

Shaping Human-Driven Validation in AI-Integrated Design Education

Jinoh Park

Department of Interior Architecture and Design, Fay Jones School of Architecture and Design, University of Arkansas, Fayetteville, United States

Soo Jeong Jo

School of Architecture, Louisiana State University, Baton Rouge, LA, United States

Accepted Manuscript

Shaping Human-Driven Validation in AI-Integrated Design Education

Generative artificial intelligence (AI) is increasingly used in design education to support early-stage tasks such as programming, ideation, and environmental analyses. Yet its educational value depends on how students learn to question, verify, revise, and justify AI-supported outputs rather than uncritical adoption. The present study compares two pedagogical cases: a text-to-program workflow in a Council for Interior Design Accreditation (CIDA)-accredited interior architecture course and a performance-based mass study workflow in a National Architectural Accrediting Board (NAAB)-accredited architecture course. Using Linguistic Informatics and Geometric Synthesis (LIGS) as a cross-modal analytical framework, this study analyzes 1,081 structured reflection responses from 61 students, 9 AI-generated draft programs, 9 revised adjacency matrices, and environmental performance outputs from 20 design options. The findings show that AI became pedagogically meaningful when students treated generated material as provisional input for judgment. In Case A, students used a Keep, Revise, Reject pattern to evaluate AI-generated programmatic suggestions against user needs, adjacency logic, code-related concerns, and spatial feasibility. In Case B, students used environmental feedback from the simulations to compare massing alternatives and interpret the trade-offs among different criteria in design and environmental performance. These findings were understood as illustrative rather than representative. Across both cases, AI-supported learning depended on human-driven validation, trust calibration, metacognitive awareness, and decision ownership. The present study offers exploratory comparative evidence for a human-centered and a performance-based approach to AI-integrated design education.

Keywords: human-centered AI, design education, trust calibration, AI literacy, environmental performance

Accepted Manuscript

Introduction

Generative artificial intelligence (AI) is reshaping design education, particularly in early-stage design tasks such as ideation, programming, and environmental analyses. In architecture and interior design programs, AI-supported tools can help students generate alternative scenarios, organize project information, compare options, and visualize performance-related feedback. Yet, the educational value of these tools depends not only on efficiency, but also on how students are guided to interpret, verify, and revise AI-supported outputs. As design programs respond to these changes, educators must integrate AI while maintaining the professional and pedagogical expectations of accredited design education, including attention to human experience, environmental performance, responsibility, context, and judgment (Buolamwini, 2024; CIDA, 2024; NAAB, 2025).

A central challenge in AI-integrated design education is not whether students use generative AI, but whether they engage with it critically. Fluent text, plausible programmatic suggestions, and optimized design options can appear authoritative, encouraging students to treat AI outputs as shortcuts rather than as materials for inquiry. This risk is pedagogically important because effortful learning depends on comparison, correction, and justification. When students accept AI-generated outputs without examining their assumptions or limitations, they may lose opportunities to develop disciplinary understanding, monitor their own knowledge, and identify what still requires verification. Recent research suggests that metacognitive support is necessary to sustain self-regulated learning, while student agency must be protected when AI systems scaffold or automate learning tasks (Darvishi et al., 2024; Xu et al., 2025).

This issue becomes more complex in design workflows because AI support often operates across different modes of reasoning. Language-based tools may assist with client interpretation,

persona development, and programmatic drafting, whereas simulation-based tools may support geometric exploration and environmental performance optimization. However, neither type of output automatically produces coherent design judgment. Students must evaluate whether AI-generated information is relevant, incomplete, biased, technically inaccurate, or inconsistent with the project context. AI literacy therefore involves more than effective prompting. It requires students to calibrate trust, recognize the limits of AI-generated outputs, and adapt them to disciplinary and project-specific criteria (Ehsan et al., 2024; Knoth et al., 2024; Liao et al., 2020).

Against this background, this study questions: How do students in accredited design education validate, revise, and justify AI-supported outputs across linguistic and geometric design workflows? A secondary research question asks: How can the Linguistic Informatics and Geometric Synthesis (LIGS) framework help compare the forms of judgment, metacognitive awareness, and decision ownership that emerge across these workflows?

Responding to these research questions, the present study examines two pedagogical cases: a text-to-program workflow in a CIDA-accredited interior architecture program (Case A) and a data-informed massing workflow in a NAAB-accredited architecture program (Case B). Rather than evaluating the technical accuracy of AI systems in general, this study examines how AI-supported workflows structure student judgment, trust calibration, and critical curation in early-stage design learning. Drawing on student reflections, AI-generated draft programs, revised adjacency matrices, and environmental performance outputs, this study uses the LIGS framework to compare how students validate, revise, and justify AI-supported decisions across disciplinary contexts. Figure 1 visualizes this framework by connecting linguistic interpretation, geometric evaluation, and human validation.

2. Literature Review and Theoretical Framework

This study draws on four intersecting areas of scholarship: accredited design education, AI-supported learning, cross-modal design reasoning, and human-centered validation. Together, these areas frame the central concern of the study: how students use AI-supported tools while retaining responsibility for design judgment. The review first examines how interior architecture and architecture accreditation define professional learning expectations. It then discusses student agency, metacognitive support, and AI literacy in generative AI learning environments. Building on these discussions, the section introduces Linguistic Informatics and Geometric Synthesis (LIGS) as an analytical framework for comparing language-based and geometry-based AI workflows. Finally, it connects trust calibration, explainability, and ethical stewardship to the study's focus on student validation.

2.1. Accredited Design Education and AI-Supported Learning

Accredited design education defines what students are expected to know, do, and justify as they prepare for professional practice. In interior architecture and interior design education, CIDA emphasizes human experience, user needs, interior environments, communication, professional practice, and context-sensitive decision-making (CIDA, 2024). In architecture education, NAAB emphasizes design integration, technical knowledge, environmental systems, research, equity, ethics, and professional responsibility (NAAB, 2025). Although these accreditation systems differ in disciplinary focus, both require students to demonstrate informed judgment rather than merely produce design outcomes.

This expectation becomes more complex as generative AI enters early-stage design learning. AI-supported tools can help students generate scenarios, develop programmatic suggestions, explore

spatial options, and compare performance-related feedback. Recent scholarship suggests that AI can support both divergent and convergent phases of design learning, especially during ideation, exploration, and early evaluation (Melker et al., 2025). If used appropriately, these tools can expand the range of possibilities that students may explore and help them refine emerging design directions.

