeds a control process that acts on it, here the decision to defer, adjust, or override (Flavell, 1979; Nelson & Narens, 1990). Suitability-governed reliance is that control decision made well: it avoids both automation bias, the tendency to over-trust the machine, and unwarranted algorithm aversion, the tendency to reject useful advice. Placed in a self-regulated learning cycle, it spans forethought, in which the learner forms an expectation of where the system is reliable, performance monitoring, in which the recommendation is checked against that expectation and against the learner's goals, and self-reflection, in which the expectation is revised after the outcome (Zimmerman, 2002). Framing appropriate reliance in these terms is what makes it teachable: monitoring accuracy and control policies can be modelled, practised, and assessed, rather than only designed into an interface. This dimension is the pivot on which good outcomes turn.

Fiduciary and responsible use

The learner can act on the ethical dimensions of algorithmic advice: data privacy, accountability for outcomes, disclosure, and awareness of automation and mis-selling risks, and can align choices with responsible-investing values. This dimension reflects the consistent finding that ethics is not an optional add-on to AI competency but one of its core components (Allen & Kendeou, 2024; Chiu et al., 2024).

A stylised illustration

A composite student, drawn from no particular person or institution, shows the dimensions working together. She receives a recommendation to concentrate her savings in a single thematic fund. Provenance awareness lets her recognise that the recommendation reflects recent returns the model was trained on, not a forecast. Suitability appraisal lets her see that the concentration conflicts with her stated goal of stability. Suitability-governed reliance lets her retain the parts of the advice that fit, such as the low-cost fund structure, while overriding the concentration. Fiduciary and responsible use lets her check what personal data the application retained and whether the recommendation carried an undisclosed incentive. None of these acts is reducible to knowing more finance; each is an act of appraising and governing an adviser.

Distinguishing the dimensions

A reasonable objection is that the dimensions, and in particular provenance awareness and suitability appraisal, may not be empirically separable, a difficulty long noted for the sub-domains of TPACK, whose factor structure has proven hard to recover (Archambault & Barnett, 2010). We treat the two as conceptually ordered but distinct. Provenance awareness is a descriptive understanding of how advice is generated, a knowledge state; suitability appraisal is an appraisal act applied to a particular output. The distinction is testable through dissociation. A learner who can accurately explain how a model produces advice yet fails to flag a seeded error in a specific recommendation displays awareness without evaluation; the reverse, flagging that something is wrong without being able to explain its source, is also possible. Demonstrating such dissociations, and recovering four factors rather than one, is part of the discriminant-validity agenda set out in Sections 9.4 and 11, and we treat the separability of the dimensions as a hypothesis to be tested rather than an assumption.

The boundary most open to challenge is not provenance and appraisal but suitability appraisal and suitability-governed reliance, since appraising fit and deciding how far to act on it are closely linked. We separate them by function and by what governs each. Suitability appraisal answers a question about the recommendation: does this fit this investor? Suitability-governed reliance answers a question about action under two-sided uncertainty: given that appraisal and the learner's confidence in her own competing judgement, how far should she defer, adjust, or override? The two dissociate in both directions. A learner may correctly judge a recommendation unsuitable yet still, correctly, defer to it because her own alternative is weaker; another may judge advice sound yet override it out of misplaced confidence in herself. Because reliance is governed by the learner's estimate of her own competence as much as by the appraisal, it is a distinct control act rather than a restatement of the appraisal, and the validation design in Section 9.4 treats whether the two load on separate factors as an explicit test rather than an assumption.

Table 1 summarises the dimensions and indicates the propositions most closely related to each. Proposition 1 is structural and applies to all four dimensions; Proposition 3 (the instructor antecedent) and Proposition 9 (incremental validity over the combined neighbours) are likewise cross-cutting rather than tied to a single row.

Table 1: Dimensions of Algorithmic Investment Literacy

Dimension	What it involves	Related proposition
Provenance awareness	Understanding how AI systems generate investment advice, and their assumptions and limits.	P2, P4
Suitability appraisal	Appraising the quality, suitability, and integrity of a specific recommendation.	P2, P4, P5
Suitability-governed reliance	Self-regulating acceptance, adjustment, or override of advice, with justification.	P2, P5, P6, P7, P8
Fiduciary and responsible use	Acting on privacy, accountability, and responsible-investing considerations.	P8

Differentiation from adjacent constructs

Construct clarity requires showing that AIL is neither a relabelling nor a subset of an existing literacy. For each neighbour we state what AIL borrows and the precise mechanism that separates it.

Table 2 makes the separation operational. For each neighbour it states a critical case, what that construct predicts the learner does, what AIL predicts, and the observable that adjudicates between them. The two rows that carry the most weight are generic AI literacy and appropriate reliance: in both, the learner's appraisal of the output's general quality or the system's reliability is held fixed, yet AIL predicts the opposite reliance decision because a personal-suitability judgement governs it. That is the behaviour no single neighbour can generate.

