

Beyond “Did You Use AI?”: A Framework for Measuring Human Contribution in AI-Assisted Writing

Kevin Mentzer
Kevin.Mentzer@Nichols.edu
Nichols College
Dudley, MA 01571, USA

Abstract

Generative artificial intelligence (AI) tools are now widely used to support academic and professional writing, creating new challenges for authorship, disclosure, and assessment. Existing institutional and publisher policies typically require disclosure of AI use but provide limited guidance for evaluating the degree of human contribution in AI-assisted work. This paper introduces the Human-AI Contribution Model (HACM), a practical framework for assessing human contribution across six dimensions: Source Origination, Directive Specificity, Iterative Steering, Human Text Inclusion, Transform Depth, and Attribution & Judgment. HACM combines these dimensions into a Human Authorship Score (HAS) that maps to four interpretable authorship classifications, ranging from AI-Created, Human-Curated to Human-Author, AI-Assisted. The paper presents the HACM framework, illustrates its use through representative writing scenarios, and discusses applications for information systems education, including assignment design, student reflection, AI-use disclosure, and policy development. By moving beyond binary “AI used / not used” distinctions, HACM provides educators and institutions with a transparent and teachable method for characterizing human-AI collaboration in writing.

Keywords: Generative AI; authorship; AI-assisted writing; academic integrity; disclosure; human-AI collaboration; information systems education.

Beyond “Did You Use AI?”: A Framework for Measuring Human Contribution in AI-Assisted Writing

Kevin Mentzer

1. INTRODUCTION

Generative artificial intelligence (AI) tools have rapidly become embedded in academic and professional writing workflows. Students use systems such as ChatGPT, Claude, and Gemini to brainstorm ideas, generate outlines, draft content, revise prose, and verify information. Faculty members increasingly encounter assignments that reflect varying levels of AI involvement, ranging from minor editorial assistance to extensive AI-generated content. Similar patterns are emerging in scholarly publishing, where authors use AI to support literature reviews, manuscript drafting, and language editing. This paper presents the Human-AI Contribution Model (HACM) as a conceptual and pedagogical framework rather than as an empirically validated measurement instrument.

Despite widespread adoption, guidance regarding AI-assisted authorship remains underdeveloped. Most institutional policies, publisher guidelines, and academic integrity statements focus on disclosure requirements and generally agree that AI systems cannot be listed as authors. However, these policies rarely provide a practical method for determining the extent of human contribution to a finished work. As a result, authorship is often treated as a binary condition: AI was either used or not used. Such an approach fails to capture the reality that human-AI collaboration exists along a continuum of involvement and responsibility.

The challenge is particularly relevant in information systems education, where instructors must evaluate student work produced with increasingly sophisticated AI tools. A student who uses AI to proofread a self-authored report represents a fundamentally different case than a student who submits largely AI-generated content with minimal oversight. Yet many existing policies lack a mechanism for distinguishing between these scenarios in a transparent and consistent manner.

A central motivation for HACM is to move classroom and policy conversations beyond the simple dichotomy of AI-used versus AI-not-used. This binary framing treats very different practices - brainstorming, outlining, drafting, rewriting, proofreading, citation checking, and full content generation - as if they belong to the same category. For instructors and students, this obscures the more important question: not merely whether AI was used, but how it was used, where it entered the writing process, and whether its use altered the human contribution being assessed. By breaking AI-assisted writing into separate axes, HACM provides a more precise vocabulary for discussing allowable and non-allowable uses of AI in relation to specific learning objectives.

Consider two students in the same information systems course. One drafts a systems analysis report independently and uses AI only to improve sentence clarity. Another enters the assignment prompt into an AI system, lightly edits the output, and submits it as final work. Both students could truthfully disclose that they “used AI,” but the instructional and authorship implications of those disclosures are very different. Without a more precise vocabulary, instructors are left to interpret these cases inconsistently.