However, AI support does not replace the reasoning that accredited design education seeks to cultivate. Research on AI assistance suggests that automated scaffolding can support learning tasks but may also weaken student agency when learners become dependent on system-generated guidance (Darvishi et al., 2024). Research on generative AI learning environments similarly emphasizes metacognitive support, because students must monitor what they understand, what remains uncertain, and what requires verification (Xu et al., 2025). This is especially important in design education, where learning develops through interpretation, comparison, revision, and justification rather than through a single correct answer.

AI literacy is therefore not limited to prompt writing or tool operation. It includes the ability to evaluate whether AI-generated outputs are relevant, reliable, incomplete, biased, or misaligned with project-specific criteria (Knoth et al., 2024). In this study, AI-supported learning is understood as a structured process of inquiry and validation, not as the automation of design work.

2.2. Cross-modal Design Reasoning and LIGS

Design students increasingly move across multiple modes of reasoning. Early design stages (Cross, 2001). Computational design scholarship further shows how digital media reshape architectural thinking by linking information, representation, and design action (Kalay, 2004). Performance-based design extends this discussion by examining how simulation and evaluation can inform form

generation and design decision-making (Oxman, 2008). This study proposes Linguistic Informatics and Geometric Synthesis (LIGS) as an analytical framework for comparing how students validate AI-supported outputs across language-based interpretation and geometry-based evaluation.

Linguistic informatics refers to the use of language-based AI tools to interpret, organize, and transform textual information into design-relevant criteria. In Case A, large language models supported client interpretation, persona development, program drafting, adjacency reasoning, and spatial prioritization. The educational value of this workflow did not lie in the generated text itself, but in how students compared that text with project requirements, code expectations, user needs, and their own design reasoning. Geometric synthesis refers to the use of computational and AI-supported tools to compare spatial and formal alternatives in relation to measurable performance criteria. In Case B, students examined massing options through environmental performance feedback during early-stage architectural design. These tools helped students connect geometric decisions with environmental implications, but they did not determine the final design decision. Students still had to judge whether a performance result was meaningful, aligned with design intent, and sufficient to justify revision (Verganti et al., 2020).

LIGS is therefore positioned as an analytical lens rather than as a pre-existing theory. It synthesizes prior work on *designerly* ways of knowing, computational design reasoning, performance-based architectural evaluation, and AI-supported design learning (Cross, 2001; Kalay, 2004; Melker et al., 2025; Oxman, 2008) rather than accept them as final solutions (Tellez, 2025). Its purpose is to compare how students move between linguistic interpretation and geometric evaluation while retaining responsibility for design decision-making.

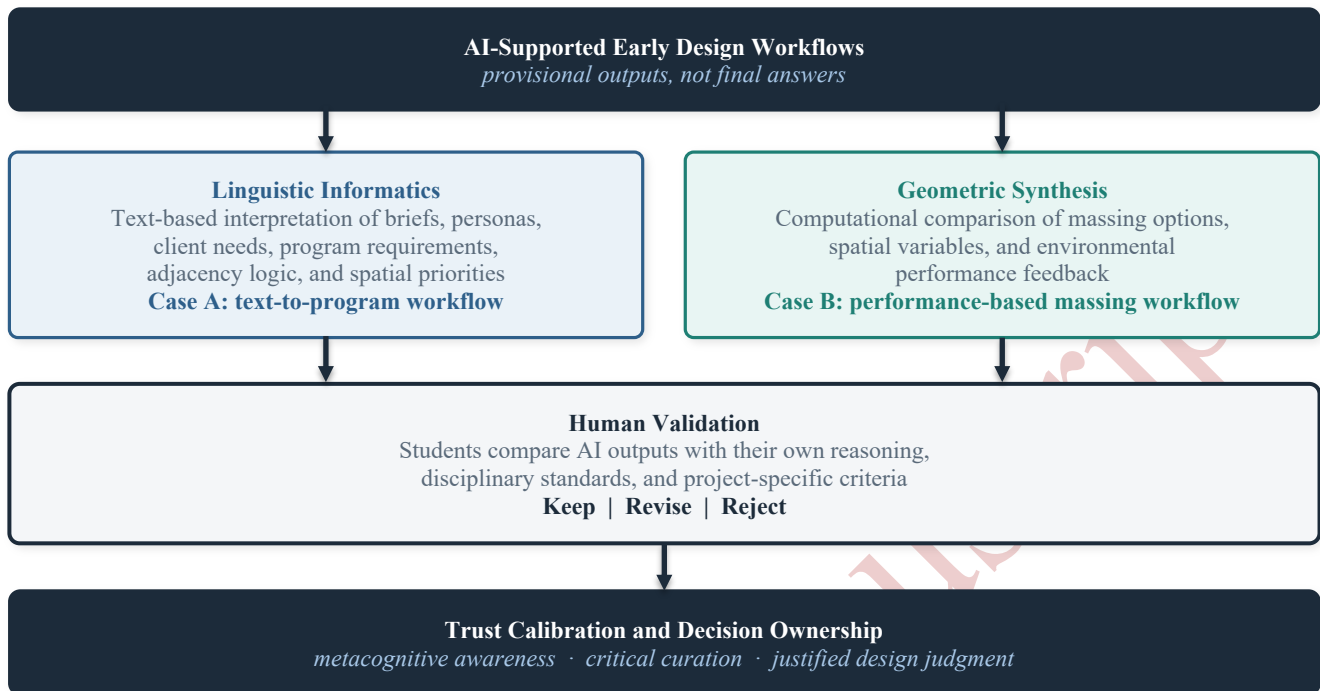


Figure 1. LIGS as a cross-modal framework for AI-integrated design education

2.3. HCAI, Trust Calibration, and Critical Validation

Human-Centered AI (HCAI) emphasize that AI systems should support, rather than displace, human agency and judgment in professional and educational settings (Park, 2026). In design education, this means students must learn not only how to use AI tools, but also how to question their outputs and understand their limitations. Explainable AI research is relevant here because users need explanations that clarify how systems work, why outputs are generated, when outputs can be trusted, and where errors may occur (Liao et al., 2020). For design students, these questions are also pedagogical, because they must determine whether a suggestion is appropriate for a specific user, space, context, and design intention. Trust calibration refers to the ability to align confidence in AI-supported outputs with their actual reliability, limits, and relevance. Overtrust may lead students to accept fluent or visually

convincing results without sufficient review. Undertrust may prevent them from using AI productively for exploration and comparison. Calibrated trust requires students to examine output critically and decide whether a suggestion by AI should be kept, revised, or rejected. Seamful design in Explainable AI (XAI) offers a useful perspective on this process. Rather than hiding system limitations, seamful approaches make gaps, mismatches, uncertainties, and breakdowns visible so that users can respond more critically (Ehsan et al., 2024). In AI-integrated design education, such seams can become learning moments. Inaccurate adjacency logic, incomplete user assumptions, unrealistic program suggestions, or performance outputs that conflict with design intent can prompt students to compare evidence, correct assumptions, and justify decisions.

