

GENERATIVE AI LITERACY AMONG INDIAN SECONDARY-SCHOOL STUDENTS

Awareness, Usage, Verification and Risk Perception

Aryaveer Singh

Founder & CEO, GENESIS Foundation | Independent Student Researcher

Submission-Ready Research Protocol Manuscript

Empirical results pending genuine primary-data collection

August 2026

Research-integrity statement

This manuscript presents a literature-grounded study protocol, complete survey instrument, coding framework, and analysis plan. It does not claim that participants have been recruited or that data have been collected. All empirical results, statistical estimates, quotations, and participant-data graphs remain intentionally unreported until genuine primary data are available.

Abstract

Generative artificial intelligence (GenAI) tools can produce fluent explanations, summaries, images, code, and answers, but fluent output is not necessarily accurate, unbiased, private, or educationally appropriate. For secondary-school students, meaningful literacy therefore extends beyond awareness or frequency of use. It includes functional understanding, purposeful use, critical verification, calibrated trust, recognition of risk, and responsible behaviour. This paper develops an academically grounded protocol for investigating these dimensions among students in Classes IX and X in India. The proposed cross-sectional, multi-site survey uses a mixed-format instrument combining objective knowledge items, self-reported competence, behavioural-frequency measures, and scenario-based verification tasks. The conceptual model treats GenAI literacy as five related constructs: functional AI understanding, applied AI competence, critical evaluation and verification, risk and ethical awareness, and responsible AI behaviour. Access, school context, language, prior instruction, and usage patterns are treated as contextual predictors, while trust is examined as a calibration variable rather than as an uncomplicated indicator of literacy. The protocol specifies a realistic independent-student sampling strategy, pilot testing, translation procedures, content and construct validation, reliability analysis, ethical safeguards for minors, and an analysis plan that accounts for clustering and multiple comparisons. It also supplies participant materials, a complete draft questionnaire, and a variable-coding framework. No empirical findings are reported because primary data have not yet been collected. The study is designed to produce transparent, reproducible evidence about how Indian secondary-school students use GenAI, how they judge its outputs, and what forms of school-based instruction could strengthen safe, critical, and responsible use.

Keywords: generative artificial intelligence; AI literacy; secondary education; verification; risk perception; India; digital literacy; research protocol

Manuscript status

Protocol complete; recruitment and data collection not yet undertaken. The Results section contains prespecified shells only, so that future analyses can be reported without inventing or retrospectively selecting outcomes.

1. Introduction

Generative artificial intelligence has moved quickly from a specialist technology to an everyday interface. Text-generating systems can respond to questions, create explanations, translate language, draft essays, write computer code, and simulate dialogue. Image and multimodal systems extend these capabilities across media. For school students, such tools may support brainstorming, practice, feedback, language access, and explanation. At the same time, the systems can produce fabricated claims and references, reproduce bias, encourage the disclosure of personal information, and make it difficult to distinguish assistance from inappropriate substitution of a student's own work. The educational question is therefore not simply whether adolescents have heard of GenAI or have used it. The more consequential question is whether they can use it with understanding, verification, restraint, and accountability.

This distinction matters because generative systems create a persuasive surface. Their output is usually coherent and often confident, even when it is incomplete or wrong. The technical process that produces a plausible sequence of words is not the same as a process that guarantees factual truth. Research on hallucination in natural-language generation and experimental work on AI-generated information show why users must evaluate claims rather than equate fluency with reliability (Ji et al., 2023; Spitale et al., 2023). Human–AI interaction research also suggests that overreliance can persist unless users are prompted to think independently and examine evidence (Buçinca et al., 2021). These problems are especially relevant to students who are still developing subject knowledge, source-evaluation routines, and norms for authorship.

AI literacy offers a useful foundation, but the construct has been defined in different ways. Early frameworks emphasized the competencies needed to understand, evaluate, communicate about, and interact with AI (Long & Magerko, 2020). Subsequent reviews identified technical, cognitive, ethical, and social dimensions, while newer GenAI-specific work gives greater attention to prompting, evaluation, content creation, limitations, and responsible use (Almatrafi et al., 2024; Annapureddy et al., 2025; Ng et al., 2021). Measurement studies further show that self-reported confidence cannot substitute for objective knowledge or demonstrated evaluation skill (Chiu et al., 2024b; Lintner, 2024). A student may feel capable because a chatbot is easy to operate, yet still accept invented citations or disclose sensitive information. Conversely, a cautious student may report low confidence while showing strong judgement. A credible study must therefore measure several dimensions and use more than one item format.

India provides a particularly important setting. The National Education Policy 2020 and the National Curriculum Framework for School Education 2023 emphasize digital capability, critical thinking, ethics, and responsible participation in a technology-rich society (Ministry of Human Resource Development, 2020; National Council of Educational Research and Training, 2023). The Central Board of Secondary Education’s current Artificial Intelligence curriculum for Classes IX and X includes AI concepts, the AI project cycle, data literacy, bias, access, and ethical concerns (Central Board of Secondary Education [CBSE], 2026a, 2026b). UNESCO’s India-focused education report similarly highlights both the opportunities and governance challenges of AI in education (UNESCO New Delhi, 2022). Yet curriculum availability does not establish what students across different schools know, how often they use public GenAI tools, or whether they verify what those tools produce.

The Indian secondary-school population is heterogeneous. Access to devices, connectivity, English-language proficiency, formal AI instruction, and adult guidance varies across regions and school contexts. Evidence from the ASER 2023 Beyond Basics survey documents uneven digital access and skill among rural youth aged 14–18, although its district-based rural sample should not be treated as a complete picture of all Indian adolescents (ASER Centre, 2024). This heterogeneity makes nationally sweeping claims inappropriate for a modest independent-student project. It also makes carefully described, multi-site evidence valuable if the limits of the sampling frame are respected.

The present paper therefore develops a complete research protocol for a study of GenAI literacy among Indian students in Classes IX and X. It preserves a narrow empirical focus—awareness, access, usage, verification, trust, risk perception, and responsible behaviour—while grounding these observations in a multidimensional literacy model. The design deliberately separates what students say they can do from what they know and how they respond to short scenarios. It also treats verification as an active practice: checking important claims against reliable independent sources, opening and inspecting citations, comparing answers with curriculum materials or knowledgeable people, and recognizing when uncertainty requires caution.

The contribution is threefold. Conceptually, the paper integrates general AI literacy, GenAI-specific competence, information evaluation, and child-centred risk principles into a framework suitable for secondary-school research. Methodologically, it supplies an original survey instrument, validation pathway, ethical protocol, and prespecified analysis plan that an independent student researcher can implement with school and adult oversight. Substantively, the future study is positioned to identify gaps between use and judgement—for example, frequent use paired with weak verification—without presuming that such gaps exist before data are collected. This protocol approach protects research integrity: it advances the paper as far as current evidence permits and stops before reporting participants or findings that do not yet exist.

2. Literature Review

2.1 From AI awareness to multidimensional literacy

Awareness is the weakest threshold of literacy. Knowing that a system called ChatGPT or another GenAI tool exists reveals little about whether a student understands why outputs vary, how prompts influence responses, or why personal information should not be entered. Long and Magerko (2020) framed AI literacy as a set of competencies that enable people to evaluate AI, communicate and collaborate with it, and use it effectively. Ng et al. (2021) organized AI literacy around knowing and understanding AI, using and applying it, evaluating and creating with it, and considering ethical issues. These accounts reject a purely operational definition and position literacy as a combination of knowledge, skill, critical judgement, and social understanding.

Later reviews show both consolidation and unresolved diversity. Laupichler et al. (2022) found that AI literacy education often combines basic concepts, applications, ethical issues, and attitudes, but that definitions and assessment practices vary. Almatrafi et al. (2024) similarly documented variation in constructs, implementation, and measurement across studies published from 2019 to 2023. Pinski and Benlian (2024) identified learning components and effects across a broad user-literacy literature. Together, these reviews support a family of related dimensions rather than one universally accepted score. They also warn against assembling an index without explaining whether its components are reflective indicators of a common trait or distinct capabilities that together form a broader competence.

GenAI introduces additional demands within broader AI literacy and competency frameworks (Chiu et al., 2024a). Unlike many predictive systems operating in the background, generative tools invite direct conversation and content production. Annapureddy et al. (2025) proposed twelve competencies for GenAI literacy, reflecting the need to understand system capabilities, craft and refine interactions, assess outputs, recognize social implications, and use generated material responsibly. UNESCO's student competency framework locates AI learning within a human-centred mindset, ethics, AI techniques and applications, and system design, with progression from understanding to applying and creating (UNESCO, 2024). For school research, the implication is clear: awareness, use, verification, and responsibility should be measured separately before they are combined, and any combined score should be theoretically and empirically justified.

2.2 Measuring AI literacy

Existing scales provide useful precedents but do not remove the need for population-specific validation. Wang et al. (2023) developed an AI literacy scale focused on user competence, and Laupichler et al. (2023) developed the Scale for the Assessment of Non-Experts' AI Literacy. Such instruments demonstrate that nontechnical populations can be studied systematically. However, a systematic review of AI literacy scales found substantial differences in conceptual coverage and evidence of validity (Lintner, 2024). A scale validated with adults, university students, or users of earlier AI applications cannot automatically be assumed to function identically among Indian students in Classes IX and X.

A central measurement problem is the difference between perceived and demonstrated competence. Chiu et al. (2024b) developed objective AI literacy test items for Grades 7–9 and used Rasch analysis with a large sample, illustrating that age-appropriate knowledge can be tested rather than inferred from self-confidence. Self-report remains valuable for frequency, intention, perceived competence, and experiences that are not directly observable in a short survey. Yet it is vulnerable to social desirability, response style, and inaccurate self-assessment. This protocol therefore combines four response modes: objective multiple-choice knowledge questions; agreement ratings for perceived applied competence and risk awareness; behavioural-frequency items for verification and responsible practice; and scenarios requiring the respondent to choose an action. Convergence across formats would strengthen interpretation; divergence would itself be substantively informative.

Measurement should also distinguish reflective and formative uses of items. Several items describing the same underlying verification tendency may form a reflective scale if they covary and show appropriate dimensionality. In contrast, a checklist of distinct risks—privacy, bias, academic integrity, misinformation, and overreliance—may be better treated as a profile or formative index because recognizing one risk does not guarantee recognizing another. Reporting only Cronbach's alpha for every group of items would therefore be methodologically weak. The analysis plan includes omega or alpha for plausible reflective scales, item-level reporting, factor analysis where justified, and separate treatment of heterogeneous knowledge and risk indicators.