Importantly, HACM is not intended to imply that a higher level of human contribution always produces a better or more meaningful final product. In many educational and professional contexts, effective use of tools is itself an important learning outcome. Students may appropriately use AI to improve clarity, explore alternatives, test ideas, or support revision, just as they may use spreadsheets, statistical software, or database tools to complete disciplinary work. The purpose of HACM is therefore not to rank human-only work above AI-assisted work, but to make visible the type and extent of human contribution so that instructors, students, reviewers, and institutions can evaluate whether the use of AI aligns with the goals of the task.

To address this gap, we propose the Human-AI Contribution Model (HACM), a practical framework for assessing the degree of human contribution in AI-assisted writing. HACM evaluates six dimensions of

authorship: Source Origination, Directive Specificity, Iterative Steering, Human Text Inclusion, Transform Depth, and Attribution & Judgment. Scores across these dimensions are combined into a Human Authorship Score (HAS) that maps to four interpretable authorship classifications ranging from AI-Created, Human-Curated to Human-Author, AI-Assisted.

The contribution of this paper is threefold. First, we introduce a multidimensional framework for discussing human-AI writing collaboration in more precise terms than binary AI-use disclosure allows. Second, we propose four authorship classifications that provide a practical vocabulary for describing the dominant relationship between human and AI contributions. Third, we offer a six-axis scoring process that can support assignment design, student reflection, faculty calibration, and disclosure practices. By moving beyond binary notions of AI usage, HACM offers educators, researchers, and editors a practical tool for discussing authorship, responsibility, and appropriate AI use in the age of generative AI.

2. RELATED WORK

The Human-AI Contribution Model (HACM) draws upon prior work in authorship attribution, transparency frameworks, and human-AI collaboration in educational settings. While these streams of research provide important foundations for understanding AI-assisted work, none offer a practical mechanism for quantifying the degree of human contribution in AI-generated or AI-assisted writing.

2.1 Authorship, Contributorship, and Disclosure

Questions surrounding authorship are not new. Academic disciplines have long relied on authorship guidelines and contributor frameworks to distinguish between substantive intellectual contributions and supporting activities. One influential example is the Contributor Roles Taxonomy (CRediT), which identifies specific contributor roles such as conceptualization, methodology, writing, supervision, and validation. Rather than treating authorship as a binary designation, CRediT recognizes that scholarly work is often the product of multiple forms of contribution.

The emergence of generative AI has complicated traditional notions of authorship. Organizations such as the Committee on Publication Ethics (COPE), ACM, IEEE, and major academic publishers generally agree that AI systems cannot be listed as authors because they cannot assume responsibility for the content they generate. Similarly, the U.S. Copyright Office has emphasized the importance of distinguishing human-authored elements from material generated by AI when evaluating works containing AI-generated content (U.S. Copyright Office, 2023). At the same time, these organizations increasingly require disclosure of AI usage during manuscript preparation. Although these policies promote transparency, they typically do not provide a systematic method for evaluating how much of a work should be attributed to human versus AI effort.

This gap suggests a need for mechanisms that move beyond simple disclosure statements and provide a more nuanced understanding of human contribution within AI-assisted workflows.

2.2 Transparency and Documentation Frameworks

A second foundation for HACM comes from the growing body of work on transparency artifacts in artificial intelligence. Datasheets for Datasets introduced structured documentation intended to describe how datasets were collected, processed, and intended to be used (Gebru et al., 2021). Model Cards extended this concept to machine learning models by documenting intended use, performance characteristics, and limitations (Mitchell et al., 2019). Subsequent work on system cards and AI governance frameworks further emphasized transparency, accountability, and responsible deployment of AI systems (Gursoy & Kakadiaris, 2022; National Institute of Standards and Technology, 2023).

These approaches share several common characteristics. They rely on standardized reporting structures, encourage consistent disclosure practices, and improve auditability across contexts. Although they focus on datasets and AI systems rather than writing processes, they demonstrate how structured documentation can improve trust, comparability, and accountability.

HACM applies this same design philosophy to authorship. Rather than documenting datasets or models, HACM documents the distribution of human and AI contributions within a written artifact through a

standardized scoring framework.

2.3 Human–AI Collaboration in Educational Contexts

Research in information systems education has increasingly examined how generative AI affects student learning, assessment, and academic integrity. Recent IS education research highlights the need for instructors to maintain oversight of AI-assisted student work and to ensure that such work still reflects human judgment and accountability (Gimpel et al., 2025; Zhang, 2025). At the same time, researchers have explored approaches for detecting AI-generated content, developing instructional guidance, and designing policies that encourage responsible AI use.