The “Keep, Revise, Reject” pattern is used in this study as an analytical expression of student validation. “Keep” describes moments when students retain outputs that align with project criteria. “Revise” describes moments when they modify outputs to improve specificity, accuracy, code alignment, spatial logic, environmental performance, or user fit. “Reject” describes moments when they discard outputs that are implausible, incomplete, biased, technically inaccurate, or inconsistent with the project context. This pattern links AI literacy to design judgment by making student decision-making visible.

2.4. Ethical Stewardship and Situated Design Judgment

AI-integrated design education also raises ethical questions. Generative AI systems can reproduce bias, normalize dominant assumptions, or obscure the cultural and social conditions embedded in datasets. Students must therefore evaluate not only whether AI outputs are useful, but also whether they are inclusive, context-sensitive, and ethically appropriate. Recent work on algorithmic harm and AI justice

emphasizes that AI systems can encode social inequities, making human oversight essential in educational and professional contexts (Buolamwini, 2024; Zhang, 2023). In relation to professional expectations concerning global context and human diversity, these concerns also raise questions about cognitive imperialism in data-driven systems (Arora, 2025; CIDA, 2024; Ofosu-Asare, 2025). AI-supported outputs may appear neutral, but they often carry assumptions about users, space, performance, value, and norms. Students need opportunities to examine whose needs are represented, whose perspectives are missing, and how design decisions may affect different users or communities. This study treats ethical stewardship as part of human-driven validation. In Case A, validation involves examining whether AI-generated programmatic suggestions reflect user needs, accessibility, adjacency, and spatial priorities. In Case B, it involves examining whether performance-based outputs support design decisions without reducing architecture to numerical optimization alone. Across both cases, students are expected to treat AI outputs as prompts for reflection rather than substitutes for judgment.

As summarized in Table 1, the literature suggests that the educational value of AI in design depends on structured human validation. Accreditation standards define the disciplinary expectations students must meet. Research on AI-supported design education explains how AI can support divergent and convergent thinking, while studies of agency, metacognition, and AI literacy clarify why students must remain responsible for comparison and judgment. Human-centered explainability and seamless XAI show why AI outputs should remain open to questioning. Ethical stewardship extends validation beyond technical accuracy to include bias, context, and human consequences. LIGS synthesizes these concerns by comparing two AI-supported design workflows. In the linguistic workflow, students validate text-based programmatic suggestions. In the geometric workflow, they validate performance-based spatial alternatives. The two cases are not treated as equivalent datasets, but as comparable

pedagogical situations in which students interpret, revise, and justify AI-supported outputs. This framework supports the study's central research question by focusing attention on trust calibration, metacognitive awareness, and decision ownership across different modes of AI-integrated design learning.

Table 1. Theoretical constructs and their analytical use in this study

Construct	Literature	Analytical use in this study
Accredited design judgment	CIDA, 2024; NAAB, 2025	Defines the disciplinary expectations used to evaluate AI-supported outputs
AI-supported design learning	Melker et al., 2025	Frames AI as support for divergent and convergent design thinking
Student agency	Darvishi et al., 2024	Explains why students must remain active decision-makers
Metacognitive support	Xu et al., 2025	Explains why students must monitor what they know, do not know, and need to verify
AI literacy	Knoth et al., 2024	Supports evaluation of AI outputs beyond prompting or tool operation
LIGS	Proposed in this study, drawing on Cross, 2001; Kalay, 2004; Melker et al., 2025; Oxman, 2008	Compares linguistic and geometric AI-supported workflows
Trust calibration	Ehsan et al., 2024; Liao et al., 2020	Explains how students decide whether to keep, revise, or reject AI outputs
Ethical stewardship	Buolamwini, 2024	Connects validation to bias, inclusion, context, and human responsibility

3. Methodology

This study uses a dual-case comparative design to examine how students validated AI-supported outputs in early-stage design education. The study did not aim to measure the technical performance of AI systems or to compare the two cases through a single quantitative metric. Instead, it examined how students interpreted, questioned, revised, and justified AI-supported outputs within two accredited design learning contexts: pre-design programming in interior architecture and performance-based

massing in architecture.

Table 2. Comparative overview of the two pedagogical cases.

Component	Case A: Linguistic Informatics	Case B: Geometric Synthesis
Disciplinary setting	CIDA-aligned interior architecture course at a U.S. public university	NAAB-aligned fifth-year Bachelor of Architecture course at a U.S. public university
Student level	Third-year interior architecture students, n = 61	Fifth-year architecture students, n = 10
Design phase	Pre-design programming and spatial organization	Early-stage massing and environmental performance studies
AI-supported workflow	LLM-supported client interpretation, persona development, program drafting, adjacency reasoning, and spatial prioritization	Autodesk Forma and Revit-supported comparison of massing options, illuminance, and environmental performance feedback
Data sources	1,081 structured reflection responses, 9 AI-generated draft programs, and 9 revised adjacency matrices	20 design options, Autodesk Forma microclimate simulations, Revit illuminance renderings, and the environmental performance indicators
Unit of analysis	Student validation action: a documented moment when an AI-generated suggestion was kept, revised, or rejected	Iterative design studies and feedback on the design options' environmental performance
Analysis focus	How students interpreted, questioned, modified, or discarded AI-generated programmatic suggestions	How students used environmental feedback to compare, revise, and justify their form-finding processes
Validation strategy	Triangulation across reflections, AI-generated draft programs, and final adjacency matrices	Comparison of performance outputs across design iterations and interpretation in relation to design decisions
Scope of claim	Qualitative evidence of student trust calibration and decision ownership in a linguistic workflow	Illustrative case-specific evidence of performance-based decision-making in a geometric workflow

The two cases were selected because they represent distinct but comparable AI-supported workflows. Case A focused on language-based interpretation during pre-design programming. Case B focused on performance-based design decision-making during the early stages of mass studies. As

summarized in Table 2, the cases differed in discipline, student level, workflow, and data format. They were therefore compared as pedagogical processes rather than as equivalent datasets.

3.1. Research Design: Cyber-Constructive Triangulation (CCT)

The study adopts Cyber-Constructive Triangulation (CCT) as an interpretive framework for examining AI-supported design processes. Rather than functioning as a formal measurement instrument, CCT provides a way to analyze the relationship among humans (student designers), machines (AI systems), and data (the social and digital information structures that inform design decisions) (Chiou et al., 2023; Stoimenova & Price, 2020). In this study, CCT is further informed by seamful design principles, which highlight gaps, uncertainties, and inconsistencies in AI-supported decision-making (Ehsan et al., 2024). Accordingly, the analysis is organized across three levels: (1) the interpretive layer, (2) the algorithmic layer, and (3) the contextual layer as shown in Figure 2.