2.3 Student use, educational opportunity, and learning risk

Educational reviews describe plausible benefits of conversational GenAI: rapid feedback, individualized explanation, examples at different difficulty levels, brainstorming, translation, and support for drafting (Lo, 2023; Tlili et al., 2023). Broader evaluations also emphasize that the same affordances create weaknesses and threats when systems are used without sound pedagogy or verification (Farrokhnia et al., 2024). Student-perception research in higher education reports enthusiasm alongside concerns about accuracy, dependency, fairness, privacy, and future skills (Chan & Hu, 2023; Chan & Lee, 2023). These studies cannot be transferred uncritically to younger students, but they identify dimensions that secondary-school research should examine. In particular, ease of access can obscure the intellectual work required to formulate a problem, judge an answer, and decide what to retain.

K–12 evidence is growing but remains uneven across age groups and settings. Klar’s (2025) mixed-method study of K–12 students found that using a chatbot interface may be easy while using it effectively for learning remains demanding. This distinction between interface fluency and learning competence is central to the present protocol. A student who can obtain a polished answer has demonstrated access and basic operation; a student who can refine a question, interrogate weaknesses, compare sources, and integrate verified information has demonstrated a more advanced educational practice.

Risks are not limited to factual error. Akgun and Greenhow (2022) identify ethical challenges for K–12 education, including privacy, bias, autonomy, transparency, and fairness. Large language models may reproduce social stereotypes or uneven representation from their data and design (Bender et al., 2021). Memorization and data-extraction research shows why users should not assume that information entered into or produced by AI systems is harmless or private (Carlini et al., 2021). UNICEF’s policy guidance for AI and children places safety, privacy, fairness, inclusion, transparency, well-being, and children’s agency at the centre of governance (UNICEF, 2021). For adolescents, responsible literacy must therefore include the practical ability to avoid sharing sensitive data, to recognize biased or harmful output, and to seek help when stakes are high.

Academic integrity adds a school-specific dimension. GenAI can assist learning when students use it to generate examples, receive feedback, or test understanding, but it can replace learning when unverified output is submitted as personal work. Scholarship on ChatGPT and integrity emphasizes the need for assessment redesign, transparent expectations, and disclosure rather than reliance on detection alone (Cotton et al., 2024). Detection tools themselves can generate false positives and false negatives, making them an unsafe substitute for clear process evidence and fair educational practice (Dalalah & Dalalah, 2023). A student survey should therefore ask about knowledge of school rules and disclosure behaviour without using accusatory wording or attempting to identify individual misconduct.

2.4 Verification, information evaluation, and calibrated trust

Verification is the bridge between receiving an output and responsibly acting on it. Generative systems can produce non-existent sources, inaccurate quotations, outdated details, and correct statements supported by weak reasoning. Students need routines that are specific enough to be actionable: identify which claims matter; look for the claim in an authoritative or primary source; open rather than merely count citations; compare multiple independent sources; examine dates and context; and consult a teacher, textbook, or subject expert when uncertainty remains. These practices connect AI literacy to established work on digital source evaluation.

Lateral reading is particularly relevant. Instead of remaining within a single page or answer, a user leaves it to investigate the source, organization, author, and corroborating evidence. Wineburg and McGrew (2019) found that expert fact-checkers read in this outward, networked manner. In a GenAI setting, lateral reading means treating the generated response as a claim set rather than as a source of record. The user searches for the underlying evidence, verifies whether named studies exist, and checks whether those studies support the claim attributed to them. This protocol operationalizes verification with both self-reported frequency and scenarios in which the safest response is to inspect independent evidence.

Trust is not equivalent to literacy. Total distrust can prevent students from benefiting from a useful tool, while unconditional trust can produce overreliance. The desired state is calibrated trust: confidence that changes according to the task, evidence, system limitations, and consequences of error. Bućinca et al. (2021) showed that cognitive forcing interventions can reduce overreliance in AI-assisted decision making, suggesting that a deliberate pause for independent reasoning can matter. In this study, trust is therefore analysed as a separate and potentially non-linear variable. High literacy may be associated with moderate, conditional trust rather than the maximum score on a general trust item.

2.5 Indian policy and curricular context

Indian education policy provides a strong rationale for studying GenAI literacy but not an empirical answer. The National Education Policy 2020 promotes critical thinking, digital literacy, coding, computational thinking, and ethical reasoning (Ministry of Human Resource Development, 2020). The National Curriculum Framework for School Education 2023 emphasizes competencies, interdisciplinary learning, and responsible engagement with technology (National Council of Educational Research and Training, 2023). These priorities align with an understanding of AI literacy that includes reflection and judgement, not simply technical familiarity.

CBSE's Artificial Intelligence subject for Classes IX and X provides the clearest curricular reference point. The current curricula address AI concepts, data literacy, the AI project cycle, bias, access, and ethics, with more advanced application and programming content in Class X (CBSE, 2026a, 2026b). However, the subject is not necessarily taken by every student, and exposure within a curriculum does not indicate mastery. Prior formal instruction must therefore be measured as a contextual variable rather than assumed from grade or board. The survey also needs to be understandable to students who have never used a public GenAI system; excluding non-users would inflate apparent access and competence.

UNESCO's mapping of government-endorsed K–12 AI curricula shows that school AI education is an international policy concern, while its guidance on GenAI in education calls for human-centred, age-appropriate, privacy-aware use and institutional safeguards (UNESCO, 2022, 2023). UNESCO New Delhi (2022) situates these debates in India and highlights the need to align innovation with inclusion and governance. The OECD (2023) similarly describes effective digital education as an ecosystem involving infrastructure, capacity, governance, evidence, and coordination. Thus, a survey of individual students should not interpret low competence as a personal deficit alone; it may reflect unequal opportunities, unclear rules, language barriers, or the absence of structured instruction.

The targeted literature review located limited peer-reviewed evidence that jointly measures GenAI use, objective understanding, verification performance, calibrated trust, risk perception, and responsible behaviour specifically among Indian Classes IX–X students. This is a bounded gap statement, not a claim that no relevant Indian study exists. The rapid development of the field and uneven indexing mean that the review should be updated immediately before journal submission. The proposed study addresses the identified gap at an appropriate scale: it seeks multi-site evidence and transparent subgroup comparisons but does not claim national representativeness.

2.6 Synthesis and study rationale

The literature supports five decisions. First, literacy is multidimensional and must not be reduced to tool awareness or usage frequency. Second, objective and scenario-based items are needed alongside self-report. Third, verification is a distinct practice that connects technical understanding with responsible action. Fourth, trust should be calibrated to evidence and consequences rather than maximized. Fifth, student behaviour is shaped by opportunity and school context, so interpretation must account for access, language, prior instruction, and institutional rules. These decisions inform the framework, instrument, and hypotheses presented below.

3. Conceptual Framework

The proposed framework defines GenAI literacy as a set of interrelated capacities that enable a secondary-school student to understand generative systems at an age-appropriate level, use them purposefully, evaluate outputs, recognize risks, and act responsibly. It is a research model rather than a claim that a single latent trait explains every behaviour. The dimensions are examined separately first; a higher-order composite will be considered only if theoretical coherence and empirical structure support it.

Table 1. Proposed constructs and measurement approach

Construct	Operational meaning	Primary evidence
Functional AI understanding	Age-appropriate knowledge of what GenAI does, why outputs can vary, and why fluent output can be inaccurate or biased.	Objective multiple-choice items
Applied AI competence	Perceived ability to formulate, refine, and compare prompts and to select appropriate learning tasks.	Agreement-scale items
Critical evaluation and verification	Routine and demonstrated ability to check important claims, sources, dates, and uncertainty using independent evidence.	Frequency items and scenarios
Risk and ethical awareness	Recognition of privacy, bias, misinformation, overreliance, copyright, disclosure, and academic-integrity concerns.	Knowledge and agreement items
Responsible AI behaviour	Reported practices that retain human judgement, protect data, follow school rules, disclose assistance, and seek help when needed.	Behavioural-frequency and intention items
Trust calibration	Conditional confidence in output that varies with evidence, task, and consequence; neither maximum trust nor total distrust.	General and task-sensitive trust items

Note. Construct labels are study-specific and informed by prior AI literacy, GenAI literacy, information-evaluation, and child-centred risk frameworks.

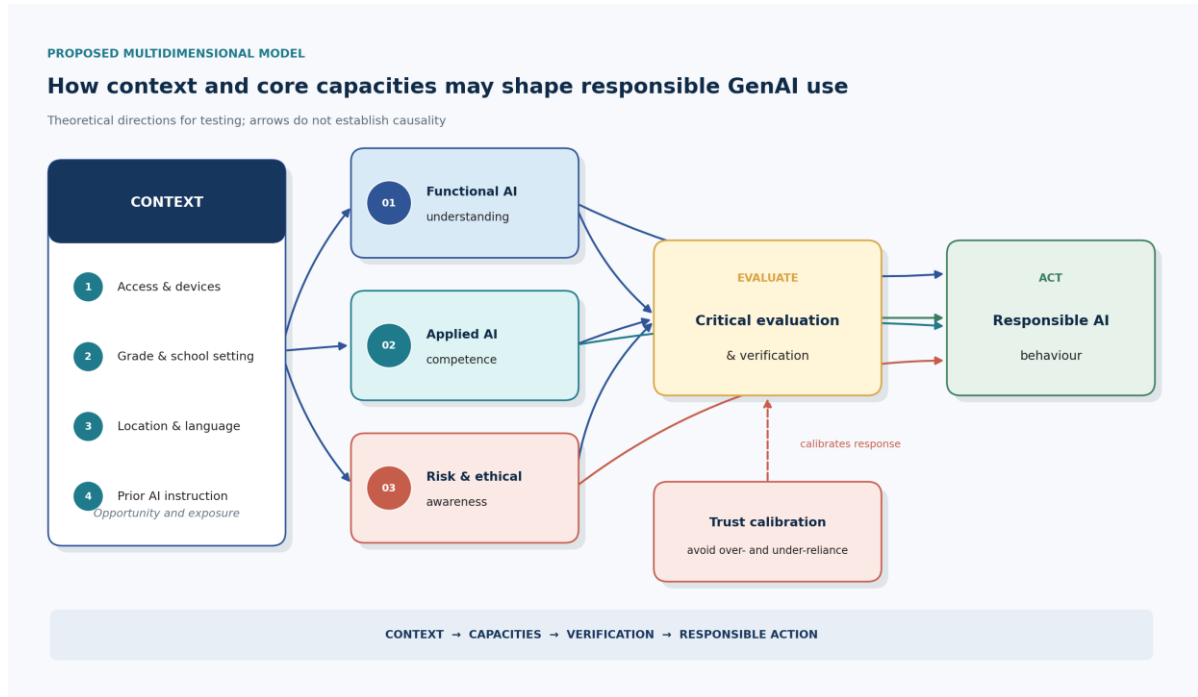


Figure 1. Proposed conceptual framework

Note. Original conceptual diagram developed for this study. Arrows represent hypothesized directional associations, not established causal effects.