Table 2: What AIL predicts that each adjacent construct cannot

Adjacent construct	What the construct predicts	What AIL predicts	Observable that separates them
Financial literacy	Evaluates the finance content, which looks sound, and accepts.	Flags an adviser-level fault, such as a fabricated statistic, and discounts.	Detection of a seeded fabrication that is invisible from the finance content alone.
Financial capability	Has access and opportunity, so transacts.	Overrides an executable, reliable-looking recommendation on suitability grounds.	Override despite full access, because the recommendation does not fit her goals.
Financial self-efficacy	Confidence drives action, regardless of appraisal accuracy.	Correct overrides track appraisal accuracy, not confidence.	Among confident learners, override correctness tracks appraisal rather than self-rated confidence.
Digital financial literacy	Operates the platform competently, so transacts.	Adjudicates the advice and overrides when it does not suit.	High operating fluency co-exists with low advice adjudication in the same learner.
Generic AI literacy	Judges the output good in itself, so relies.	Vetoes it on personal suitability despite high generic quality.	Same quality appraisal, opposite reliance decision.
Appropriate reliance (automation trust)	System is reliable in general, so relies.	Overrides the reliable system because the output violates a personal constraint.	Override of a reliable-system output for person-fit, which reliability calibration cannot license.
Metacognition and self-regulated learning	Monitors her own knowledge well, but has no purchase on the adviser's suitability.	Appraises the external adviser using finance-suitability content.	High self-regulation with low finance knowledge still fails the suitability appraisal.
Robo-advisor adoption	Predicts uptake and continued use of the tool, so more adoption reads as success.	Predicts governance of the tool's advice, including declining or overriding it when it does not suit.	Adoption intensity does not predict override accuracy; a heavy user can be a poor adjudicator.

From financial literacy, AIL borrows the substantive knowledge base, but it is separated by its object. Financial literacy locates competence in what the individual knows; AIL locates competence in how the individual appraises and governs an external adviser. A learner can score highly on a financial-literacy test and still lack AIL if she cannot tell whether a chatbot's confident recommendation is sound.

From financial capability, AIL borrows the relational view of competence, but it is separated by the counterpart in the relation. Capability concerns the fit between a person and financial institutions and products (Sherraden, 2013); AIL concerns the fit between a person and an algorithmic adviser whose reasoning is opaque. From financial self-efficacy, AIL borrows the recognition that confidence matters, but it is separated by the requirement that confidence be calibrated to the quality of advice rather than felt in general (Al Rahahleh, 2023). High self-efficacy with low evaluative capacity is not AIL; it is the overconfident reliance the construct is designed to prevent.

From digital financial literacy, AIL borrows operational fluency and risk awareness, but it is separated by the difference between operating a platform and adjudicating its advice. Safe use of a service (OECD/INFE, 2024) does not entail the capacity to judge whether the service's algorithmic recommendation is suitable. From generic AI literacy, AIL borrows the evaluate-and-ethics architecture, but it is separated by a domain contextualisation that is constitutive rather than cosmetic. In investing, evaluation is not the generic appraisal of an AI output; it is a suitability judgement, whether a specific recommendation fits this learner's goals, risk tolerance, horizon and constraints, which has no analogue in a domain-neutral framework. Ethics is not generic AI ethics but the finance-specific matter of fiduciary duty, undisclosed incentives, mis-selling and responsible-investing alignment. Even awareness is domain-shaped, requiring an understanding of why financial prediction is hard, through market non-stationarity and model training cut-offs, that a generic account of how AI works does not supply. The domain changes what each dimension is, not merely where it is applied (Ng et al., 2021; Chiu et al., 2024). From TPACK, which is an instructor-side account of the knowledge needed to teach with technology (Koehler et al., 2013; Ning et al., 2024), AIL is separated by side and by object: TPACK concerns what the teacher knows in order to teach, whereas AIL is a competency the learner acquires.

From self-regulated learning and metacognition, AIL borrows the monitoring-and-control form of its central dimension, but it is separated by object and by content. Metacognition and self-regulated learning are domain-general and turned inward: their object is the learner's own knowledge and performance, and the monitoring at issue is monitoring of oneself (Flavell, 1979; Zimmerman, 2002). AIL turns the same form outward, onto an external adviser whose reasoning the learner cannot introspect, so the monitored quantity is the reliability and suitability of another agent's output rather than the learner's own recall or effort. It also supplies content that general self-regulation lacks, the financial-suitability knowledge without which reliability monitoring cannot become a sound accept-or-override decision. A capable self-regulated learner who does not grasp how a robo-advisor works or what makes a portfolio suitable will still fail to appraise the advice, which is why prior self-regulatory skill is an enabling antecedent of AIL rather than the thing itself. AIL is therefore metacognitive in form but neither reducible to metacognition nor derivable from it without the adviser as its object and finance as its content.

What is genuinely new

The sharpest objection is no longer that AIL is digital financial literacy renamed, but that its central mechanism, suitability-governed reliance, is simply appropriate reliance imported from the human-automation trust literature (Dietvorst et al., 2015; Lee & See, 2004). We accept the lineage and locate the difference precisely, and we do not overstate it. That literature does recognise a purpose component of trust, the fit between an automated aid and the operator's goals (Lee & See, 2004); what it does not do is treat a private, individually indexed suitability

judgement as the standard to which reliance is calibrated, because its benchmark for appropriate reliance is the system's verifiable accuracy on a defined task. Two further differences separate AIL. First, the reliance literature theorises reliance as a behavioural tendency of operators and consumers and seeks to engineer it through system design and interface transparency, whereas AIL constitutes reliance calibration as a teachable, assessable competency the learner acquires: the move is from aligning the system to educating the person. Second, it exercises that capability in a domain where the suitability standard is personal, unverifiable, and consequential. These are the differences Section 6.2 defends and Proposition 9 puts to an incremental-validity test.