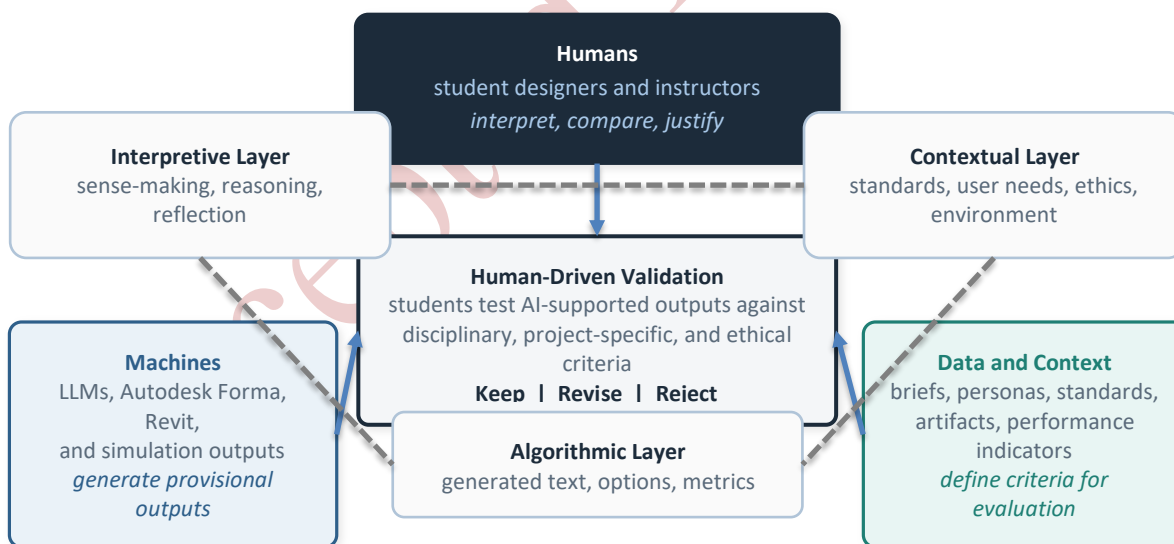


Figure 2. Cyber-Constructive Triangulation (CCT) for AI-supported design learning

3.2. Case Analysis

Each case was analyzed within its own disciplinary and pedagogical context. In Case A, qualitative data were collected through two structured reflection surveys, with supporting artifacts including AI-generated draft programs and revised adjacency matrices. The analysis focused on student validation actions in response to AI-generated programmatic suggestions. Students were encouraged to “question the AI” by comparing AI outputs with user needs, project requirements, adjacency logic, and design intent, reflecting explainability-oriented and user-centered approaches to AI interaction (Liao et al., 2020).

The researchers reviewed the reflection responses and supporting artifacts to identify documented moments in which students retained, modified, or discarded AI-supported outputs. These actions were grouped into the Keep, Revise, Reject pattern. “Keep” referred to suggestions retained because they aligned with project criteria. “Revise” referred to suggestions modified to improve accuracy, specificity, adjacency logic, code alignment, user fit, or spatial feasibility. “Reject” referred to suggestions discarded because they were incomplete, implausible, technically inaccurate, biased, or inconsistent with the project context. This pattern was constructed inductively from comparisons between AI-generated draft materials and students’ final programming decisions. It was used as an analytical category for describing trust calibration, not as a prescriptive design method.

In Case B, quantitative data were collected through AI-supported and cloud-based design platforms. The dataset included daylight potential, indoor illuminance levels, winter unlit areas, photovoltaic potential, adjusted temperature, and wind flow patterns. The analysis focused on how students used these performance indicators to compare alternative massing options and engage simulation feedback during the design brainstorming and refinement. Because Case B involved a small

sample, these performance outputs were interpreted as illustrative evidence of performance-informed design learning rather than as representative quantitative findings.

3.3. Cross-Case Analytical Procedure

The cross-case analysis proceeded in three stages. First, the findings from each case were mapped onto the CCT layers: interpretive, algorithmic, and contextual. Second, the cases were compared through the LIGS framework to examine how students validated AI-supported outputs across linguistic and geometric workflows. Third, the comparison focused on shared indicators of human-driven validation, including questioning, verification, revision, rejection, justification, and decision ownership. The two cases were not treated as equivalent datasets. Their comparability lies in their shared pedagogical structure: both introduced AI-supported outputs during early-stage design learning and required students to evaluate those outputs before making design decisions. The analysis therefore emphasizes process comparability rather than statistical equivalence.

3.4. Validation Strategy and Ethical Considerations

The study used methodological triangulation to strengthen the credibility of the analysis. In Case A, student reflections were compared with AI-generated draft programs and revised adjacency matrices. In Case B, environmental performance outputs stored in AI-supported cloud platforms were compared across design iterations and interpreted in relation to students' design decisions. Across both cases, the analysis examined whether AI-supported outputs were accepted, modified, or rejected through human judgment. For Case A, student reflection data were collected and analyzed under an approved institutional review protocol. For Case B, the study examined course-generated design artifacts and

aggregated environmental performance outputs retrieved from cloud-based design platforms, including daylight potential, indoor illuminance levels, winter unlit areas, photovoltaic potential, adjusted temperature, and wind flow patterns. The Case B analysis did not report individually identifiable student information. Across both cases, student work was anonymized or de-identified where applicable. The scope of the study was limited to the pedagogical workflows and educational contexts described here. Each case was first analyzed by the author responsible for the corresponding pedagogical context. During the final synthesis stage, both authors jointly reviewed the case-specific findings, compared the analytical categories across the two cases, and checked whether the cross-case interpretation was supported by the available evidence. This process was used to strengthen analytical consistency across the linguistic and geometric workflows, rather than to establish formal intercoder reliability.

4. Results

This section presents findings from the two pedagogical cases. The results are organized around the LIGS framework: Case A examines linguistic informatics through AI-supported programming, while Case B examines geometric synthesis through performance-based massing. Across both cases, the findings show that AI-supported outputs became educationally meaningful when students used them as provisional materials for verification, revision, and design judgment.