More broadly, research on human interaction with automation has long emphasized that automated systems vary in the level of support they provide and the degree of human oversight they require (Parasuraman et al., 2000). This perspective is useful for AI-assisted writing because the central issue is not simply whether automation is present, but how responsibility, control, and judgment are distributed between the human and the system.

Recent work suggests that binary distinctions between AI-generated and human-generated writing may be insufficient for describing contemporary writing practices. Paullet, Pinchot, Kinney, and Stewart (2025) found that hybrid human–AI writing presents challenges for detection-based approaches, particularly when students combine AI-generated content with their own revisions. These findings highlight the limitations of treating AI involvement as a simple yes-or-no condition.

HACM differs from detection-oriented approaches in both purpose and logic. Detection tools attempt to infer whether text was generated by AI, often from the final artifact alone. HACM instead asks writers and evaluators to document the writing process and classify the distribution of human and AI contribution. This distinction is important because hybrid writing may not be reliably reducible to textual traces, especially when human revision and AI editing occur across multiple drafts.

Educational assessment provides another important foundation for HACM. Rubrics are widely used to evaluate complex constructs through multiple criteria and are valued for their transparency, consistency, and interpretability. Prior work in information systems education has demonstrated how structured rubrics and evaluation frameworks can improve instructional consistency and provide meaningful feedback to learners (Lending et al., 2022; McHenry, 2022). HACM extends this tradition by treating authorship as a multidimensional construct that can be assessed through a set of clearly defined criteria. Taken together, the literature suggests that effective governance of AI-assisted writing requires more than disclosure alone. Existing work provides mechanisms for documenting contributions, increasing transparency, and evaluating complex behaviors, but it does not offer a standardized method for quantifying the degree of human contribution within AI-assisted writing workflows. HACM addresses this gap through a six-axis framework that measures human involvement across the writing process and translates those measurements into an interpretable authorship classification.

3. THE HUMAN-AI CONTRIBUTION MODEL (HACM)

3.1 Framework Overview

The Human–AI Contribution Model (HACM) is a multidimensional framework designed to measure the degree of human contribution within AI-assisted writing. Rather than treating AI usage as a binary condition, HACM conceptualizes authorship as a continuum of human and AI involvement. The model evaluates six dimensions of the writing process and combines them into a single Human Authorship Score (HAS) ranging from 0 to 100.

The six axes were selected to represent the major forms of contribution that occur throughout the lifecycle of a written artifact, from ideation and task specification to drafting, revision, and verification. Collectively, they provide a structured mechanism for assessing who contributed what to a finished work and how responsibility for the content is distributed between human and AI participants.

3.2 The Six HACM Axes

The six HACM axes represent distinct dimensions of authorship and responsibility:

- Source Origination (SO): The extent to which core ideas, claims, datasets, and intellectual direction originate from the human author.
- Directive Specificity (DS): The degree to which the human constrains and directs AI behavior through prompts, instructions, and requirements.
- Iterative Steering (IS): The amount of human involvement in revising, refining, and redirecting AI outputs.
- Human Text Inclusion (HT): The proportion of the final artifact directly authored by the human.
- Transform Depth (TD): The extent to which AI modifies or rewrites human-generated content.
- Attribution & Judgment (AJ): The degree to which the human verifies facts, evaluates claims, and assumes responsibility for the final artifact.

Table 1 presents the scoring rubric used to evaluate each axis. Detailed scoring anchors for each axis are provided in Appendix A.