The crux, though, is the fusion of two judgements that the source literatures keep apart. The automation literature asks one question, how reliable is the system, and calibrates reliance to the answer. Investing forces a second, independent question that the first cannot answer: does this particular recommendation fit this particular person, given goals, risk tolerance, horizon, liquidity needs, and constraints such as Shariah compliance or ethical exclusions. A robo-advisor can be reliable in general and still return advice that is unsuitable for a given investor, and a generative model can be articulate and confident while being wrong for reasons only a suitability check would catch. AIL requires the learner to hold both judgements at once and to let suitability govern reliance: to override a reliable-looking recommendation because it does not fit, and to accept a plainly generated one because, on inspection, it does. Neither the automation literature, which has no concept of personal suitability, nor financial literacy, which has no concept of appraising an external adviser, contains this joint operation. The move from designing systems for appropriate reliance to educating people for it, in a domain where reliability and suitability must be judged together, is the mechanism that makes the construct new.

The ground-truth defence

One objection deserves a direct answer, because the whole construct rests on it. A reviewer may argue that the correctness of AI advice in a personal-advice task already includes suitability, so that reliability and suitability collapse into one judgement and AIL adds nothing to appropriate reliance. The reply is a difference in ground truth. The human-automation reliance literature operationalises correctness as accuracy against a defined task with a known, external ground truth, for example whether a classifier labelled an image correctly. Suitability has no such external ground truth. It is indexed to the individual's own goals, risk tolerance, horizon, liquidity needs and constraints, information the system may never receive and cannot verify, and two investors receiving the identical, technically correct recommendation can be correctly served and badly served by it. Because suitability is private to the person and not a property of the output, it cannot be folded into a reliability estimate, and a competency defined by making that private judgement govern reliance is not reducible to calibrating reliance to system accuracy. This is why the two signature rows of Table 2 are possible at all: reliability is held fixed while the suitability-governed decision moves.

Propositions and framework

The framework separates one structural hypothesis from a set of theoretical propositions. Proposition 1 is structural: a measurement claim about the construct's dimensionality and its separability from neighbouring constructs. It is the hypothesis a validation study tests first, and the rest of the framework presupposes it. A second structural proposition, Proposition 9, makes the incremental-validity claim that AIL predicts reliance-decision quality beyond the additive combination of its component literacies; it too is a measurement claim, tested directly rather

than tied to any single dimension. The remaining propositions share a common logic: a pedagogical or learner condition develops AIL; AIL, exercised through suitability-governed reliance under the defining condition of opaque and consequential advice, produces decision and learning outcomes that ordinary financial knowledge would not; and the strength of that effect depends on conditions that make calibration more or less demanding. Three of these carry the theoretical weight, and a reader who wants the core claims should look to them: P2, that experiential design develops the competency; P5, that AIL improves reliance calibration and decision quality; and P8, that the competency transfers beyond the classroom. The remaining four specify the antecedents and boundaries that the core presupposes: P3 and P4 name the instructor and learner conditions that enable development, and P6 and P7 name the task and learner conditions that strengthen or weaken the central effect. We keep the outcomes distinct throughout, since knowing, judging, and doing are not interchangeable. Figure 1 presents the framework, and the propositions follow.

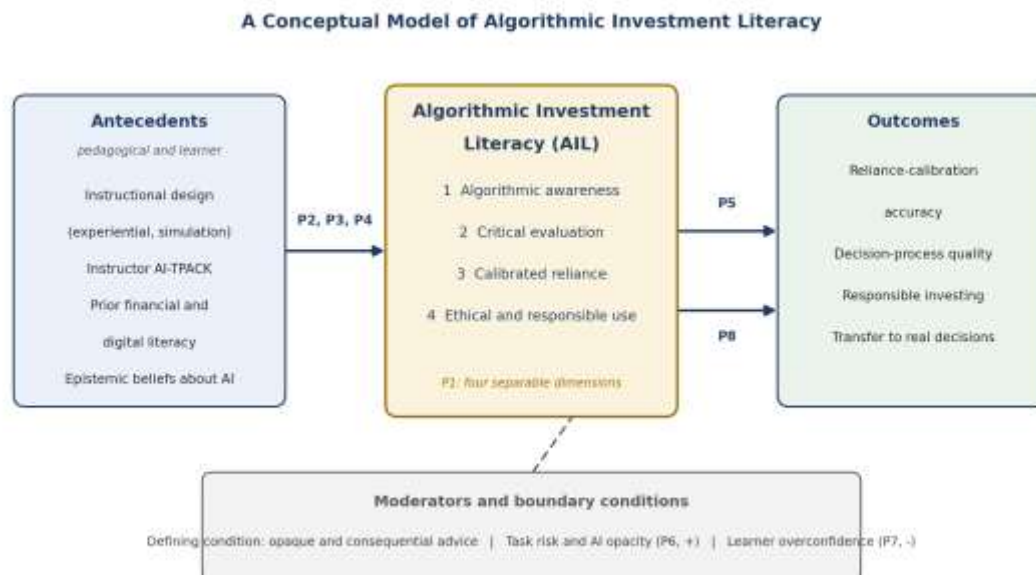


Figure 1. A conceptual model of Algorithmic Investment Literacy: pedagogical and learner antecedents, the four-dimensional construct, outcomes, and boundary conditions.