4.1. Case A: Linguistic Informatics and Programmatic Validation

Case A showed that AI-supported programming helped students begin the pre-design process more quickly but also required careful human validation. Analysis of 1,081 structured reflection responses

from 61 students and supporting programming artifacts indicated that students used AI most productively when they generated outputs as drafts rather than final program decisions. The most frequent keywords in the revision process were “Program” (n = 30), “Adjacency” (n = 8), and “Square Footage” (n = 8). These patterns suggest that students used AI to generate early programmatic structure, but still had to correct spatial relationships, technical requirements, and dimensional assumptions. The spaces most frequently requiring revision included Workstations (n = 16) and Wellness Rooms (n = 10), both of which required more specific attention to user experience, spatial function, and professional expectations. Several AI-generated drafts produced plausible but incomplete programs. For example, some drafts omitted key back-of-house (BOH) functions such as legal document storage, IT closets, and plumbing requirements for catering pantries. Other suggestions required correction because they produced unrealistic square footage, weak adjacency logic, or code-related concerns. These findings indicate that AI was useful for initiating programmatic thinking, but insufficient for completing it without a critical review.

Table 3. Summary of Case A validation actions

Validation action	Description	Example from Case A	Pedagogical significance
Keep	Students retained AI-generated suggestions that aligned with project criteria	High-level themes such as flexibility, well-being, and user-centered support	AI helped generate useful starting points for concept and program development
Revise	Students modified AI outputs to improve accuracy, specificity, or feasibility	An oversized 20 × 20 ft Managing Partner’s office was revised to a more appropriate 15 × 15 ft layout	Students learned to compare AI suggestions with spatial standards and project needs
Reject	Students discarded outputs that were implausible, incomplete, inaccurate, or noncompliant	Unrealistic spatial arrangements, outdated	Students practiced critical judgment by refusing AI

workplace assumptions, ADA
or fire-safety conflicts

outputs that failed professional
criteria

As explained in Table 3, the Keep, Revise, reject pattern shows that students did not accept AI outputs passively. Instead, they used AI-generated material as a basis for comparison and correction. This pattern also suggests a shift in student effort. AI reduced some of the initial burden of generating program ideas, but increased the need for verification, explanation, and justification. In this sense, AI functioned less as a source of answers than as a prompt for metacognitive awareness and decision ownership.

4.2. Case B: Performance-Driven Geometric Synthesis

In Case B, 10 students created 20 design options using the Revit Illuminance rendering and Forma cloud systems (Figure 3). In this case, AI-supported simulations provided real-time environmental feedback during early-stage design exploration. The resulting comparisons made daylight access, solar energy potential, and related environmental factors more visible during the massing process. Rather than functioning only as post hoc evaluation tools, these simulations informed design iteration while alternative forms were still being explored.

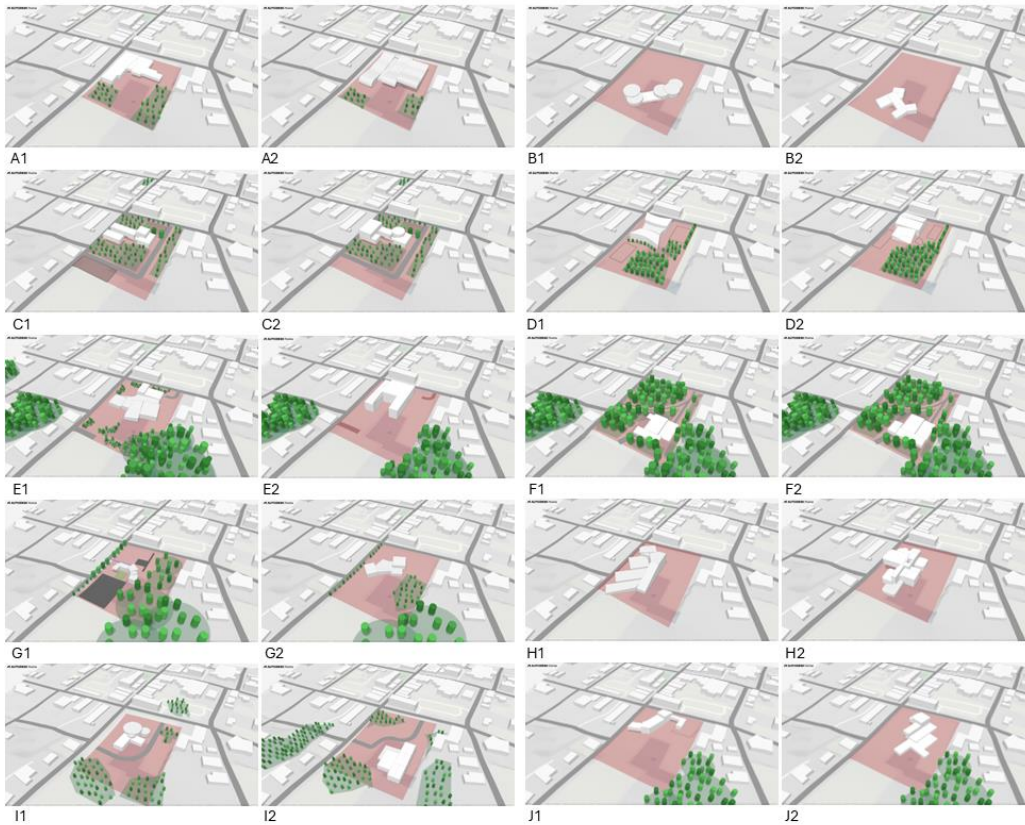


Figure 3. The tested mass design options in Autodesk Forma (Same alphabet indicates the same designer).

4.2.1. Quantitative Impacts

Across the options compared, the AI-supported massing studies were associated with substantial environmental performance enhancements. Well-lit space, measured through Vertical Sky Component (VSC), ranged from 17% to 96% of the building area, while unlit space in winter decreased from 54% to 10%. Heat stress was reduced by 8%, and the area with wind conditions favorable for sitting increased from 0% to 49% of the total site area. In addition, the potential solar energy harvested through photovoltaic (PV) systems increased fivefold across the iterative design process. At the individual designer level, winter unlit areas decreased by an average of 13%, while well-lit spaces increased by an average of 11%. The potential for solar energy generation increased by an average of

49%. Table 4 summarizes the environmental performance improvements identified across student design iterations.

Table 4. Environmental performance enhancement of each student.

	Unlit space (less than 1 hr) reduction in winter	Daylight potential (above VSC 27) improvement	Potential solar energy increase
A	1%	5%	10%
B	5%	10%	83%
C	21%	2%	68%
D	3%	8%	153%
E	5%	1%	11%
F	31%	44%	79%
G	15%	18%	7%
H	38%	12%	22%
I	9%	7%	20%
J	5%	5%	37%
Average	13%	11%	49%

4.2.2. Integrative Design Thinking

Traditional building performance simulation tools often require relatively developed design inputs, including floor plans and façade information, and are therefore commonly used to evaluate a design after major decisions have already been made. In Case B, however, AI-supported simulations were introduced during the brainstorming and massing stages, allowing environmental feedback to inform form-finding earlier in the design process. This early integration appeared to support more integrative design thinking by shifting attention from sculptural form alone toward the environmental consequences of design decisions. As students revised their proposals, they returned to the simulation platform to examine how changes in massing affected daylight access, thermal conditions, wind

comfort, and solar potential. In this sense, the simulation workflow supported iterative comparison rather than one-time evaluation. Figure 4 illustrates one example of this iterative process.