Figure 1 positions contextual conditions as antecedents of opportunity and exposure, not as components of literacy. Functional understanding, applied competence, and risk awareness may each support verification and responsible behaviour. Verification provides a central pathway from knowledge to action because it determines whether generated material is accepted, revised, rejected, or investigated. Trust calibration is shown as a conditional influence on this pathway. Because the planned design is cross-sectional, arrows express theoretical direction only; causal claims will not be made from observed associations.

4. Aim, Objectives, Research Questions, and Hypotheses

4.1 Aim

To investigate the awareness, access, usage, functional understanding, verification practices, trust calibration, risk perception, and responsible GenAI behaviour of Indian secondary-school students in Classes IX and X within a transparently bounded multi-site sample.

4.2 Objectives

1. Describe awareness of and access to GenAI tools, frequency and purposes of use, and sources of guidance.
2. Measure functional AI understanding, applied competence, critical evaluation and verification, risk and ethical awareness, and responsible behaviour using mixed item formats.
3. Examine the relationship between usage frequency and the quality of verification, rather than assuming frequent use indicates stronger literacy.
4. Assess whether trust is calibrated to task and evidence and whether it is associated with verification and responsible behaviour.

5. Compare results by class, school type, location, language, device access, and prior formal AI instruction while reporting effect sizes and uncertainty.
6. Identify school-based supports that students report needing for safe, critical, and educationally useful GenAI use.

4.3 Research questions

- RQ1. What are the levels and patterns of GenAI awareness, access, frequency of use, purposes of use, and formal instruction in the participating Classes IX–X sample?
- RQ2. What are the distributions of functional AI understanding, applied AI competence, critical evaluation and verification, risk and ethical awareness, responsible AI behaviour, and trust calibration?
- RQ3. How are usage frequency, applied competence, objective understanding, and verification performance related?
- RQ4. How are risk and ethical awareness, trust calibration, and responsible AI behaviour related?
- RQ5. Do literacy dimensions differ by class, school type, location, language, device access, or prior formal AI instruction?
- RQ6. Which contextual and literacy variables independently predict verification performance and responsible AI behaviour after prespecified covariates and school clustering are considered?
- RQ7. What benefits, concerns, and desired forms of school support do students describe in optional open-ended responses?

4.4 Hypotheses

- H1. Higher functional AI understanding will be positively associated with stronger critical evaluation and verification performance.
- H2. Higher risk and ethical awareness will be positively associated with more responsible AI behaviour.
- H3. Students reporting formal AI instruction will have higher objective understanding and verification performance than students reporting no formal instruction, after adjustment for grade, access, and school context.
- H4. Greater usage frequency will be positively associated with applied AI competence; its association with verification performance will be tested without prespecifying a positive direction.
- H5. Verification performance will mediate part of the association between functional understanding and responsible behaviour, if sample size, temporal caution, and model fit permit an exploratory indirect-effect analysis.

Trust calibration is treated as an exploratory moderator because a simple linear hypothesis would be misleading. Both indiscriminate trust and indiscriminate distrust may reflect poor calibration. The primary analysis will therefore examine task-sensitive trust patterns and, if justified, quadratic or category-based relationships specified before the analysis is run.

5. Methodology

5.1 Design and reporting status

The study is designed as a cross-sectional, multi-site, school-based survey with quantitative and optional qualitative components. It is a protocol at the time of writing. No school has been described as recruited, no participant has been counted, and no response has been analysed. The cross-sectional design is appropriate for estimating distributions and associations within a bounded sample, but it cannot establish temporal order or causality. Future reporting should follow relevant observational-study and survey-reporting guidance and should distinguish prespecified from exploratory analyses.

The intended unit of observation is the individual student, with students nested within schools. The design is not nationally representative. Its value depends on transparent documentation of the recruitment frame, participation rates, school characteristics, and missingness. Any eventual article must name the states or territories and the number and type of participating schools only after permission and disclosure-risk review.

5.2 Setting and target population

The target population is students enrolled in Classes IX and X in participating Indian secondary schools. These grades are selected because students are old enough to encounter independent digital-tool use and academic-authorship expectations, while still being within compulsory school structures where adult permission and safeguarding are central. Eligible schools may include government and private institutions and, where feasible, urban, semi-urban, and rural settings. The feasible sampling frame will depend on schools that provide institutional permission and a supervised route to students.

Students do not need to have used GenAI to participate. Non-use is substantively important because it may reflect lack of access, lack of interest, restrictions, uncertainty, or deliberate choice. The questionnaire uses a short definition and examples but avoids requiring a specific commercial brand. Students who have not used a tool bypass detailed usage items and complete knowledge, scenario, risk, and intention items that do not require personal experience.

5.3 Sampling and sample size

A stratified purposive cluster-sampling approach is realistic for an independent student researcher. The researcher should first seek adult-supervised contact with approximately 8–12 schools that vary, where feasible, by management type and location. Within participating schools, intact Classes IX and X sections may be invited rather than asking teachers to select individual students. Inviting whole eligible sections reduces discretionary selection but does not eliminate school-level self-selection. The resulting sample must be described as a participating-school sample, not as a sample of all Indian secondary-school students.

A provisional target of approximately 500 completed surveys is proposed, aiming for reasonable balance by class and diversity across schools. This is not a claimed enrolment figure. Before fieldwork, a formal power analysis should be documented using the smallest effect of practical interest and the planned primary model. For orientation, detecting a small correlation of about $|r| = .15$ with two-sided alpha .05 and 80% power requires roughly 347 independent observations; clustering, nonresponse, incomplete surveys, and subgroup analyses justify inflation. If fewer schools participate or the achieved sample is materially smaller, the analysis should be simplified and uncertainty emphasized rather than compensated for with stronger claims.

Inclusion criteria are enrolment in Class IX or X at a participating school, comprehension of an approved survey language, required guardian permission, student assent, and voluntary participation. Prespecified exclusion criteria are absent required consent or assent, duplicate records identified without retaining direct identifiers, surveys missing most substantive items, impossible completion patterns based on pilot-derived timing, and failure of a plainly worded attention check. Exclusions must be reported with counts and reasons; they must not be chosen after viewing group differences.

5.4 Instrument development

The questionnaire in Appendix A is an original, study-specific instrument informed by published AI literacy frameworks, measurement research, GenAI competency work, digital source-evaluation research, UNESCO guidance, UNICEF’s child-centred principles, and the Indian curricular context. It does not reproduce proprietary scale items verbatim. Original wording allows alignment with Classes IX–X and the study’s specific constructs, but it also means that validity and reliability must be established before substantive score interpretation.

The instrument contains background and access questions; awareness and usage items; eight objective knowledge questions; five applied-competence items; six verification-frequency items; two verification scenarios; eight risk and ethical-awareness items; six responsible-behaviour items; four trust-calibration items; and three optional open-ended prompts. Behaviour items use a five-point Never-to-Always scale; perceived competence and awareness items use a five-point Strongly-disagree-to-Strongly-agree scale. Consistent anchors reduce cognitive switching. Estimated administration time is 15–20 minutes, to be confirmed during pilot testing.

Figures 3 and 4 make the instrument architecture transparent. They summarize the number and allocation of drafted questionnaire items; they do not contain participant responses, pilot estimates, or study findings.

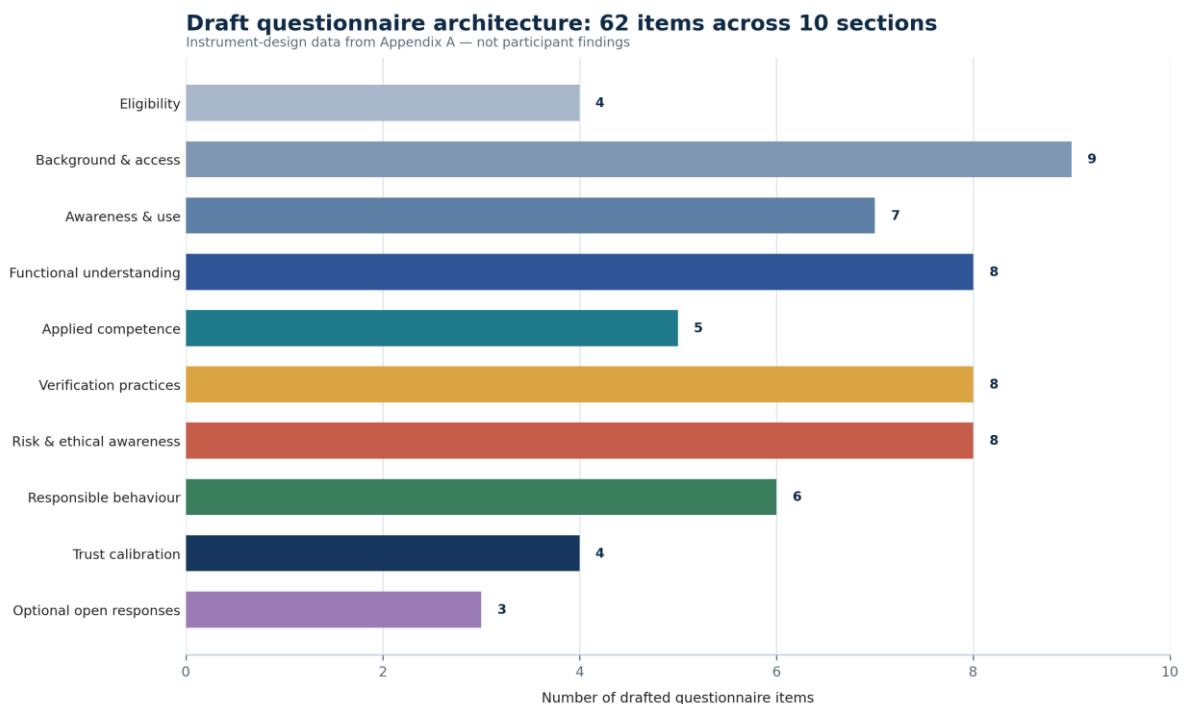


Figure 3. Draft questionnaire item allocation by section

Note. Counts are derived directly from Appendix A (62 drafted items across 10 sections). This is instrument-design information, not participant data.