P1. Algorithmic Investment Literacy is a distinct, multidimensional competency that is empirically separable from financial literacy, financial capability, financial self-efficacy, digital financial literacy, generic AI literacy, appropriate reliance (trust in automation), and metacognition and self-regulated learning, and its four dimensions are recoverable as distinct but correlated factors. This is the structural claim that a validation study tests first; its discriminant-validity half is sharpened by Proposition 9, which asks not only whether AIL is separable from each neighbour but whether it predicts reliance decisions beyond the neighbours combined. (antecedent: construct structure; outcome: discriminant validity and dimensionality)

P2. When instruction requires learners to appraise opaque, consequential machine advice and to justify accept, adjust, or override decisions, experiential and simulation-based design develops AIL, and in particular its suitability-governed-reliance dimension, more effectively than transmissive instruction that conveys the same content without contextualised practice on advice. (antecedent: instructional design exercised on opaque, consequential advice; outcome: competency development, especially suitability-governed reliance)

P3. Instructor AI-related technological pedagogical content knowledge is positively associated with students' AIL development, because instructors who possess it design authentic appraisal tasks, model suitability-governed reliance by thinking aloud through their own accept, adjust, and override decisions, and give targeted feedback on evaluation, whereas instructors without it tend to default to demonstrating tools. (*antecedent: instructor AI-TPACK; mechanism: task design, modelling, and feedback; outcome: competency development*)

P4. Prior financial literacy and digital literacy are enabling antecedents of AIL but are individually insufficient to produce it: students who score highly on both yet receive no instruction in appraising machine advice will not outperform on an advice-appraisal task, and the effect of AIL instruction on appraisal performance is not fully mediated by prior literacies. (*antecedent: prior literacies; outcome: competency development; falsifiable prediction: a residual effect of instruction beyond prior literacies*)

P5. AIL improves the process quality of algorithm-assisted investment decisions and the accuracy of learners' reliance calibration, reducing both automation bias and algorithm aversion. (*antecedent: AIL; mediator: suitability-governed reliance; outcome: reliance-calibration accuracy and decision-process quality*)

P6. The relationship between AIL and decision quality is stronger when task risk, complexity, and AI opacity are higher. (*moderator: task risk and opacity; strengthens AIL to outcome*)

P7. Learner overconfidence weakens the translation of AIL into suitability-governed reliance. (*moderator: overconfidence; weakens AIL to reliance*)

P8. AIL transfers beyond the classroom to learners' real algorithm-assisted financial decisions, and this transfer is carried by the durable metacognitive disposition of suitability-governed reliance, together with the responsible-use orientation, because these generalise across advisers and tools whereas evaluation skills tied to one interface do not. (*antecedent: AIL, principally its calibrated-reliance and responsible-use dimensions; outcome: transfer of suitability-governed reliance to real-world decisions*)

P9. AIL predicts the quality of suitability-governed reliance decisions beyond the additive combination of generic AI literacy, financial-suitability knowledge, and appropriate-reliance calibration: a learner high on all three components who has never been taught to let suitability govern reliance will still follow reliable-looking but unsuitable advice, so the joint operation explains variance that the three components entered together do not. This is the incremental-validity claim that separates the construct from the conjunction of its parts, and it is the sharpest test of whether AIL is more than a synthesis. (*antecedent: the joint operation, over and above its component literacies; outcome: reliance-decision quality; falsifiable prediction: a significant AIL or suitability-by-appraisal interaction term after the three components are entered additively as controls*)

Scope and boundary conditions

AIL applies where two conditions hold together: advice is generated by an artificial-intelligence system whose reasoning is at least partly opaque, and the decision is consequential enough that reliance carries cost. Two contrasting cases sharpen the boundary. A student who uses a generative model to rebalance a real retirement portfolio is squarely within scope, because the advice is opaque and the stakes are real. A student who uses a calculator with transparent, deterministic formulae to compute compound interest is outside scope, because nothing is

opaque and there is no adviser to appraise. The construct is therefore not a general claim about technology in finance education; it is specific to the appraisal of opaque, consequential machine advice.

Two further boundaries matter for the university setting. First, AIL presupposes a threshold of financial and digital literacy; below that threshold, developing AIL is premature, which is why the learning progression in the next section begins with foundational capability. Second, the construct is defined at the level of the individual learner and does not, in its present form, extend to institutional or regulatory responses to algorithmic advice, which are distinct questions with their own literatures.

The framework is meant to travel, but not unchanged. The balance among the dimensions, and the emphasis of the progression, will shift with context. Where regulation mandates suitability assessment and advice disclosure, the ethical and responsible-use dimension can lean on an existing rulebook, and instruction can spend more time on evaluation and calibration; where retail participation is newer and consumer protections thinner, provenance awareness and suitability appraisal carry more of the load, because the learner has fewer external guardrails. Markets also differ in the tools that dominate, from full robo-advisers to general-purpose chatbots, which changes what opacity looks like in practice. The construct is defined at a level that holds across these settings; its instructional weighting is a local decision, and we treat that weighting as an empirical question rather than a fixed prescription.