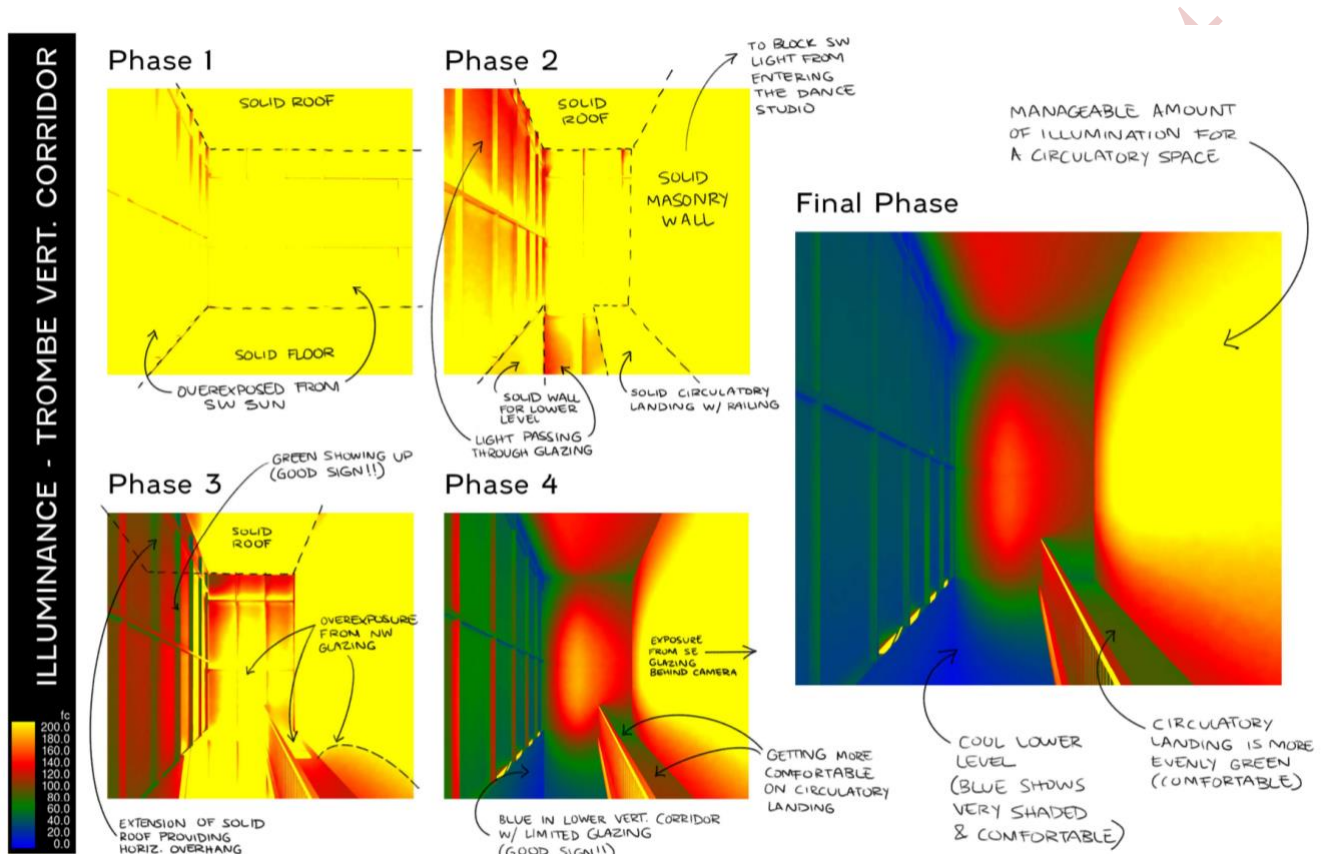


Figure 4. Iterative material studies using Illuminance rendering in Revit cloud systems (student: A1)

4.3. The Verification Gap

Across both cases, the main shared finding was a verification gap (Table 5). AI-supported systems accelerated early-stage design work, but their outputs still required students to evaluate plausibility, relevance, accuracy, and contextual fit. In Case A, the gap appeared when AI-generated program drafts seemed coherent but required correction for missing back-of-house functions, inappropriate square footage, adjacency conflicts, and code-related concerns. In Case B, the gap appeared when

performance feedback made environmental criteria visible but still required interpretation in relation to design intent and competing priorities. This gap was not only a technical limitation. It functioned as a pedagogical space where students had to exercise judgment. In the linguistic workflow, students validated AI-generated text by deciding what to keep, revise, or reject. In the geometric workflow, students validated performance feedback by comparing massing options and interpreting trade-offs. In both cases, AI was most educationally valuable when students treated its outputs as provisional inputs rather than authoritative decisions.

Table 5. Cross-case synthesis of human-driven validation

Dimension	Case A: Linguistic Informatics	Case B: Geometric Synthesis	Shared finding
AI-supported output	Draft programs, personas, adjacency suggestions	Massing options, daylight, PV, thermal, and wind indicators	AI made early-stage design information more visible
Main validation task	Check textual and programmatic plausibility	Interpret environmental feedback and design trade-offs	Students had to compare outputs with design criteria
Primary risk	Plausible but incomplete program suggestions	Over-reading performance metrics as final design answers	AI outputs required human interpretation
Student judgment	Keep, Revise, Reject	Compare, select/adjust, justify	Validation supported design decision-making

These findings respond to this study's research question by showing how students validated, revised, and justified AI-supported outputs across linguistic and geometric workflows. LIGS helped compare the two cases without treating them as equivalent datasets. The comparison shows that AI-integrated design learning depends less on whether AI produces useful outputs by itself and more on whether students are required to question, compare, and justify those outputs within disciplinary and ethical design contexts.

5. Discussion

This study examined how students in accredited design education validated, revised, and justified AI-supported outputs across linguistic and geometric design workflows. The findings show that AI became educationally meaningful when students treated generated material as a draft for further development and refinement rather than as a final design decision. In Case A, students evaluated programmatic suggestions in relation to user needs, adjacency logic, code-related concerns, and spatial feasibility. In Case B, students interpreted environmental feedback to compare massing options and justify design revisions. Across both cases, learning occurred through the human work of questioning, comparison, correction, and justification.