Axis	Score				
	0	1	2	3	4
SO	No human ideas originate the work	Minimal human idea input	Shared idea formation	Human-dominated idea formation	Human solely originates all ideas & structure
DS	No constraints	Basic instruction	Moderate guidance	High detail	Fully constrained, precise directives
IS	No iteration	Reactive minor edits	Some targeted feedback	Multi-round revision	Human leads iterative transformations
HT	0% human text	<25%	25–50%	50–75%	>75% human text
TD	Deep model rewrite	Major rewrite	Mixed	Light revision	Proofreading only
AJ	No factual verification	Minimal	Partial checking	Systematic checking	Full human responsibility

Table 1. HACM Scoring Rubric for Assessing Human Contribution in AI-Assisted Writing

The six axes were selected to provide a parsimonious but comprehensive representation of human authorship. The intent is not to capture every possible feature of writing, but to identify the smallest set of dimensions needed to distinguish materially different human-AI authorship patterns. Together, the axes capture ideation, instruction, revision, textual contribution, transformation, and accountability. Additional dimensions were considered during model development, including creativity, stylistic voice, and citation management, but these concepts overlap substantially with the six core constructs: creativity is treated primarily through Source Origination, stylistic voice through Human Text Inclusion and Transform Depth, and citation checking through Attribution & Judgment. Conversely, collapsing dimensions would obscure important distinctions between different forms of human contribution. HACM therefore adopts a six-axis structure that balances conceptual completeness with practical usability.

3.3 Human Authorship Score (HAS)

Once each axis has been scored on a scale from 0 to 4, the values are combined into a weighted Human Authorship Score (HAS). The HAS provides a continuous measure of human contribution ranging from 0 to 100.

$$HAS = 25 \times (0.25SO + 0.15DS + 0.20IS + 0.20HT + 0.10TD + 0.10AJ)$$

The weighting scheme reflects the view that conceptual origination, iterative oversight, and textual contribution represent the most significant indicators of authorship. Source Origination receives the

highest weight because intellectual ownership begins with the generation of ideas, arguments, and research direction. Iterative Steering and Human Text Inclusion are also weighted heavily because they reflect direct human engagement in shaping the artifact. Directive Specificity, Transform Depth, and Attribution & Judgment remain important but receive slightly lower weights because they serve supporting roles within the authorship process.

The weights are conceptually motivated rather than empirically derived. They reflect the authors' judgment that idea origination, iterative control, and human-authored text are especially important indicators of authorship in writing contexts. However, the weighting scheme should be viewed as provisional and adjustable rather than definitive. Instructors, journals, programs, or institutions may choose to use the unweighted axis scores, adjust the weights for a particular assignment or policy context, or rely primarily on the resulting authorship classification. Future empirical work should examine whether alternative weighting schemes produce more reliable or valid classifications across contexts.

3.4 Authorship Classifications

The classifications are intended to be the primary interpretive output of HACM. The numerical HAS provides a structured pathway for reaching a classification, but the classification itself is often more useful for communication, policy, and disclosure. Instructors, students, and reviewers may find the category labels easier to apply than the precise numerical score, particularly in contexts where the available evidence is based on self-report or incomplete documentation.

Scores near classification thresholds should be interpreted with caution. For example, a HAS of 49 and a HAS of 51 may fall into different categories, but the underlying workflows may be substantively similar. In such cases, instructors or reviewers should examine the individual axis scores and supporting documentation rather than relying only on the numerical cutoff. The classifications should therefore be understood as interpretive bands rather than rigid adjudication rules.

To improve interpretability, HACM maps HAS values to four authorship classifications.

HAS Range	Classification	Primary Authorship Pattern	Typical Disclosure Implication
0–24	AI-Created, Human-Curated	AI originates most ideas, structure, and prose; human contribution is limited to selection, acceptance, or light curation.	Strong disclosure required; generally inappropriate for claims of human scholarly authorship without substantial revision.
25–49	AI-Led Co-Creation	Human provides topic, constraints, or sources, but AI leads drafting, organization, and much of the final expression.	Clear disclosure required; suitable for low-stakes drafting or formative work, not high-stakes authorship without additional human revision and verification.
50–74	Human-Led Co-Creation	Human originates the core framing and directs revisions; AI assists with drafting, organization, or prose refinement.	Disclosure recommended; may be appropriate where AI-assisted drafting is permitted and human responsibility is retained.
75–100	Human-Author, AI-Assisted	Human is primary author of ideas, structure, and prose; AI provides limited support such as editing, proofreading, or formatting.	Brief disclosure generally sufficient where required; consistent with human authorship and AI editorial assistance.