From construct to curriculum: progression, task, and assessment

Because AIL is defined as a competency, it must be developable in sequence and assessable in practice. This section translates the construct into a learning progression, a worked classroom task, an aligned rubric, and an item pool with a validation plan, so that AIL is usable in a course next term, not only in a research programme.

A learning progression

Following the higher-education pathway logic of AI-literacy competency frameworks (Chee et al., 2025), we specify three levels that map onto a typical undergraduate-to-postgraduate sequence, shown in Table 3.

Table 3: A university-level learning progression for AIL

Level	Typical stage	Target capability
Foundational	Early undergraduate; non-finance electives	Recognise AI in financial tools, describe at a high level how advice is generated, and identify obvious red flags.
Applied	Upper-level finance and investment courses	Critically evaluate advice against goals and risk, run simulations, and justify accept or override decisions with evidence.
Integrative	Capstone and postgraduate	Manage a full advice-assisted decision process under uncertainty, address ethics and accountability, and mentor peers.

An assessment blueprint aligned to the four dimensions

Assessment should align with the four dimensions and use complementary evidence, since no single instrument can capture a competency that spans knowing, judging, and doing. A self-report scale, adapted from validated AI-literacy questionnaires (Ng et al., 2024), can capture perceived competency across the four subscales. Scenario and performance tasks, in which students appraise realistic robo-advisor or generative outputs that have been seeded with errors,

can measure suitability appraisal directly. Simulation-embedded behavioural traces, capturing accept and override decisions and their justifications in a trading-laboratory exercise (Marriott et al., 2015; Yuan, 2026), can measure suitability-governed reliance as behaviour rather than as self-report. Reflective artefacts, marked with a rubric, can evidence the ethical dimension. Together these approaches make AIL assessable without reducing it to a single test score.

A worked classroom task and rubric

To make the blueprint concrete, we set out one assessable task that an instructor can adopt directly. In the seeded-error advice-appraisal task, each student receives a realistic robo-advisor output for a stated investor profile, together with a generative-AI rationale for that output, into which three flaws have been planted: an over-concentration that contradicts the profile's stated goal of stability, a fabricated or unsupported claim (for example, a fund or statistic the model has invented), and a stale-data assumption (advice that ignores a training cut-off or a regime change). Working individually or in pairs, students must (a) explain in plain language how the recommendation was generated, (b) identify and justify each flaw, (c) decide, for each element, whether to accept, adjust, or override it, and (d) state what evidence would change their decision. The task can be run on paper, inside a trading-simulation laboratory, or against a live tool whose output is captured, and it maps directly onto the four dimensions: (a) exercises provenance awareness, (b) exercises suitability appraisal, (c) exercises suitability-governed reliance, and (d) surfaces the metacognitive monitoring that calibration requires. A reflective addendum, asking who is accountable if the advice is followed and what data the tool retained, exercises fiduciary and responsible use.

Table 4 gives a rubric that scores each dimension across three performance levels, so that the task yields dimension-level evidence rather than a single mark.

Table 4: Rubric for the seeded-error advice-appraisal task

Dimension	Novice	Developing	Proficient
Provenance awareness	Cannot describe how the advice was produced, or treats it as a forecast.	Gives a partial account of data and assumptions.	Explains data, assumptions, and limits, including training cut-off and non-stationarity.
Suitability appraisal	Misses the planted flaws, or accepts the advice wholesale.	Detects some flaws but cannot justify them or link them to the profile.	Detects all three flaws and justifies each against the investor profile and evidence.
Suitability-governed reliance	Accepts or rejects the advice indiscriminately, without justification.	Adjusts some elements but reasoning is inconsistent or unjustified.	Accepts, adjusts, or overrides each element with justification, avoiding both bias and blanket aversion.
Fiduciary and responsible use	Ignores accountability, privacy, and disclosure.	Notes one ethical issue in passing.	Addresses accountability, data retention, disclosure, and responsible-investing alignment.

Illustrative item pool, validation plan, and the limits of self-report

To make the construct empirically actionable, and to meet the reasonable expectation that a competency should be measurable, we sketch an initial item pool and a validation plan. No data are reported here; this is a proposed instrument for subsequent testing, offered so that the construct can be examined rather than only asserted. Two illustrative self-report items per

dimension indicate the intended content. For provenance awareness: *I can explain in simple terms how a robo-advisor arrives at its recommendation; and I understand which market conditions a generative model is likely to handle poorly.* For suitability appraisal: *when an application recommends a portfolio, I check whether it fits my goals before accepting it; and I can tell when AI-generated financial advice makes a claim it cannot support.* For suitability-governed reliance: *I adjust how much I rely on automated advice according to how much I trust it in the situation; and I can decide when to override an algorithm's investment recommendation.* For fiduciary and responsible use: *I consider what personal data an investing application collects before I use it; and I think about who is accountable when automated advice leads to a loss.*

Self-report alone is, however, insufficient for this construct, and for a reason internal to the theory. The learners most likely to lack AIL are, on the evidence, the overconfident (Piehlmaier, 2022), and overconfident learners are precisely those who over-rate their own capacity to appraise advice, the well-documented pattern in which those least skilled at a task are least able to recognise their own deficit (Kruger & Dunning, 1999). Because AIL is itself a metacognitive competency, poor monitoring is not incidental to the deficit but part of it, so a self-report scale would be most inflated exactly where the true competency is lowest. For this reason the validation design treats the performance task of Section 9.3 and simulation-embedded behavioural traces as the primary evidence of the suitability-appraisal and suitability-governed-reliance dimensions, and uses the self-report scale as a convergent, not a defining, measure.