Composition of the questionnaire's core literacy measures

Item allocation only — not a distribution of student responses

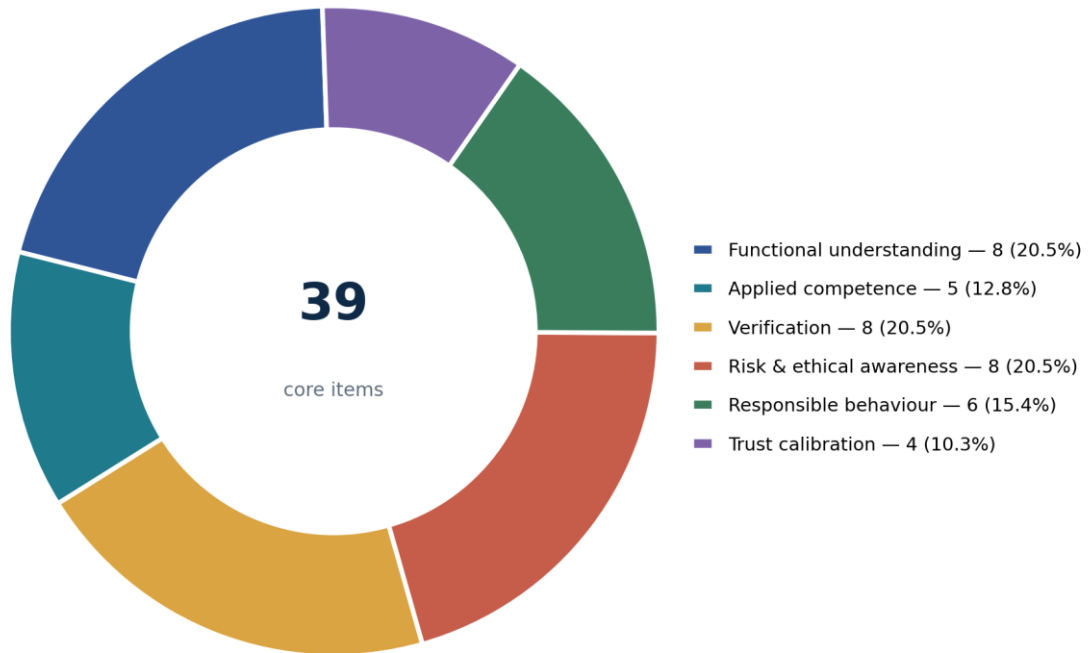


Figure 4. Composition of core and related construct items

Note. The chart shows the allocation of 39 drafted items across six construct domains. Percentages describe questionnaire composition only and do not represent student responses.

The English version should be translated into Hindi if the participating schools require it. Translation should use forward translation by a fluent bilingual educator, review by a second bilingual reviewer, reconciliation, and back-translation of high-stakes or conceptually difficult items. Cognitive interviews should test whether phrases such as ‘independent source,’ ‘personal data,’ ‘fabricated citation,’ and ‘school rules’ are understood as intended. Other languages should be added only when competent translation and review are available; automatic translation alone is insufficient for a research instrument.

Table 2. Questionnaire blueprint

Questionnaire section	Content	Planned analytic use
Background/access	Grade, age band, gender, school context, language, device/connectivity, formal instruction	Description and covariate definition
Awareness/usage	Awareness, ever use, frequency, purposes, tool categories, guidance	Usage profile; skip logic for non-users
Functional understanding	Eight objective items	Item difficulty; prespecified knowledge score if coherent
Applied competence	Five agreement items	Candidate reflective scale
Verification	Six frequency items and two scenarios	Behaviour profile plus performance score

Questionnaire section	Content	Planned analytic use
Risk/ethics	Eight knowledge/agreement items	Profile; subscale only if structure supports it
Responsible behaviour	Six frequency/intention items	Candidate reflective scale
Trust calibration	Four task-sensitive items	Item profile; non-linear exploratory analysis
Open responses	Benefit, concern, requested school support	Inductive coding with deductive framework map

Note. Final item retention and scale composition depend on expert review, cognitive testing, pilot performance, and the main-sample measurement model.

5.5 Pilot testing, validity, and reliability

Instrument development will proceed through three gates. First, a panel of approximately three to five reviewers—ideally including a secondary-school educator, an AI or computing educator, a measurement-informed researcher, and a child-safeguarding or ethics adviser—will judge relevance, clarity, age appropriateness, and construct coverage. Item-level content-validity indices may be calculated if reviewers use a four-point relevance scale, but numerical indices will supplement, not replace, documented comments and revision decisions.

Second, cognitive interviews with about 6–10 students from the target grades but outside the main sample will use think-aloud and probing questions. Students will explain what they believe each item asks, how they chose a response, and whether any term is confusing or uncomfortable. This stage is crucial because an apparently simple item can be interpreted as a test of morality or computing vocabulary rather than the intended construct. Interview notes will be de-identified and used to revise wording and response options.

Third, a pilot survey with approximately 25–40 students from schools or sections not included in the main analysis will assess completion time, missingness, floor and ceiling effects, distractor performance, response distributions, and preliminary internal consistency. This pilot range is practical for usability testing but too small for definitive factor analysis. Items will not be deleted solely to increase alpha. Revision will consider content importance, comprehension, item-total behaviour, and redundancy.

In the main sample, measurement evidence will be examined before hypothesis testing. For candidate reflective scales, dimensionality will be assessed using exploratory or confirmatory factor analysis depending on sample size and the strength of prior specification; ordinal estimators will be preferred for Likert items. McDonald's omega will be the primary internal-consistency estimate, with alpha reported for comparison. Objective items will be reported individually and may be evaluated with KR-20, item-response, or Rasch methods if assumptions and sample size are adequate. Known-groups or convergent evidence may be explored through prior instruction and related constructs, but no single correlation will be labelled proof of validity.

5.6 Recruitment and data-collection procedure

1. Obtain adult research mentorship, school-level permission, and an ethics/safeguarding determination appropriate to the institution and jurisdiction before contacting students.
2. Provide schools with the protocol, survey, participant information, guardian permission form, student assent form, data-management plan, and a plain-language explanation that participation will not affect grades.

3. Distribute guardian materials through the school using the approved consent model. Do not use assumed consent unless explicitly permitted by the responsible authority.
4. Administer the survey in a supervised setting or approved secure online environment. Teachers should not hover over individual responses, and the student researcher should not collect names, roll numbers, personal emails, phone numbers, or account identifiers.
5. Present information and assent before the survey. Students may decline or skip any nonessential question without penalty. A non-user route must remain available.
6. Export the data to an access-restricted, encrypted location; remove platform metadata not needed for analysis; create a processing copy with indirect identifiers minimized; and maintain a reproducible cleaning log.
7. Report the number of schools approached, schools participating, students eligible, permissions returned, assents provided, completed surveys, and exclusions, subject to disclosure-risk limits.

PLANNED RESEARCH WORKFLOW

A staged route from instrument design to defensible evidence

Each stage has a documentary output and a decision gate

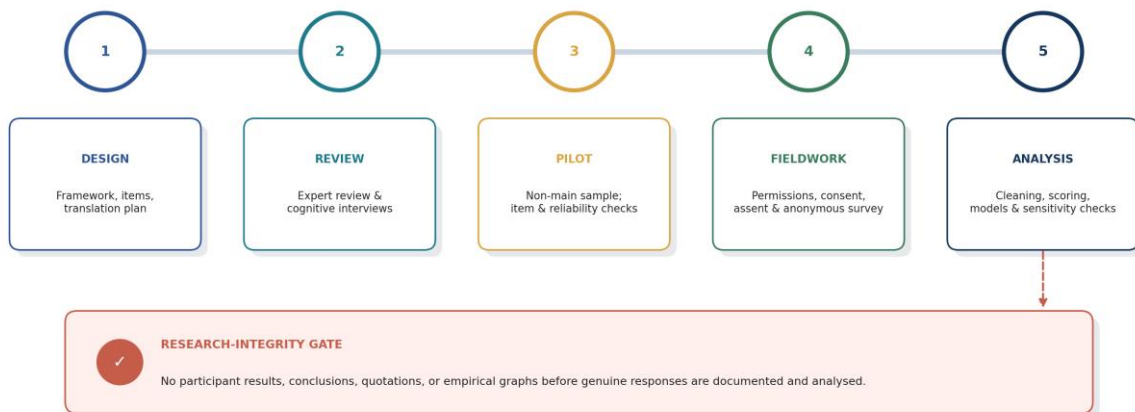


Figure 2. Planned research workflow and evidence gate

Note. Original workflow diagram developed for this protocol.

5.7 Ethical safeguards and data management

Participants are minors, so school permission, guardian permission where required, student assent, voluntariness, privacy, and minimization of risk are non-negotiable. This manuscript does not claim ethics approval. Data collection must not start until a responsible adult mentor and participating institutions determine the appropriate ethics or institutional-review route and approve the materials. If an institutional review board or ethics committee is unavailable, this absence does not justify informal recruitment; the researcher should work through the school's authorized safeguarding and research processes and document the decision pathway.

The survey is designed to avoid direct identifiers. Age is collected in bands, location is coded broadly, and school identity should be replaced by a study code in the analytic file. Combination-based disclosure risk must still be considered: a rare combination of school, grade, language, and demographic characteristics can identify a student even without a name. Small cells should be suppressed or combined in public tables. Open-ended responses must be checked for names and other identifying details before quotation. Direct quotations should be used only with approval and should be lightly de-identified without changing meaning.

The data-management plan should specify where data are stored, who has access, how transfers occur, how long raw and processed data are retained, and when they are securely deleted. A practical plan is to restrict access to the student researcher and named adult mentor; store data in encrypted, password-protected institutional or family-approved storage; keep consent records separate from anonymous survey data; retain de-identified analysis data only for the period needed for assessment or publication; and delete identifiable administrative records according to school or ethics requirements. Exact dates must be inserted before recruitment.

The questionnaire asks about practices, not confessions. It does not ask students to paste AI conversations, reveal accounts, identify teachers, or describe punishable incidents. Wording emphasizes that there are no grade consequences and that honest responses, including non-use, are useful. If a student asks for help about a concerning AI output during administration, the researcher should follow the school's existing safeguarding route rather than attempting to investigate or counsel the student independently.

5.8 Variable scoring and data-quality rules

Scoring decisions will be finalized after pilot revision and frozen in a dated analysis protocol before the main outcome models are run. Objective items will be coded 1 for the keyed response and 0 otherwise; 'not sure' will be treated as incorrect for knowledge scoring but retained as a distinct response in item-level tables. Likert items will retain their ordered 1–5 values. Negatively worded items, if retained after cognitive testing, will be reverse-coded so that higher values consistently indicate stronger competence, awareness, verification, or responsibility. Appendix E provides the preliminary codebook.