Table 2. HACM Authorship Classifications Based on Human Authorship Score (HAS)

These classifications provide a common vocabulary for describing the relationship between human and AI contributions. Rather than focusing solely on whether AI was used, the classifications communicate the dominant authorship pattern represented by the writing process.

3.5 Applying HACM

The following section demonstrates the application of HACM through a series of representative writing scenarios. These examples illustrate how differing levels of human involvement affect individual axis scores, resulting HAS values, and final classifications.

4. ILLUSTRATIVE APPLICATION OF HACM

To demonstrate how HACM distinguishes among different forms of human-AI collaboration, Table 3 presents four representative writing scenarios. The scenarios are ordered from lowest to highest Human Authorship Score (HAS) and illustrate how differences in human origination, prompting, revision, textual contribution, transformation, and verification affect the final authorship classification.

Case	Scenario	Representative Axis Scores	HAS	Classification
1	Minimal human input: human enters a generic prompt such as "Write a paragraph on blockchain" and submits the output with little or no revision.	SO=1, DS=0, IS=0, HT=0, TD=0, AJ=0	6	AI-Created, Human-Curated
2	AI-led draft: human provides a topic, sources, or broad direction, while AI generates most of the structure and prose.	SO=2, DS=2, IS=1, HT=1, TD=1, AJ=1	35	AI-Led Co-Creation
3	Human-led co-creation: human provides a detailed outline, argument flow, and examples, while AI assists with drafting and refinement.	SO=3, DS=3, IS=3, HT=2, TD=2, AJ=2	65	Human-Led Co-Creation
4	Human-authored draft with AI editing: human writes the substantive content and uses AI for proofreading, clarity, or formatting.	SO=4, DS=3, IS=3, HT=4, TD=4, AJ=3	89	Human-Author, AI-Assisted

Table 3. Representative HACM Scoring Scenarios

The four cases show why a binary distinction between "AI used" and "AI not used" is insufficient. In all four scenarios, AI is involved in the writing process, but the nature of that involvement differs substantially. In the lowest-scoring case, the human provides only a generic prompt and accepts the AI output with little or no revision or verification. In the highest-scoring case, the human originates the ideas, drafts the substantive text, and uses AI only for limited editing or proofreading. These two workflows both involve AI, but they represent fundamentally different authorship relationships.

The middle cases illustrate the more ambiguous zone in which many educational and professional writing practices now occur. A writer may provide a topic, sources, or general direction while relying on AI to generate most of the prose, resulting in an AI-Led Co-Creation classification. Alternatively, the writer may originate the argument, provide a detailed outline, and use AI to assist with drafting or refinement, resulting in a Human-Led Co-Creation classification. HACM makes these distinctions explicit by assigning separate scores to the different forms of human contribution rather than relying on a single disclosure statement.

The scenarios in Table 3 can also be used as calibration exercises for instructors or departments developing AI-use policies. Faculty members can independently score the same sample workflow, compare axis scores, and discuss why they assigned different values. These conversations can reveal where expectations differ, such as whether AI-assisted outlining should affect Source Origination, whether grammar-level editing should affect Transform Depth, or how much verification is needed for a high Attribution & Judgment score. In this way, HACM can support not only student self-reflection but also faculty and department-level alignment around acceptable and unacceptable uses of generative AI.

In instructional settings, these classifications can help faculty communicate expectations for appropriate AI use. For example, an instructor might permit Human-Author, AI-Assisted work for formal writing assignments, allow Human-Led Co-Creation for drafts or formative exercises, and prohibit AI-Created, Human-Curated submissions for graded work requiring independent reasoning. In scholarly settings, the same classification structure can support more precise disclosure by clarifying whether AI served primarily as an editor, drafting assistant, or content generator.

5. DISCUSSION AND FUTURE WORK

HACM is intended to function as both a summative classification tool and a formative reflection tool. As a summative tool, it helps instructors, editors, and institutions classify the degree of human contribution in AI-assisted writing. As a formative tool, it can help writers understand how different practices—such as originating ideas, writing substantive prose, revising AI outputs, and verifying claims—affect authorship and disclosure. This dual use is especially important in educational settings, where the goal is not only to detect inappropriate AI use but also to teach students how to use AI responsibly.