A validation plan would proceed in two stages. First, content validity, with a panel of experts rating item relevance to produce a content-validity index, followed by cognitive interviews with students to refine wording. Second, construct validity, with exploratory factor analysis on one sample and confirmatory factor analysis on an independent sample to test whether the four dimensions are recoverable, tests of measurement invariance across finance and non-finance students, and discriminant-validity tests against financial literacy, financial capability, financial self-efficacy, digital financial literacy, generic AI literacy, appropriate reliance (trust in automation), and metacognition and self-regulated learning. Two further tests are decisive for the construct. To guard against circularity, because the self-report scale adapts the Ng et al. (2024) affective-behavioural-cognitive-ethical instrument, the validation must show that AIL items load on factors separate from their generic AI-literacy parent items rather than reproducing the generic structure. And to test Proposition 9, an incremental-validity analysis must show that AIL predicts reliance-decision quality after generic AI literacy, financial-suitability knowledge, and appropriate-reliance calibration are entered additively as controls. The behavioural and scenario measures would triangulate the self-report scale, so that the competency is not assessed by perception alone, and the rubric of Section 9.3 would itself be validated through inter-rater reliability and by checking that rubric scores align with the item pool and the behavioural traces. This plan operationalises Propositions 1 and 9 and gives the construct a direct route from theory to measurement.

Discussion and contributions

The paper makes four theoretical contributions. To the financial-literacy literature, it argues that the binding constraint on good decisions has shifted from the individual's stock of knowledge to the individual's capacity to appraise an external adviser, and it names the competency that this shift requires (Goyal & Kumar, 2021; Niszczoła & Abbas, 2023). To the financial-capability literature, it extends the relational view of competence to a new counterpart, the algorithmic adviser, showing that the fit that matters is no longer only between a person and

financial institutions but between a person and an opaque advice-giving system (Sherraden, 2013). To the AI-literacy literature, it shows where the domain-neutral frameworks are incomplete rather than merely instantiating them: generic evaluation asks whether an output is sound in itself and has no vocabulary for whether it suits a particular person, so AIL adds a personal-suitability judgement that governs reliance, a competency the generic frameworks cannot express (Ng et al., 2021; Chiu et al., 2024). To investment and finance education, it supplies a defined and assessable target that connects experiential pedagogy to a clear competency, rather than treating simulations as ends in themselves (Bakoush, 2022; Marriott et al., 2015). The single move running through all four is the relocation of competence from knowing to judging.

The practical implications follow directly and address the audiences finance-education research must serve. For curriculum and course design, AIL can be treated as an explicit learning outcome and sequenced across a programme rather than left implicit. For assessment design, evaluation can move beyond knowledge recall toward the appraisal and calibration tasks that AIL comprises, using the worked task and rubric in Section 9.3 as a starting template. For instructor training and professional development, the construct implies that teaching AIL depends on the instructor's own AI-related technological pedagogical content knowledge, which becomes a target for faculty development (Ning et al., 2024). For the choice and design of financial-technology tools used in teaching, it implies favouring tools whose reasoning can be surfaced and interrogated, so that students can practise appraisal rather than mere use. For an era in which the least literate learners may be the most willing to follow algorithmic advice (Piehlmaier, 2022), making AIL a deliberate objective is a matter of protecting students, not merely upgrading a syllabus.

Two limitations bound these contributions. The framework is developed at a level of generality that abstracts from national and institutional context, and the balance among the dimensions may differ across programme types, regulatory settings, and cultures, so the propositions should be tested before the construct is assumed to travel. The framework is also conceptual; although Section 9 offers a concrete route to measurement and a classroom-ready task, the dimensionality of the construct and its separation from trust in automation and the adjacent literacies remain to be demonstrated empirically. We present AIL, accordingly, as a construct to be tested rather than a finding to be applied.

Research agenda

The propositions are more valuable as falsifiable hypotheses than as claims, and each implies a design. First, instrument development and validation of an AIL scale, including content validity, exploratory and confirmatory factor analysis, and measurement invariance across disciplines, with rubrics and behavioural measures to complement self-report. Second, discriminant-validity and dimensionality studies that test whether AIL is separable from financial literacy, financial capability, financial self-efficacy, digital financial literacy, generic AI literacy, appropriate reliance (trust in automation), and metacognition and self-regulated learning, and whether it explains reliance-decision quality beyond those neighbours combined (Proposition 9), and whether its four dimensions are recoverable as distinct factors, a direct test of Proposition 1. Third, design-based and quasi-experimental classroom studies comparing instructional designs and instructor preparation for AIL gains, addressing Propositions 2 and 3, using pre-test and post-test and, where possible, longitudinal designs. Fourth, mediation and moderation studies of the calibrated-reliance mechanism and its boundary conditions, addressing Propositions 5 to 7, using vignette experiments, think-aloud protocols, and

simulation-embedded traces. Fifth, longitudinal and transfer studies that follow students from the classroom into real financial decisions, addressing Proposition 8 and connecting the construct to downstream financial well-being. Across these designs, the construct becomes tractable to the extent that suitability-governed reliance can be measured behaviourally rather than only reported.