Scale scores will be calculated only when a respondent answers a prespecified proportion of items, provisionally at least 80%. If that threshold is met, the mean of completed items can preserve the original response metric; otherwise the scale score will be missing. Scenario items will be scored separately and combined only if they measure a coherent verification construct. A global GenAI literacy score will not be reported by default. It will be considered only if higher-order structure and interpretability are supported; dimension profiles are likely to be more educationally useful.

Quality checks include range checks, internal skip-logic checks, duplicate review using non-identifying technical patterns approved in advance, completion-time flags derived from the pilot, and a single transparent attention check. A flagged survey is not automatically invalid. Exclusion requires a prespecified combination of evidence, and sensitivity analyses will compare results with and without borderline cases. Outliers in legitimate responses will not be removed merely because they affect results.

5.9 Data-analysis plan

Analysis will begin with a reproducible flow diagram and a table comparing available characteristics of included and excluded surveys. Descriptive statistics will report counts, percentages, means and standard deviations when distributions justify them, and medians and interquartile ranges otherwise. All estimates will include confidence intervals where meaningful. Usage categories and risk items will be shown individually to avoid concealing important patterns within a total score.

Psychometric analysis precedes substantive group comparisons. Objective-item difficulty and discrimination will be examined. Candidate Likert scales will be assessed for dimensionality and internal consistency. If a scale is revised after examining the main sample, the change will be labelled exploratory and cross-validation or sensitivity analysis will be considered. Measurement invariance across class or language will be tested only if subgroup sizes and model quality are sufficient; otherwise group comparisons will be framed cautiously at the item or observed-score level.

Bivariate analyses will use chi-square tests for categorical variables, t tests or analysis of variance for approximately continuous scores when assumptions are reasonable, and nonparametric alternatives when they are not. Effect sizes—such as Cramér’s V, standardized mean differences, rank-biserial correlations, or eta squared—will accompany p values. Spearman or Pearson correlations will be chosen based on measurement and distribution. Multiple testing within families of subgroup comparisons will be controlled using the Benjamini–Hochberg false-discovery-rate procedure, while primary hypotheses will remain clearly distinguished from exploratory analyses.

For verification performance and responsible behaviour, hierarchical regression models will enter contextual variables first (grade, school type, location, language, device access, formal instruction), followed by usage and literacy variables. Generalized or ordinal models will be used if outcome distributions require them. Because students are nested within schools, cluster-robust standard errors or multilevel models will be used when the number and distribution of schools permit; with too few clusters, the limitation will be stated and school fixed effects or conservative descriptive comparisons considered. Complex mediation or moderation will not be attempted if the achieved sample is inadequate.

Missing data will be described by item and subgroup. If missingness is minimal, complete-case analysis may be acceptable for a given model. If it is substantial and plausibly compatible with missing at random, multiple imputation may be used for covariates and scale items, with the imputation model documented. Outcomes will not be imputed without strong justification. Sensitivity analyses will compare complete-case and imputed estimates. Statistical software, package versions, code, scoring rules, and deviations from the plan will be archived with the project where permissions allow.

Optional open-ended responses will undergo a concise qualitative content analysis. An initial deductive frame will include benefits, accuracy/verification, privacy, bias/fairness, academic integrity, overreliance, access, and requested school support. New categories may be added inductively. A second coder should independently code a subset if feasible; disagreements will be discussed and the codebook revised. Counts will describe code prevalence, but short student comments will not be treated as representative narratives. Quotations, if approved, will be de-identified and used sparingly.

Table 3. Prespecified analysis matrix

Question	Variables	Primary analysis	Safeguard
----------	-----------	------------------	-----------

Question	Variables	Primary analysis	Safeguard
RQ1	Awareness, access, use, purposes, instruction	Frequencies, cross-tabulations, confidence intervals	None; descriptive
RQ2	Dimension items and validated scores	Item distributions; scale summaries; psychometrics	Scale construction precedes comparison
RQ3 / H1 / H4	Understanding, competence, use, verification	Correlations; hierarchical regression; scenario analysis	Grade, access, school context, instruction
RQ4 / H2	Risk awareness, trust, responsibility	Correlations; linear/ordinal regression; non-linear trust tests	Context and usage covariates
RQ5 / H3	Subgroup variables and literacy outcomes	Chi-square; t/ANOVA or nonparametric tests; adjusted models	FDR correction; clustering
RQ6 / H5	Predictors and exploratory pathway	Multilevel/cluster-robust models; indirect effect only if supported	No causal language
RQ7	Optional open text	Content coding and cautious frequency summaries	De-identification; coder audit

Note. Exact model forms will be frozen after pilot revision but before the main outcomes are inspected.

6. Results

EMPIRICAL RESULTS PENDING

Primary data have not been collected. This section contains reporting shells only. No participant count, percentage, mean, correlation, regression coefficient, p value, confidence interval, quotation, or data-based graph should be inserted until it is calculated from genuine, documented responses using the prespecified procedures.

6.1 Participant flow and sample characteristics

Future report: identify the school-recruitment frame, schools approached and participating, eligible students, guardian permissions, assents, returned surveys, exclusions with reasons, and final analytic samples. Describe grade, age band, gender response, school type, location category, primary survey language, device access, and formal AI instruction. Suppress small cells that create disclosure risk.

Table 4. Participant flow and analytic sample—reporting shell

Flow element	n	Reporting note
Schools approached	—	Pending documented recruitment log
Schools participating	—	Pending school permission
Eligible students invited	—	Pending class lists/approved counts
Guardian permissions and student assents	—	Pending consent process
Surveys submitted	—	Pending data collection
Surveys excluded	—	Report each prespecified reason
Final analytic sample	—	Pending cleaning and eligibility checks

Note. Em dash indicates that no genuine value is currently available; it is not zero.

6.2 Descriptive findings

Future report: present awareness, access, ever-use, frequency, purposes, and formal instruction, followed by item-level and validated dimension summaries. Report non-use as a meaningful category. Use denominators that reflect skip patterns and missing responses. Where percentages are compared, provide confidence intervals and raw counts.

Table 5. Core descriptive indicators—reporting shell

Indicator	Planned statistic	Estimate	95% CI
Aware of at least one GenAI tool	n/N (%)	—	—
Ever used a GenAI tool	n/N (%)	—	—
Uses GenAI at least weekly	n/N (%)	—	—
Functional understanding	Mean (SD) or median [IQR]	—	—
Applied competence	Mean (SD) or median [IQR]	—	—
Verification performance	Mean (SD) or median [IQR]	—	—
Risk and ethical awareness	Item profile / justified score	—	—
Responsible behaviour	Mean (SD) or median [IQR]	—	—

Note. Scale rows may be retained only after the planned measurement evaluation. All estimates are pending genuine data.

6.3 Psychometric and hypothesis-testing results

Future report: document item retention, factor structure, omega and alpha where appropriate, objective-item performance, and any deviations from the scoring plan. Then report each hypothesis with effect estimates, uncertainty, model diagnostics, adjusted and unadjusted results, and a clear distinction between prespecified and exploratory analysis. Avoid interpreting statistical significance as educational importance.

Participant-data figures reserved for later insertion:

- Graph A: distribution of validated literacy-dimension scores, with sample size and uncertainty.
- Graph B: verified claim-checking frequency by usage-frequency group, using observed denominators.
- Graph C: adjusted associations with verification performance, reported as coefficients or marginal effects with 95% confidence intervals.
- Graph D: subgroup comparisons only where cell sizes, measurement comparability, and disclosure safeguards are adequate.

Graphing rule

Instrument-design charts may report transparent item counts. Participant-result charts must not be created until every empirical mark is traceable to a cleaned dataset and reproducible analysis script.

7. Discussion Framework

Because no primary data are available, this section does not offer findings. It specifies how future results should be interpreted. The first principle is to compare dimensions rather than label students globally as literate or illiterate. A profile of frequent use, high perceived competence, low objective understanding, and inconsistent verification would indicate an interface-fluency gap. A profile of low use but strong scenario judgement could indicate limited access or deliberate caution rather than weak literacy. Interpretation should therefore begin with patterns across measures and their uncertainty.

The second principle is to distinguish association from cause. If formal instruction is associated with stronger verification, the result would be consistent with educational benefit, but students who choose AI courses or attend participating schools may differ in other ways. If frequent users show stronger applied competence, practice may contribute, but confident students may also choose to use GenAI more often. Cross-sectional mediation would be treated as a model of covariance, not proof of a developmental mechanism. Longitudinal or experimental research would be required for causal claims.

The third principle is calibration. A high trust score should not automatically be celebrated, and a low score should not automatically be interpreted as critical thinking. The most defensible evidence would show that students increase scrutiny when claims are high stakes, unfamiliar, or unsupported, while retaining efficient use for low-risk tasks. Scenario responses and verification behaviour should therefore receive greater interpretive weight than a single general trust statement.

The fourth principle is contextual fairness. Differences by school type, location, language, or device access should not be framed as inherent deficits. They may reflect infrastructure, curricular exposure, adult guidance, institutional rules, or the language and examples used in the instrument. Measurement limitations must be considered before subgroup ranking. Small differences should be reported with effect sizes and confidence intervals, and labels that stigmatize schools or communities should be avoided.

The fifth principle is triangulation with prior scholarship. Results should be compared with general AI literacy frameworks, objective-measurement research, K–12 studies, and Indian policy documents. Agreement with previous research would strengthen plausibility but would not erase local sampling limitations. Disagreement could reflect age, curriculum, language, timing, tool change, or measurement. The literature search should be updated near submission because GenAI research and school policy change rapidly.

Finally, future discussion should address negative, null, and mixed results. If an expected association is not observed, the explanation should not be invented after the fact. Plausible alternatives—restricted range, insufficient power, measurement error, confounding, or genuinely weak association—should be evaluated against diagnostics and reported as uncertainty. The paper's contribution is credible evidence, not confirmation of every hypothesis.

8. Educational and Policy Implications

The framework implies that school AI education should move beyond one-time warnings and tool demonstrations. A compact curriculum can teach a repeatable workflow: define the learning goal; decide whether GenAI is appropriate; write and refine a prompt; identify claims requiring verification; check authoritative and independent sources; protect personal data; and disclose assistance according to school rules. This workflow connects functional understanding to actual student behaviour and can be practised across subjects.