5.1 Instructor Use

Instructors can use HACM to clarify expectations for AI-assisted work before an assignment begins, guide student reflection during the writing process, and evaluate submitted work after completion. Rather than issuing a broad statement such as “AI may be used” or “AI may not be used,” instructors can specify which HACM classifications are acceptable for a particular assignment. For example, an instructor might require final research papers to fall within the Human-Author, AI-Assisted classification, while allowing Human-Led Co-Creation for early-stage brainstorming, outlining, or draft development. Conversely, AI-Created, Human-Curated submissions may be prohibited for assignments intended to assess independent analysis, writing ability, or conceptual understanding.

A central instructional value of HACM is that it gives students and faculty a more precise vocabulary for discussing AI use. Many classroom policies currently frame AI use as a binary condition: either AI was used or it was not. This framing is increasingly inadequate because it treats proofreading, brainstorming, outlining, drafting, rewriting, citation checking, and full content generation as if they were equivalent. By separating AI-assisted writing into six axes, HACM helps instructors and students distinguish among different types of AI involvement and discuss which forms are appropriate for a given assignment.

For example, an instructor may decide that AI support for grammar and clarity is acceptable because it affects Transform Depth but does not substantially reduce Source Origination or Human Text Inclusion. The same instructor may prohibit AI-generated argumentation because it would lower Source Origination and Iterative Steering, even if the student later edited the prose. In this way, HACM supports more nuanced policy language: instead of asking only “Was AI used?” instructors can ask “Which part of the writing process did AI influence, and did that influence compromise the learning objective being assessed?”

HACM can also be incorporated into student reflection and disclosure practices. Students may be asked to submit a brief HACM reflection alongside their assignment, identifying the AI tools used, describing the purpose of that use, estimating scores for each of the six axes, and explaining the resulting authorship classification. This reflection shifts the instructional conversation away from whether AI was used and toward how it was used, what the student contributed, and whether the final submission reflects the learning outcomes of the assignment. Such reflection can also help students develop more deliberate AI-use strategies by showing how prompting, revision, original writing, and verification affect authorship. Appendix B provides a sample student reflection template that instructors can adapt for classroom use.

In this instructional use, the six axes and the authorship classifications serve related but distinct purposes. The axes are most useful as a planning and reflection tool because they help instructors and students discuss specific forms of AI involvement before or during an assignment. The classifications are most useful as a communication tool because they summarize the overall authorship pattern after the work is completed. This distinction helps avoid treating the HAS as a purely mechanical score while preserving the pedagogical value of the framework.

Finally, HACM can support faculty calibration and program-level policy development. Instructors teaching multiple sections of the same course can score sample workflows together to develop shared expectations for acceptable AI use. Departments may also use HACM classifications to define common policy language across assignments while still allowing individual instructors to adjust expectations based on course level, task type, or learning objective. In this way, HACM provides a common vocabulary for AI-assisted writing without requiring a single universal rule for all assignments.

5.2 Disclosure and Policy Use

HACM also provides a practical structure for disclosure. Rather than requiring a generic statement that AI was used, HACM allows writers to describe the nature and extent of AI involvement. A low HAS classification would require more detailed disclosure because the AI contributed substantially to the artifact. A high HAS classification may require only a brief statement that AI was used for editing, proofreading, or formatting. This approach aligns disclosure expectations with the actual level of AI contribution. Appendix C provides sample disclosure statements aligned with each HACM classification.

At the policy level, HACM classifications can help institutions distinguish between acceptable AI assistance, reportable AI collaboration, and uses that may conflict with assignment or publication expectations. Rather than adopting a single institution-wide rule for all AI use, programs can identify which classifications are acceptable for different types of work. For example, Human-Author, AI-Assisted work may be appropriate for most formal submissions, while Human-Led Co-Creation may be appropriate for drafts, brainstorming, or scaffolded assignments. AI-Created, Human-Curated work may require explicit permission or may be disallowed when the purpose of the task is to assess independent reasoning or writing ability.