Conclusion

The student who receives an opaque, confident, machine-generated portfolio recommendation faces a problem that investment education has not yet named. It is not that she knows too little finance; a machine can now supply the finance. It is that she has not been taught to understand, evaluate, calibrate, and ethically use the advice a machine gives her. Algorithmic Investment Literacy names that competency, specifies its four dimensions, distinguishes it from its neighbours by the mechanism of suitability-governed reliance on opaque advice under consequential stakes, and offers a progression, a classroom-ready task and rubric, an assessment blueprint, and a testable framework. The invitation to the field is to measure it, to teach it, and to test whether learners who possess it decide better.

Declaration of competing interest

The authors declare no competing financial or personal interests.

Funding

This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Declaration of generative AI and AI-assisted technologies in the writing process

During the preparation of this work the authors used generative AI tools in order to assist with drafting, literature scanning, and figure preparation. After using these tools, the authors reviewed and edited the content as needed and take full responsibility for the content of the publication. All substantive claims and citations were reviewed and verified by the authors, and the intellectual contribution, including the construct, its dimensions, and the propositions, is the authors' own.

CRedit authorship contribution statement

Fahmi Abdul Rahim: Conceptualization, Methodology, Writing - original draft, Writing - review and editing, Supervision. Nur Hafidzah Idris: Conceptualization, Validation, Writing - review and editing. Amirudin Mohd Nor: Validation, Writing - review and editing. Rohaiza Kamis: Validation, Writing - review and editing. Mohd Hafiz Bakar: Validation, Writing - review and editing. Faezah Othman: Validation, Writing - review and editing.

Data availability

No data were used for the research described in this article.

References

- Al Rahahleh, N. (2023). Determinants of the financial capability: The mediating role of financial self-efficacy and financial inclusion. *International Journal of Economics and Financial Issues*, 13(6), 15-29. <https://doi.org/10.32479/ijefi.14531>
- Allen, L. K., & Kendeou, P. (2024). ED-AI Lit: An interdisciplinary framework for AI literacy in education. *Policy Insights from the Behavioral and Brain Sciences*, 11(1), 3-10. <https://doi.org/10.1177/23727322231220339>
- Archambault, L. M., & Barnett, J. H. (2010). Revisiting technological pedagogical content knowledge: Exploring the TPACK framework. *Computers & Education*, 55(4), 1656-1662. <https://doi.org/10.1016/j.compedu.2010.07.009>
- Aristei, D., & Gallo, M. (2025). Financial literacy, robo-advising, and the demand for human financial advice: Evidence from Italy. *Journal of Behavioral and Experimental Finance*, 49, 101125. <https://doi.org/10.1016/j.jbef.2025.101125>
- Bakoush, M. (2022). Evaluating the role of simulation-based experiential learning in improving satisfaction of finance students. *The International Journal of Management Education*, 20(3), 100690. <https://doi.org/10.1016/j.ijme.2022.100690>
- Bandura, A. (1997). *Self-efficacy: The exercise of control*. W. H. Freeman.
- Burton, J. W., Stein, M.-K., & Jensen, T. B. (2020). A systematic review of algorithm aversion in augmented decision making. *Journal of Behavioral Decision Making*, 33(2), 220-239. <https://doi.org/10.1002/bdm.2155>
- Chee, H., Ahn, S., & Lee, J. (2025). A competency framework for AI literacy: Variations by different learner groups and an implied learning pathway. *British Journal of Educational Technology*, 56(5), 2146-2182. <https://doi.org/10.1111/bjet.13556>
- Chiu, T. K. F., Ahmad, Z., Ismailov, M., & Sanusi, I. T. (2024). What are artificial intelligence literacy and competency? A comprehensive framework to support them. *Computers and Education Open*, 6, 100171. <https://doi.org/10.1016/j.cao.2024.100171>
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114-126. <https://doi.org/10.1037/xge0000033>
- Downen, T., Kim, S., & Lee, L. (2024). Algorithm aversion, emotions, and investor reaction: Does disclosing the use of AI influence investment decisions? *International Journal of Accounting Information Systems*, 52, 100664. <https://doi.org/10.1016/j.accinf.2024.100664>
- Filiz, I., Judek, J. R., Lorenz, M., & Spiwoks, M. (2023). The extent of algorithm aversion in decision-making situations with varying gravity. *PLOS ONE*, 18(2), e0278751. <https://doi.org/10.1371/journal.pone.0278751>
- Flavell, J. H. (1979). Metacognition and cognitive monitoring: A new area of cognitive-developmental inquiry. *American Psychologist*, 34(10), 906-911. <https://doi.org/10.1037/0003-066X.34.10.906>
- Goyal, K., & Kumar, S. (2021). Financial literacy: A systematic review and bibliometric analysis. *International Journal of Consumer Studies*, 45(1), 80-105. <https://doi.org/10.1111/ijcs.12605>
- Koehler, M. J., Mishra, P., & Cain, W. (2013). What is technological pedagogical content knowledge (TPACK)? *Journal of Education*, 193(3), 13-19. <https://doi.org/10.1177/002205741319300303>
- Koskelainen, T., Kalmi, P., Scornavacca, E., & Vartiainen, T. (2023). Financial literacy in the digital age: A research agenda. *Journal of Consumer Affairs*, 57(1), 507-528. <https://doi.org/10.1111/joca.12510>