Verification instruction should be explicit. Students can be given short generated answers containing a mixture of accurate, misleading, and unsupported claims and asked to mark what they would check, where they would look, and what evidence would change their judgement. Teachers can model lateral reading, source opening, date checking, and comparison with textbooks or primary documents. Assessment should reward the verification trail and revision decisions, not only the final polished response.

School policies should separate permitted assistance, required disclosure, restricted uses, and prohibited substitution in age-appropriate language. Blanket bans may drive use outside guided settings, while unrestricted use can weaken authorship and privacy norms. Clear examples are more useful than vague statements: brainstorming questions may be allowed; submitting generated prose as original work may not be; personal data should not be entered; and generated citations must be opened and checked. Policies should also specify how students can ask for clarification without being presumed dishonest.

Teacher capacity and infrastructure are essential. Students cannot be expected to evaluate AI in a vacuum when educators lack time, examples, or approved tools. Professional learning should include system limitations, bias, privacy, assessment design, and verification activities. Schools also need equitable alternatives for students without devices or accounts. If GenAI-supported assignments are used, access should be provided within supervised, privacy-aware arrangements or the assignment should be achievable without the tool.

At policy level, the study can inform curriculum implementation by identifying which competencies are weak within participating contexts, but it cannot rank states or represent India nationally. Findings should be used to generate locally testable interventions, such as a short verification module or a privacy-and-disclosure lesson, followed by pre–post evaluation. Coordination among boards, schools, teachers, families, researchers, and child-safety authorities is necessary because literacy, governance, and access are inseparable.

9. Limitations of the Proposed Study

- The purposive participating-school sample will limit population inference and may overrepresent schools already interested in technology or research.
- Cross-sectional data cannot establish causal direction among instruction, use, understanding, verification, trust, and responsibility.
- Self-reported behaviours may be influenced by recall and social desirability, although objective and scenario items reduce reliance on self-report alone.
- An original instrument requires validation; scores cannot be treated as established measures until content, response-process, structural, reliability, and relationship evidence are evaluated.

- Language translation and differing school terminology may affect item interpretation. Measurement equivalence must not be assumed.
- Rapid changes in tools, interfaces, policies, and public awareness may shorten the useful life of brand-specific examples and comparison benchmarks.
- School clustering and a modest number of sites may reduce statistical precision and constrain multilevel or subgroup analysis.
- Students without access can answer knowledge and scenario items but cannot report experienced use, creating intentional skip-pattern differences that require careful denominators.
- Optional open-ended responses can enrich interpretation but will not provide the depth of interviews or ethnographic observation.

10. Future Research

A logical next stage is a pre–post evaluation of a short school-based GenAI literacy intervention built around verification, privacy, and responsible authorship. Randomized or carefully matched designs could test whether instruction improves scenario performance and actual source-checking rather than confidence alone. Delayed follow-up would show whether routines persist. Qualitative interviews and classroom observation could examine why students accept or reject output and how teachers shape norms.

Broader research should include more states, school boards, languages, and students with different accessibility needs. Probability-based or carefully weighted sampling would be needed for stronger population estimates. Longitudinal studies could examine how literacy changes with exposure and curriculum. Tool-comparison studies should focus on task and interface features rather than treating one commercial product as stable. Researchers should also develop age-appropriate performance tasks that assess verification using trace data while protecting privacy.

Future measurement research should test whether the five proposed dimensions form distinct but related scales and whether item functioning is comparable across language, class, and school context. Trust calibration deserves dedicated study because it may be non-linear and task dependent. Responsible behaviour should also be examined through authentic school assignments and disclosure practices, with safeguards that do not turn research into surveillance.

11. Conclusion

GenAI literacy among secondary-school students cannot be inferred from awareness, access, or frequent use. It requires age-appropriate understanding of how generative systems behave, competence in using them for legitimate purposes, active verification of important claims, recognition of privacy and ethical risks, calibrated trust, and responsible conduct. Indian policy and curricular developments make these competencies increasingly relevant, but they do not substitute for empirical evidence about students' actual knowledge and practices.

This paper provides a complete and ethically bounded protocol for producing that evidence among Classes IX and X students. Its mixed-format questionnaire, conceptual model, validation pathway, sampling strategy, safeguards for minors, and prespecified analysis plan are designed to make future findings interpretable and reproducible. The boundary is equally important: no participant characteristics, statistics, results, quotations, or empirical graphs are reported before genuine data collection. The next valid step is not to write a persuasive Results section; it is to obtain appropriate adult and institutional oversight, pilot the instrument, recruit transparently, and analyse real responses according to the documented plan.

References

- Akgun, S., & Greenhow, C. (2022). Artificial intelligence in education: Addressing ethical challenges in K–12 settings. *AI and Ethics*, 2(3), 431–440.
<https://doi.org/10.1007/s43681-021-00096-7>
- Almatrafi, O., Johri, A., & Lee, H. (2024). A systematic review of AI literacy conceptualization, constructs, and implementation and assessment efforts (2019–2023). *Computers and Education Open*, 6, 100173.
<https://doi.org/10.1016/j.caeo.2024.100173>
- Annapureddy, R., Fornaroli, A., & Gatica-Perez, D. (2025). Generative AI literacy: Twelve defining competencies. *Digital Government: Research and Practice*, 6(1), Article 13, 1–21.
<https://doi.org/10.1145/3685680>
- ASER Centre. (2024). Annual Status of Education Report 2023: Beyond Basics.
<https://asercentre.org/aser-2023-beyond-basics/>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 610–623). Association for Computing Machinery.
<https://doi.org/10.1145/3442188.3445922>
- Buçinca, Z., Malaya, M. B., & Gajos, K. Z. (2021). To trust or to think: Cognitive forcing functions can reduce overreliance on AI in AI-assisted decision-making. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW1), Article 188, 1–21.
<https://doi.org/10.1145/3449287>
- Carlini, N., Tramèr, F., Wallace, E., Jagielski, M., Herbert-Voss, A., Lee, K., Roberts, A., Brown, T., Song, D., Erlingsson, Ú., Oprea, A., & Raffel, C. (2021). Extracting training data from large language models. In *30th USENIX Security Symposium (USENIX Security 21)* (pp. 2633–2650). USENIX Association.
<https://www.usenix.org/conference/usenixsecurity21/presentation/carlini-extracting>
- Central Board of Secondary Education. (2026a). Artificial intelligence (417): Class IX curriculum, session 2026–2027.
https://cbseacademic.nic.in/web_material/Curriculum27/sec/417-AI-IX.pdf
- Central Board of Secondary Education. (2026b). Artificial intelligence (417): Class X curriculum, session 2026–2027.
https://cbseacademic.nic.in/web_material/Curriculum27/Sec/417-AI-X.pdf
- Chan, C. K. Y., & Hu, W. (2023). Students’ voices on generative AI: Perceptions, benefits, and challenges in higher education. *International Journal of Educational Technology in Higher Education*, 20, Article 43.
<https://doi.org/10.1186/s41239-023-00411-8>
- Chan, C. K. Y., & Lee, K. K. W. (2023). The AI generation gap: Are Gen Z students more interested in adopting generative AI such as ChatGPT in teaching and learning than their Gen X and millennial generation teachers? *Smart Learning Environments*, 10, Article 60.
<https://doi.org/10.1186/s40561-023-00269-3>
- Chiu, T. K. F., Ahmad, Z., Ismailov, M., & Sanusi, I. T. (2024a). 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.caeo.2024.100171>
- Chiu, T. K. F., Chen, Y., Yau, K. W., Chai, C.-S., Meng, H., King, I., Wong, S., & Yam, Y. (2024b). Developing and validating measures for AI literacy tests: From self-reported to objective measures. *Computers and Education: Artificial Intelligence*, 7, 100282.
<https://doi.org/10.1016/j.caeai.2024.100282>
- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239.
<https://doi.org/10.1080/14703297.2023.2190148>
- Dalalah, D., & Dalalah, O. M. A. (2023). The false positives and false negatives of generative AI detection tools in education and academic research: The case of ChatGPT. *The International Journal of Management*

- Education, 21(2), 100822.
<https://doi.org/10.1016/j.ijme.2023.100822>
- Farrokhnia, M., Banihashem, S. K., Noroozi, O., & Wals, A. (2024). A SWOT analysis of ChatGPT: Implications for educational practice and research. *Innovations in Education and Teaching International*, 61(3), 460–474.
<https://doi.org/10.1080/14703297.2023.2195846>
- Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y. J., Madotto, A., & Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), Article 248, 1–38.
<https://doi.org/10.1145/3571730>
- Klar, M. (2025). Using ChatGPT is easy, using it effectively is tough? A mixed methods study on K–12 students' perceptions, interaction patterns, and support for learning with generative AI chatbots. *Smart Learning Environments*, 12, Article 32.
<https://doi.org/10.1186/s40561-025-00385-2>
- Laupichler, M. C., Aster, A., Haverkamp, N., & Raupach, T. (2023). Development of the 'Scale for the assessment of non-experts' AI literacy'—An exploratory factor analysis. *Computers in Human Behavior Reports*, 12, 100338.
<https://doi.org/10.1016/j.chbr.2023.100338>
- Laupichler, M. C., Aster, A., Schirch, J., & Raupach, T. (2022). Artificial intelligence literacy in higher and adult education: A scoping literature review. *Computers and Education: Artificial Intelligence*, 3, 100101.
<https://doi.org/10.1016/j.caeai.2022.100101>
- Lintner, T. (2024). A systematic review of AI literacy scales. *npj Science of Learning*, 9, Article 50.
<https://doi.org/10.1038/s41539-024-00264-4>
- Lo, C. K. (2023). What is the impact of ChatGPT on education? A rapid review of the literature. *Education Sciences*, 13(4), 410.
<https://doi.org/10.3390/educsci13040410>
- 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>
- Ministry of Human Resource Development. (2020). National Education Policy 2020. Government of India.
https://static.pib.gov.in/WriteReadData/userfiles/NEP_Final_English_0.pdf
- National Council of Educational Research and Training. (2023). National Curriculum Framework for School Education 2023.
https://ncert.nic.in/pdf/focus-group/NCF-SE_2023EN.pdf
- 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>
- OECD. (2023). OECD Digital Education Outlook 2023: Towards an effective digital education ecosystem. OECD Publishing.
<https://doi.org/10.1787/c74f03de-en>
- Pinski, M., & Benlian, A. (2024). AI literacy for users—A comprehensive review and future research directions of learning methods, components, and effects. *Computers in Human Behavior: Artificial Humans*, 2(1), 100062.
<https://doi.org/10.1016/j.chbah.2024.100062>
- Spitale, G., Biller-Andorno, N., & Germani, F. (2023). AI model GPT-3 (dis)informs us better than humans. *Science Advances*, 9(26), eadh1850.
<https://doi.org/10.1126/sciadv.adh1850>
- Tlili, A., Shehata, B., Adarkwah, M. A., Bozkurt, A., Hickey, D. T., Huang, R., & Agyemang, B. (2023). What if the devil is my guardian angel: ChatGPT as a case study of using chatbots in education. *Smart Learning Environments*, 10, Article 15.
<https://doi.org/10.1186/s40561-023-00237-x>