5.1. From AI Output to Human Validation

The findings suggest that AI-supported design education should be understood as a process of human validation. AI tools helped students generate early-stage material more efficiently, but this material still required disciplinary interpretation. In the linguistic workflow, fluent AI-generated text often appeared plausible but required correction for missing functions, unrealistic square footage, weak adjacency logic, or insufficient user specificity. In the geometric workflow, the simulation data made the environmental condition and performance more visible, while students still had to interpret trade-offs among daylight, solar potential, wind conditions, thermal comfort, and design intent. This distinction matters because AI-supported outputs can create an illusion of completeness. Text may sound coherent, and performance metrics may appear objective, but neither automatically produces final design judgment. The educational value of AI therefore lies less in automation and more in how it prompts students to ask better questions: What is missing? What is inaccurate? What conflicts with the brief?

What should be kept, revised, or rejected? This finding extends explainability-oriented approaches to AI interaction by showing that explanation needs in design education are not only technical, but also pedagogical and disciplinary (Ehsan et al., 2024; Liao et al., 2020).

5.2. LIGS as a Comparative Analytical Lens

The LIGS framework helped compare the two cases without treating them as equivalent datasets. Linguistic informatics described how students worked with AI-generated language, including briefs, personas, program suggestions, and adjacency reasoning. Geometric synthesis described how students worked with simulation-based feedback, including massing options and environmental performance indicators. These workflows differed in data type, design task, and disciplinary context, but both required students to translate AI-supported material into situated design decisions. LIGS is therefore not presented as a universal model of AI-integrated design education. Rather, it functions as a comparative analytical lens for examining how students move between AI output and human judgment across different modalities. Its contribution is to show that validation is not limited to checking technical correctness. It also involves interpreting relevance, disciplinary fit, user appropriateness, environmental trade-offs, and design intent.

5.3. Pedagogical Implications: Trust Calibration, Metacognition, and Decision Ownership

Pedagogically, these findings suggest that AI instruction should center on structured comparison rather than tool operation alone. Students need opportunities to evaluate AI-generated material against their own reasoning, disciplinary standards, and project-specific criteria. This is especially important when AI allows students to bypass the effortful process of learning why a decision is appropriate.

The Keep, Revise, reject pattern offers one practical way to structure this process. It makes student judgment visible by requiring students to identify which suggestions are useful, which require modification, and which should be discarded. This pattern supports trust calibration because students learn not to overtrust or undertrust AI outputs. Instead, they evaluate AI suggestions in relation to evidence, context, and design intention. In this sense, the pattern extends AI literacy beyond prompt writing or tool operation by making evaluation, contextual adaptation, and justification central to student learning (Knoth et al., 2024).

This process also supports metacognitive awareness and decision ownership. Students must recognize what they know, what they do not yet understand, and what requires further verification. Instructors can use AI-generated material as a teaching prompt by asking students to explain why a suggestion should be accepted, revised, or rejected. This finding aligns with recent research showing that AI assistance must be structured to preserve student agency and that generative AI learning environments require metacognitive support (Darvishi et al., 2024; Xu et al., 2025). AI can therefore function as a pedagogical sparring partner, not because it provides correct answers, but because it creates opportunities for students to practice judgment.

5.4. Ethical Stewardship as Situated Validation

The findings also show that ethical stewardship should be integrated into validation rather than treated as a separate concern. In Case A, ethical validation involved examining whether AI-generated program suggestions adequately reflected user needs, accessibility, inclusivity, and spatial priorities. In Case B, ethical validation involved interpreting performance data without reducing architecture to numerical optimization alone. In both cases, students needed to ask whether AI-supported outputs were

appropriate for the specific users, spaces, and design contexts under consideration. This approach aligns human-centered AI with design education. Human-centered validation requires students to examine not only whether output is useful, but also whether it is incomplete, biased, decontextualized, or misaligned with human needs. Ethical stewardship therefore extends the meaning of validation. It includes technical review, disciplinary judgment, contextual sensitivity, and responsibility for the consequences of design decisions, especially because AI systems can reproduce bias, normalize dominant assumptions, or obscure unequal representation (Buolamwini, 2024).

5.5. Limitations and Case Comparability

It should be noted that this research is an exploratory study with the following limitations. First, the two cases should not be read as equivalent datasets. Case A provided richer qualitative evidence through student reflections, AI-generated draft programs, and revised adjacency matrices. Case B provided a smaller set of performance-based design comparisons from ten students and twenty design options. The Case B findings are therefore illustrative rather than representative. Second, the cases differed in discipline, student level, workflow, data format, and analytical emphasis. Their comparability lies in process, not scale: both cases introduced AI-supported outputs during early-stage design and required students to evaluate them before deciding. The study therefore emphasizes process comparability rather than statistical equivalence. Third, LIGS and CCT were used as interpretive frameworks, not as validated measurement instruments. Future research should test these frameworks across larger samples, additional institutions, and more standardized coding procedures. Further studies could also examine whether the Keep, Revise, Reject pattern supports student learning across other design disciplines, different AI tools, and longer project timelines. The present study contributes a

validation-centered account of AI-integrated design pedagogy. Rather than treating human oversight as a corrective step after AI use, it positions validation as the core learning activity through which students develop trust calibration, metacognitive awareness, and responsible design judgment.

6. Conclusion

This study examined how students in accredited design education validated, revised, and justified AI-supported outputs across linguistic and geometric design workflows. Using LIGS as a cross-modal analytical framework, the study compared a text-to-program workflow in a CIDA-aligned interior architecture course with a performance-based massing workflow in a NAAB-aligned architecture course. The findings show that AI became pedagogically meaningful when students treated generated material as a basis for inquiry rather than as a final design solution. Across the two cases, AI supported early-stage design learning in different but comparable ways. In Case A, students used the Keep, Revise, reject pattern to evaluate AI-generated programmatic suggestions against user needs, adjacency logic, code-related concerns, and spatial feasibility. In Case B, students used environmental performance feedback to compare massing options and interpret trade-offs among daylight, solar potential, thermal conditions, wind, and design intent. These findings suggest that the educational value of AI lies not in replacing disciplinary expertise, but in prompting students to practice trust calibration, metacognitive awareness, and decision ownership.

The study contributes a validation-centered account of AI-integrated design pedagogy. It shows that students need structured opportunities to compare AI outputs with disciplinary standards, project-specific criteria, and human-centered design concerns. In this sense, AI literacy in design education should extend beyond tool operation or prompt writing to include critical curation, contextual

interpretation, and ethical stewardship. Students must learn not only how to generate or simulate design options, but also how to evaluate whether those options are accurate, inclusive, relevant, and appropriate for the specific design context.