- Kruger, J., & Dunning, D. (1999). Unskilled and unaware of it: How difficulties in recognizing one's own incompetence lead to inflated self-assessments. *Journal of Personality and Social Psychology*, 77(6), 1121-1134. <https://doi.org/10.1037/0022-3514.77.6.1121>
- Küper, A., & Krämer, N. (2025). Psychological traits and appropriate reliance: Factors shaping trust in AI. *International Journal of Human-Computer Interaction*, 41(7), 4115-4131. <https://doi.org/10.1080/10447318.2024.2348216>
- Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50-80. https://doi.org/10.1518/hfes.46.1.50_30392
- Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1-16). Association for Computing Machinery. <https://doi.org/10.1145/3313831.3376727>
- Lyons, A. C., & Kass-Hanna, J. (2021). A methodological overview to defining and measuring digital financial literacy. *Financial Planning Review*, 4(2), e1113. <https://doi.org/10.1002/cfp2.1113>
- Marriott, P., Tan, S. M., & Marriott, N. (2015). Experiential learning: A case study of the use of computerised stock market trading simulation in finance education. *Accounting Education*, 24(6), 480-497. <https://doi.org/10.1080/09639284.2015.1072728>
- Nelson, T. O., & Narens, L. (1990). Metamemory: A theoretical framework and new findings. In G. H. Bower (Ed.), *The psychology of learning and motivation* (Vol. 26, pp. 125-173). Academic Press. [https://doi.org/10.1016/S0079-7421\(08\)60053-5](https://doi.org/10.1016/S0079-7421(08)60053-5)
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- Ng, D. T. K., Wu, W., Leung, J. K. L., Chiu, T. K. F., & Chu, S. K. W. (2024). Design and validation of the AI literacy questionnaire: The affective, behavioural, cognitive and ethical approach. *British Journal of Educational Technology*, 55(3), 1082-1104. <https://doi.org/10.1111/bjet.13411>
- Ning, Y., Zhang, C., Xu, B., Zhou, Y., & Wijaya, T. T. (2024). Teachers' AI-TPACK: Exploring the relationship between knowledge elements. *Sustainability*, 16(3), 978. <https://doi.org/10.3390/su16030978>
- Niszczoła, P., & Abbas, S. (2023). GPT has become financially literate: Insights from financial literacy tests of GPT and a preliminary test of how people use it as a source of advice. *Finance Research Letters*, 58, 104333. <https://doi.org/10.1016/j.frl.2023.104333>
- OECD/INFE. (2024). *OECD/INFE survey instrument to measure digital financial literacy*. OECD Publishing. <https://doi.org/10.1787/548de821-en>
- Parasuraman, R., & Riley, V. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39(2), 230-253. <https://doi.org/10.1518/001872097778543886>
- Pattanaik, S., & Hota, S. L. (2025). Exploring the nexus of robotic advisory, digital fluency, and financial literacy in the future investment landscape: A systematic review. *Folia Oeconomica Stetinensia*, 25(1), 240-258. <https://doi.org/10.2478/fofi-2025-0012>
- Piehlmaier, D. M. (2022). Overconfidence and the adoption of robo-advice: Why overconfident investors drive the expansion of automated financial advice. *Financial Innovation*, 8(1), 14. <https://doi.org/10.1186/s40854-021-00324-3>
- Schmidt, D. A., Baran, E., Thompson, A. D., Mishra, P., Koehler, M. J., & Shin, T. S. (2009). Technological pedagogical content knowledge (TPACK): The development and validation

- of an assessment instrument for preservice teachers. *Journal of Research on Technology in Education*, 42(2), 123-149. <https://doi.org/10.1080/15391523.2009.10782544>
- Shanmuganathan, M. (2020). Behavioural finance in an era of artificial intelligence: Longitudinal case study of robo-advisors in investment decisions. *Journal of Behavioral and Experimental Finance*, 27, 100297. <https://doi.org/10.1016/j.jbef.2020.100297>
- Sherraden, M. S. (2013). Building blocks of financial capability. In J. M. Birkenmaier, M. S. Sherraden, & J. C. Curley (Eds.), *Financial capability and asset development: Research, education, policy, and practice* (pp. 3-43). Oxford University Press.
- Xiao, J. J., Huang, J., Goyal, K., & Kumar, S. (2022). Financial capability: A systematic conceptual review, extension and synthesis. *International Journal of Bank Marketing*, 40(7), 1680-1717. <https://doi.org/10.1108/IJBM-05-2022-0185>
- Yuan, G. (2026). Synergistic integration of artificial intelligence and gamification in university finance trading simulation labs. *Simulation & Gaming*. Advance online publication. <https://doi.org/10.1177/10468781251377585>
- Zerilli, J., Bhatt, U., & Weller, A. (2022). How transparency modulates trust in artificial intelligence. *Patterns*, 3(4), 100455. <https://doi.org/10.1016/j.patter.2022.100455>
- Zimmerman, B. J. (2002). Becoming a self-regulated learner: An overview. *Theory Into Practice*, 41(2), 64-70. https://doi.org/10.1207/s15430421tip4102_2