- UNESCO. (2022). K–12 AI curricula: A mapping of government-endorsed AI curricula.
<https://unesdoc.unesco.org/ark:/48223/pf0000380602>
- UNESCO. (2023). Guidance for generative AI in education and research.
<https://unesdoc.unesco.org/ark:/48223/pf0000386693>
- UNESCO. (2024). AI competency framework for students.
<https://unesdoc.unesco.org/ark:/48223/pf0000391105>
- UNESCO New Delhi. (2022). State of the Education Report for India 2022: Artificial intelligence in education: Here, there and everywhere.
<https://unesdoc.unesco.org/ark:/48223/pf0000382661>
- UNICEF. (2021). Policy guidance on AI for children (Version 2.0).
<https://www.unicef.org/innocenti/reports/policy-guidance-ai-children>
- Wang, B., Rau, P.-L. P., & Yuan, T. (2023). Measuring user competence in using artificial intelligence: Validity and reliability of artificial intelligence literacy scale. *Behaviour & Information Technology*, 42(9), 1324–1337.
<https://doi.org/10.1080/0144929X.2022.2072768>
- Wineburg, S., & McGrew, S. (2019). Lateral reading and the nature of expertise: Reading less and learning more when evaluating digital information. *Teachers College Record*, 121(11), 1–40.
<https://doi.org/10.1177/016146811912101102>

Appendix A. Draft Student Questionnaire

Pilot version—do not field without approval

This instrument must undergo expert review, cognitive interviews, translation review where applicable, and pilot testing. School permission, the required guardian permission process, and student assent must be in place before administration.

Working title: Student Generative AI Literacy Questionnaire (SGAILQ–IX/X)

Estimated time: 15–20 minutes, to be confirmed in the pilot.

Definition shown to students: In this survey, ‘generative AI’ means a computer tool that can create new text, images, audio, code, or other content in response to instructions. Examples may include AI chatbots, writing assistants, image generators, or AI features inside search and study tools. You do not need to have used such a tool to participate.

A1. Eligibility and voluntary participation

- E1. I am currently in: ☐ Class IX ☐ Class X ☐ Neither [end survey]
- E2. I have read the student information and understand that participation is voluntary: ☐ Yes ☐ No [end survey]
- E3. I agree to take part: ☐ Yes ☐ No [end survey]
- E4. Where guardian permission is required, the school has confirmed that it has been received: ☐ Yes ☐ No / not confirmed [end survey]

A2. Background and access

- B1. Age band: ☐ 13 or younger ☐ 14 ☐ 15 ☐ 16 ☐ 17 or older ☐ Prefer not to say
- B2. Gender: ☐ Girl ☐ Boy ☐ Another term / self-description ☐ Prefer not to say
- B3. School management type [researcher-coded from school record, not asked if identifiable]: ☐ Government ☐ Government-aided ☐ Private unaided ☐ Other / unclear
- B4. School location [researcher-coded]: ☐ Rural ☐ Semi-urban ☐ Urban ☐ Other / unclear
- B5. Language used for this survey: ☐ English ☐ Hindi ☐ Other approved translation: _____
- B6. At home, how often can you use an internet-connected device for schoolwork? ☐ Never ☐ Rarely ☐ Sometimes ☐ Often ☐ Whenever I need it
- B7. Which devices can you usually use for schoolwork? Select all: ☐ Shared smartphone ☐ Own smartphone ☐ Tablet ☐ Laptop/desktop ☐ School device ☐ None ☐ Other
- B8. Have you received a lesson, workshop, or formal instruction about AI or generative AI? ☐ Yes ☐ No ☐ Not sure
- B9. Does your school have a rule or guidance about using generative AI for assignments? ☐ Yes ☐ No ☐ Not sure

A3. Awareness and use

- U1. Before today, had you heard of a generative AI tool? ☐ Yes ☐ No ☐ Not sure
- U2. Have you ever used a generative AI tool yourself? ☐ Yes ☐ No ☐ Not sure. [If No/Not sure, skip to A4.]
- U3. During the past 30 days, how often have you used generative AI? ☐ Not at all ☐ Less than weekly ☐ 1–2 days per week ☐ 3–5 days per week ☐ Almost every day

- U4. Where have you used it? Select all: ☐ At home ☐ At school ☐ Tuition/coaching ☐ Library/community setting ☐ Somewhere else
- U5. What have you used it for? Select all: ☐ Explain a concept ☐ Summarize text ☐ Brainstorm ideas ☐ Draft or improve writing ☐ Solve maths/science problems ☐ Translate ☐ Write/debug code ☐ Create images/media ☐ Prepare for tests ☐ Personal curiosity/entertainment ☐ Other
- U6. Who has guided you in using it? Select all: ☐ Teacher ☐ Parent/guardian ☐ Sibling/relative ☐ Friend/classmate ☐ Online video/site ☐ The tool itself ☐ No one ☐ Other
- U7. When using GenAI for schoolwork, how often do you tell the teacher or reader that you used it when rules require disclosure? ☐ Never ☐ Rarely ☐ Sometimes ☐ Often ☐ Always ☐ Not applicable / do not know the rule

A4. Functional AI understanding

Choose the best answer. 'Not sure' is a valid response.

- K1. Which statement best describes how a text-generating AI usually creates an answer? ☐ It searches a perfect database of facts ☐ It predicts and generates language patterns from training and instructions ☐ A human secretly writes every answer ☐ Not sure
- K2. A GenAI answer sounds confident. What does that tell you about accuracy? ☐ It must be correct ☐ It is probably checked by an expert ☐ Confidence of wording does not guarantee accuracy ☐ Not sure
- K3. The same prompt is entered twice and produces somewhat different answers. Is this possible? ☐ Yes ☐ No ☐ Only if the internet stops ☐ Not sure
- K4. Which action can improve an AI response but cannot guarantee it is true? ☐ Giving clearer context and asking for reasoning or sources ☐ Typing only in capital letters ☐ Accepting the first answer immediately ☐ Not sure
- K5. Why might an AI output show unfair bias? ☐ Training data and design choices can contain unequal patterns ☐ Computers cannot reflect human patterns ☐ Bias occurs only when a user misspells a word ☐ Not sure
- K6. Which information should generally not be entered into a public GenAI tool? ☐ A made-up practice question ☐ Personal passwords, private records, or identifying details ☐ A general request for study tips ☐ Not sure
- K7. An AI gives a detailed citation. Which statement is safest? ☐ Detailed citations always exist ☐ The citation may be real or fabricated and should be checked ☐ A long title proves the study is genuine ☐ Not sure
- K8. Who remains responsible for checking and appropriately using an AI-assisted school answer? ☐ The student/user and relevant adults under school rules ☐ The AI alone ☐ Nobody if the answer is fluent ☐ Not sure

A5. Applied AI competence

For each statement choose: 1 Strongly disagree, 2 Disagree, 3 Neither agree nor disagree, 4 Agree, 5 Strongly agree.

- C1. I can state a clear goal and give enough context when asking a GenAI tool for help. 1 2 3 4 5
- C2. I can improve a prompt after seeing an unhelpful answer. 1 2 3 4 5
- C3. I can decide which school tasks are suitable or unsuitable for GenAI assistance. 1 2 3 4 5
- C4. I can compare two AI responses and explain which is more useful. 1 2 3 4 5
- C5. I can use GenAI to support my thinking without letting it replace all of my own work. 1 2 3 4 5

A6. Verification practices

For each statement choose: 1 Never, 2 Rarely, 3 Sometimes, 4 Often, 5 Always. If you have never used GenAI, choose what you would be likely to do and mark the response as hypothetical.

- V1. I check important factual claims in a reliable source outside the AI response. 1 2 3 4 5
- V2. I open a source or citation instead of trusting it because it is listed. 1 2 3 4 5
- V3. I compare an AI answer with a textbook, teacher explanation, or another authoritative source. 1 2 3 4 5
- V4. I check whether information is current enough for the question. 1 2 3 4 5
- V5. I look for signs that the answer may be uncertain, incomplete, or one-sided. 1 2 3 4 5
- V6. I change, reject, or stop using an AI answer when evidence does not support it. 1 2 3 4 5
- V7. Scenario: An AI gives a confident answer to a science question but gives no evidence. What is the best next step? ☐ Submit it because it sounds clear ☐ Check the key claim in the textbook or a reliable scientific source ☐ Ask the AI to make it sound more confident ☐ Remove difficult words ☐ Not sure
- V8. Scenario: An AI lists three research papers for an assignment. What is the best way to check them? ☐ Assume they are real because the titles are detailed ☐ Search each title/author/DOI in a reliable scholarly or publisher source and inspect whether it supports the claim ☐ Count how many references appear ☐ Copy them before they disappear ☐ Not sure

A7. Risk and ethical awareness

Choose: 1 Strongly disagree, 2 Disagree, 3 Neither agree nor disagree, 4 Agree, 5 Strongly agree.

- R1. A GenAI output can contain convincing but false information. 1 2 3 4 5
- R2. AI output can reproduce unfair stereotypes or leave out some groups' experiences. 1 2 3 4 5
- R3. Entering personal or confidential information into a public AI tool can create privacy risk. 1 2 3 4 5
- R4. Depending on GenAI for every step of an assignment can reduce opportunities to practise thinking and writing. 1 2 3 4 5
- R5. School rules about authorship and disclosure still apply when AI helped create part of the work. 1 2 3 4 5
- R6. Images, text, or code produced with AI can raise copyright and attribution questions. 1 2 3 4 5
- R7. The seriousness of checking an AI answer should increase when an error could have greater consequences. 1 2 3 4 5
- R8. Equal access to devices, connectivity, language support, and guidance affects who can benefit from GenAI. 1 2 3 4 5

A8. Responsible behaviour

Choose: 1 Never, 2 Rarely, 3 Sometimes, 4 Often, 5 Always. Non-users may answer what they intend to do and mark responses as hypothetical.