This study is limited by its exploratory dual-case design and by the different forms of evidence generated in the two cases. Case A provided richer qualitative evidence, while Case B provided a smaller set of illustrative performance-based comparisons. Accordingly, LIGS should be understood as a comparative analytical framework rather than as a fully validated model. Future research should test this framework across larger samples, additional institutions, longer project timelines, and more standardized coding procedures.

Ultimately, the value of AI in architecture and interior design education depends on whether students remain active interpreters and responsible decision-makers within increasingly automated design environments. AI can expand what students generate and compare, but design learning still depends on the human capacity to question, revise, justify, and act ethically.

Jinoh Park

Department of Interior Architecture and Design, Fay Jones School of Architecture and Design, University of Arkansas, Fayetteville, United States

jinohp@uark.edu

ORCID iD of Author 1

<https://orcid.org/0000-0003-1725-6435>

Jinoh Park is an assistant professor of interior architecture and design at the University of Arkansas. His research examines AI-integrated design education, neuroarchitecture, human-centered spatial experience, and evidence-based learning environments. Drawing on professional design practice and interdisciplinary study, he investigates how emerging technologies can strengthen human judgment, representation, and decision-making in architectural and interior design education.

Soo Jeong Jo

School of Architecture, Louisiana State University, Baton Rouge, LA, United States

soojol@lsu.edu

ORCID ID: 0000-0003-0802-6363

Soo Jeong Jo is an assistant professor of architecture at Louisiana State University and a licensed architect in France and South Korea. She received her Ph.D. at Virginia Tech. Her research focuses on the early stages of high-performance building design based on computational simulations. Prior to her research and teaching, Soo Jeong was a practicing architect in Paris, New York City, and Seoul, and participated in various award-winning projects.

References

- Arora, P. (2025). Creative data justice: A decolonial and indigenous framework to assess creativity and artificial intelligence. *Information, Communication & Society*, 28(13), 2231–2247.
<https://doi.org/10.1080/1369118X.2024.2420041>
- Buolamwini, J. (2024). *Unmasking AI: My mission to protect what is human in a world of machines*. Random House.
- Chiou, L.-Y., Hung, P.-K., Liang, R.-H., & Wang, C.-T. (2023). Designing with AI: An Exploration of Co-Ideation with Image Generators. In *Proceedings of the 2023 ACM Designing Interactive Systems Conference* (pp. 1941–1954). Association for Computing Machinery.
<https://doi.org/10.1145/3563657.3596001>
- Council for Interior Design Accreditation. (2024). *CIDA Professional Standards 2024*.
<https://cida.org/professional-standards>
- Cross, N. (2001). Designerly ways of knowing: Design discipline versus design science. *Design Issues*, 17(3), 49–55. <https://doi.org/10.1162/074793601750357196>
- Darvishi, A., Khosravi, H., Sadiq, S., Gašević, D., & Siemens, G. (2024). Impact of AI assistance on student agency. *Computers & Education*, 210, 104967.
<https://doi.org/10.1016/j.compedu.2023.104967>
- Ehsan, U., Liao, Q. V., Passi, S., Riedl, M. O., & Daumé, H. (2024). Seamful XAI: Operationalizing seamful design in explainable AI. *Proceedings of the ACM on Human-Computer Interaction*, 8(CSCW1), 1–29. <https://doi.org/10.1145/3637396>
- Kalay, Y. E. (2004). *Architecture's new media: Principles, theories, and methods of computer-aided design*. MIT Press.

- Knoth, N., Tolzin, A., Janson, A., & Leimeister, J. M. (2024). AI literacy and its implications for prompt engineering strategies. *Computers and Education: Artificial Intelligence*, 6, 100225. <https://doi.org/10.1016/j.caeai.2024.100225>
- Liao, Q. V., Gruen, D., & Miller, S. (2020). Questioning the AI: Informing design practices for explainable AI user experiences. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1–15). Association for Computing Machinery. <https://doi.org/10.1145/3313831.3376590>
- Melker, S., Gabrils, E., Villavicencio, V., Faraon, M., & Rönkkö, K. (2025). Artificial intelligence for design education: A conceptual approach to enhance students' divergent and convergent thinking in ideation processes. *International Journal of Technology and Design Education*, 35(5), 1871–1899. <https://doi.org/10.1007/s10798-025-09964-3>
- National Architectural Accrediting Board (NAAB). (2025). *2020 Conditions for Accreditation (May 2025 Revised Edition)*. https://higherlogicdownload.s3.amazonaws.com/NAAB/21e8eae7-e532-47c0-bff1-4111ca0d4fb0/UploadedImages/PDFs/2020_NAAB_Conditions_for_Accreditation_Rev_2025_0501.pdf
- Ofosu-Asare, Y. (2025). Cognitive imperialism in artificial intelligence: Counteracting bias with Indigenous epistemologies. *AI & SOCIETY*, 40(4), 3045–3061. <https://doi.org/10.1007/s00146-024-02065-0>
- Oxman, R. (2008). Performance-based design: Current practices and research issues. *International Journal of Architectural Computing*, 6(1), 1–17. <https://doi.org/10.1260/147807708784640090>

- Park, J. (2026). Artificial intelligence in design education: A systematic review and CIDA-Aligned framework for interior architecture and design. *Enquiry The ARCC Journal for Architectural Research*, 23(1), 60–80. <https://doi.org/10.17831/enqarcc.v23i1.1329>
- Stoimenova, N., & Price, R. (2020). Exploring the nuances of designing (with/for) artificial intelligence. *Design Issues*, 36(4), 45–55. https://doi.org/10.1162/desi_a_00613
- Tellez, F. A. (2025). Reflecting on the integration of generative AI in design education: Lessons from the field. *Voces y Silencios. Revista Latinoamericana de Educación*, 16(2), 169–191. <https://doi.org/10.18175/VyS16.2.2025.9>
- Verganti, R., Vendraminelli, L., & Iansiti, M. (2020). Innovation and design in the age of artificial intelligence. *Journal of Product Innovation Management*, 37(3), 212–227. <https://doi.org/10.1111/jpim.12523>
- Xu, X., Qiao, L., Cheng, N., Liu, H., & Zhao, W. (2025). Enhancing self-regulated learning and learning experience in generative AI environments: The critical role of metacognitive support. *British Journal of Educational Technology*, 56(5), 1842–1863. <https://doi.org/10.1111/bjet.13599>
- Zhang, L. (2023). The ethical turn of emerging design practices. *She Ji: The Journal of Design, Economics, and Innovation*, 9(3), 311–329. <https://doi.org/10.1016/j.sheji.2023.09.002>