- BHV1. I keep my own judgement and revise AI-assisted work rather than submitting raw output. 1 2 3 4 5
- BHV2. I avoid sharing passwords, private records, contact details, or other sensitive information with public AI tools. 1 2 3 4 5
- BHV3. I follow the teacher's or school's rule for permitted use and disclosure. 1 2 3 4 5
- BHV4. I identify AI assistance in my work when disclosure is required. 1 2 3 4 5

BHV5. I stop and seek help from a trusted adult or teacher when an output is concerning, high stakes, or beyond my ability to check. 1 2 3 4 5

BHV6. I use verified information and remain responsible for the final work I submit. 1 2 3 4 5

A9. Trust calibration

Choose: 1 Strongly disagree, 2 Disagree, 3 Neither agree nor disagree, 4 Agree, 5 Strongly agree.

T1. I would trust a simple AI suggestion more when I can quickly check it against reliable evidence. 1 2 3 4 5

T2. I would use extra caution when an AI answer affects an important decision or is outside my subject knowledge. 1 2 3 4 5

T3. If reliable sources disagree with an AI answer, I would lower my confidence in the answer. 1 2 3 4 5

T4. I can benefit from GenAI without assuming that every response is correct. 1 2 3 4 5

A10. Optional open responses

O1. What is one useful way generative AI could support students' learning?

O2. What is one concern you have about students using generative AI?

O3. What should schools teach or provide so students can use generative AI safely and responsibly?

Thank you. Please submit the survey without writing your name, roll number, email address, phone number, or AI account details.

Appendix B. Participant Information Sheet

Study title: Generative AI Literacy Among Indian Secondary-School Students: Awareness, Usage, Verification and Risk Perception

Researcher: Aryaveer Singh, Founder & CEO, GENESIS Foundation | Independent Student Researcher, working under a named adult mentor and participating-school oversight to be inserted before use.

Why is this study being done?

The study aims to understand what students in Classes IX and X know about generative AI, how they use it, how they check its answers, and what risks and responsible practices they recognize. The study is not a test for grades, and it is useful to hear from students who have never used generative AI.

What will I do?

If you take part, you will complete a survey expected to take about 15–20 minutes. It includes multiple-choice questions, rating scales, short situations, and three optional written questions. You will not be asked to log in to an AI tool or paste any private conversation.

Do I have to take part?

No. Participation is voluntary. You may say no, skip any nonessential question, or stop before submitting. Your decision will not affect your grades, school standing, or access to learning activities. Once an anonymous survey has been submitted and cannot be linked to you, it may not be possible to remove an individual response.

Are there risks or benefits?

The survey presents no more than minimal expected risk, but a question may make you uncertain about your own technology use. You may skip it. There is no guaranteed personal benefit. The study may help schools plan clearer guidance on verification, privacy, and responsible use.

How will privacy be protected?

Do not provide your name, roll number, personal email, phone number, account name, password, or private AI conversation. Responses will be stored in an access-restricted location and reported in groups. Small groups will not be published where they could identify a student. Optional comments will be checked and de-identified before any approved quotation.

Who can answer questions?

Before distribution, insert the adult mentor's name and school-approved contact route, the responsible school contact, and the approved complaint or ethics contact. The student researcher should not be the only contact for a study involving minors.

Fields to complete before use

Adult mentor: _____ | School contact: _____ | Ethics/safeguarding contact: _____ |
Approved data-retention date: _____ | Survey languages: _____

Appendix C. Guardian Permission Form

Please read the Participant Information Sheet before signing. Return this form through the school-approved process. Keep a copy if requested by the school.

- ☐ I have read and understood the information about the study.
- ☐ I understand that my child's participation is voluntary and will not affect grades or school standing.
- ☐ I understand that the survey does not ask for the student's name, roll number, personal account details, or private AI conversations.
- ☐ I understand that my child may skip nonessential questions or stop before submitting.
- ☐ I understand how responses will be stored, reported, and retained, once the final approved dates and contacts are inserted.
- ☐ I have had an opportunity to ask questions through the school-approved contact route.
- ☐ I give permission for my child to be invited to decide whether to participate. I understand that my child must also assent.

Student's name or school code (as required by the school; store separately from survey data):

Guardian's name: _____ Signature: _____ Date: _____

Approved contact for questions:

Data-handling note: Signed permission forms must be stored separately from anonymous survey responses and accessed only by authorized adults under the approved procedure.

Appendix D. Student Assent Form

Please tick each box if it is true for you. Ask an adult if anything is unclear.

- ☐ I understand that this survey is about students' knowledge and use of generative AI.
- ☐ I understand what I will be asked to do and that the survey takes about 15–20 minutes.
- ☐ I understand that taking part is my choice and that saying no will not affect my grades or school standing.
- ☐ I know I may skip nonessential questions or stop before I submit the survey.
- ☐ I know I should not write my name, roll number, personal email, phone number, password, account name, or private AI conversation in the survey.
- ☐ I have had a chance to ask questions.
- ☐ I agree to take part.

Student signature or approved assent mark: _____ Date: _____

Researcher/administrator confirmation: I explained the study in age-appropriate language and observed a voluntary assent decision.

Name: _____ Signature: _____ Date: _____

Appendix E. Variable Coding Framework

This codebook is preliminary. Version-control every change after expert review and pilot testing. Freeze the final codebook before primary outcome analysis. Use -99 for system missing only in raw exports if required by software; retain labelled missing categories in analysis files rather than mixing them with valid zeros.

Table E1. Preliminary variable codebook

Variable	Source	Label	Coding	Use
class	B/E1	Grade	1=IX; 2=X	Eligibility; covariate
age_band	B1	Age band	1= ≤ 13 ; 2=14; 3=15; 4=16; 5= ≥ 17 ; 9=prefer not	Describe; covariate if justified
gender	B2	Gender response	1=girl; 2=boy; 3=self-description; 9=prefer not	Describe; combine/suppress small cells
school_type	B3	School management	1=government; 2=aided; 3=private; 4=other	Cluster/context
location	B4	School location	1=rural; 2=semi-urban; 3=urban; 4=other	Context
survey_lang	B5	Survey language	1=English; 2=Hindi; 3+=approved language	Measurement/context
device_access	B6	Home device access	1=never ... 5=whenever needed	Ordinal access indicator
formal_ai	B8	Formal AI instruction	0=no; 1=yes; 8=not sure	H3 predictor
school_rule	B9	Aware of school AI rule	0=no; 1=yes; 8=not sure	Context
aware_genai	U1	Prior awareness	0=no; 1=yes; 8=not sure	RQ1
ever_use	U2	Ever used GenAI	0=no; 1=yes; 8=not sure	RQ1; skip logic
use_freq	U3	Use in past 30 days	0=none; 1= $\leq$ weekly; 2=1–2d/wk; 3=3–5d/wk; 4=almost daily	RQ1/H4
purpose_*	U5	Purpose selections	0=not selected; 1=selected; NA=skipped non-user	Multiple binary items
knowledge_*	K1–K8	Objective items	1=keyed answer; 0=other/not sure	Item-level; provisional sum/mean
knowledge_score	Derived	Functional understanding	Mean or sum if $\geq 80\%$ complete and coherent	H1/H3
competence_*	C1–C5	Applied competence items	1=strongly disagree ... 5=strongly agree	Candidate reflective scale
competence_score	Derived	Applied competence	Mean if $\geq 80\%$ complete and validated	RQ2/H4
verify_*	V1–V6	Verification frequency	1=never ... 5=always; plus hypothetical flag	Candidate behaviour scale/profile
scenario_*	V7–V8	Verification scenarios	1=keyed action; 0=other/not sure	Performance outcome
verify_score	Derived	Verification performance	Prespecified combination only if coherent	Primary outcome
risk_*	R1–R8	Risk/ethical items	1=strongly disagree ... 5=strongly agree	Profile; scale only if supported
responsible_*	BHV1–BHV6	Responsible behaviours	1=never ... 5=always; plus hypothetical flag	Candidate reflective scale

Variable	Source	Label	Coding	Use
responsible_score	Derived	Responsible behaviour	Mean if $\geq 80\%$ complete and validated	Primary outcome
trust_*	T1–T4	Task-sensitive trust	1=strongly disagree ... 5=strongly agree	Profile/non-linear analysis
open_*	O1–O3	Optional text	De-identified text plus multi-label codes	RQ7 content analysis
school_id	Admin	Anonymous school code	Random/project code; no school name in analysis file	Clustering
pilot_flag	Admin	Pilot vs main	0=main; 1=pilot	Exclude pilot from main analysis
quality_flag	Derived	Data-quality review	0=none; 1=review; 2=exclude under prespecified rule	Sensitivity analysis

Note. Asterisks indicate a family of separately stored variables. Do not collapse multiple-response selections into one string if they will be analysed quantitatively.

E1. Proposed scoring keys

Objective knowledge keyed answers: K1 predicts/generates language patterns; K2 confident wording does not guarantee accuracy; K3 yes; K4 clearer context and requests for reasoning/sources can improve but not guarantee truth; K5 training data and design choices can contain unequal patterns; K6 personal passwords/private records/identifying details; K7 citations may be real or fabricated and should be checked; K8 the student/user and relevant adults remain responsible under school rules.

Verification scenarios: V7 check the key claim in a textbook or reliable scientific source; V8 search and inspect each citation in a reliable scholarly or publisher source. Distractors and ‘not sure’ score 0 for the provisional performance score but remain visible in item-level reporting.

No reverse-scored items are currently proposed. If pilot revisions add negatively worded items, the codebook must identify them explicitly. Hypothetical responses from non-users should be flagged and analysed separately from experienced behaviour in sensitivity analyses.

E2. Qualitative coding starter framework

Table E2. Preliminary open-response coding framework

Code	Include when response mentions
LEARN_SUPPORT	Explanation, practice, feedback, examples, brainstorming, translation
EFFICIENCY	Saving time or organizing work without necessarily improving learning
ACCURACY	Wrong, fabricated, incomplete, outdated, or misleading output
VERIFY	Checking sources, citations, dates, textbooks, teachers, or independent evidence
PRIVACY	Personal data, accounts, tracking, confidentiality
BIAS_FAIR	Stereotypes, unequal representation, unfairness, access inequality
INTEGRITY	Authorship, copying, disclosure, school rules, assessment fairness
OVERRELIANCE	Reduced effort, practice, judgement, creativity, or dependency
ACCESS	Devices, connectivity, cost, language, disability access, adult support

Code	Include when response mentions
SCHOOL_SUPPORT	Lessons, teacher guidance, approved tools, examples, policies, help routes
OTHER	Relevant content not covered; memo and consider new code

Note. Multiple codes may be applied. Add exclusions and examples after pilot responses; do not invent example quotations.