5.3 Limitations and Future Work

Several limitations should be acknowledged. First, HACM relies partly on self-report unless supported by documentation such as prompt logs, drafts, revision histories, or version-control records. Writers may intentionally or unintentionally overstate their human contribution. Second, the weighting scheme should be viewed as provisional. Different disciplines, institutions, or assignment contexts may value the six axes differently. Third, the model has not yet been empirically validated across large samples of student or professional writing. These limitations do not undermine HACM's usefulness as an initial framework, but they indicate the need for further testing.

Future work should evaluate HACM in classroom and scholarly contexts. Possible validation studies include inter-rater reliability testing, comparison of HACM classifications with expert judgments of authorship, and experimental studies that manipulate the degree of AI involvement in writing workflows. Inter-rater reliability studies are particularly important because the model depends on consistent interpretation of axis scores across instructors, students, and reviewers. Researchers could also examine whether HACM-based reflections improve student understanding of responsible AI use. Finally, future versions of the model may incorporate optional instrumentation, such as text-similarity measures, prompt-history analysis, or citation-auditing tools, to support more objective scoring.

A further area for future work is adapting HACM to AI-assisted programming assignments. Although this paper focuses on writing, the six-axis structure may transfer to code-based artifacts with modified scoring anchors. Source Origination could capture whether the student designed the algorithm, data model, or solution strategy; Directive Specificity could reflect the precision of prompts, requirements, constraints, or test cases; Iterative Steering could capture debugging, refactoring, and revision; Human Text Inclusion could be reinterpreted as human-authored code, queries, scripts, or configuration; Transform Depth could distinguish syntax-level correction from architectural rewriting; and Attribution & Judgment could include verification of functionality, security, documentation, and edge cases. Because many information systems courses involve programming, analytics, databases, and systems development, validating HACM for code-based work represents an important extension of the model.

6. CONCLUSIONS

Generative AI has made authorship more difficult to define, disclose, and assess. Treating AI use as a binary condition does not reflect the range of ways students, faculty, and professionals now use AI in

writing workflows. A student who asks AI to proofread a self-authored paper and a student who submits a largely AI-generated draft have both “used AI,” but the authorship implications of those uses are fundamentally different.

This paper introduced the Human-AI Contribution Model (HACM), a six-axis framework for evaluating human contribution in AI-assisted writing. By scoring Source Origination, Directive Specificity, Iterative Steering, Human Text Inclusion, Transform Depth, and Attribution & Judgment, HACM produces a Human Authorship Score that maps to four interpretable authorship classifications. These classifications provide a practical vocabulary for describing human-AI collaboration and for aligning disclosure expectations with the actual level of AI involvement.

For information systems educators, HACM offers a way to move beyond detection-oriented responses to generative AI and toward more transparent, teachable, and defensible assessment practices. The model can support assignment design, student reflection, faculty calibration, and institutional policy development. Rather than asking only whether AI was used, HACM encourages instructors and students to ask how AI was used, what the human contributed, and who remains responsible for the final work.

Future research should validate HACM across instructional, scholarly, and code-based contexts, test inter-rater reliability, refine axis weights, and examine whether HACM-based reflection improves students’ understanding of responsible AI use. As generative AI tools continue to evolve, authorship frameworks must remain adaptable. HACM provides an initial scaffold for that work: practical enough for immediate educational use, but structured enough to support future empirical refinement.

7. AUTHOR AI USE DISCLOSURE

Generative AI tools were used to support the development of this manuscript, including brainstorming, outlining, drafting selected passages, revising prose, condensing sections for conference length, and refining examples. The author originated the central research problem, conceptual framework, six-axis model, scoring logic, authorship classifications, case structure, interpretations, and final editorial decisions. The author reviewed, revised, and approved all AI-assisted content and assumes full responsibility for the accuracy, citations, claims, and conclusions presented in the paper.

Using the Human-AI Contribution Model, this manuscript is classified as **Human-Author, AI-Assisted**. The estimated axis scores are: Source Origination = 4, Directive Specificity = 3, Iterative Steering = 4, Human Text Inclusion = 3, Transform Depth = 3, and Attribution & Judgment = 4. Applying the HACM formula produces the following Human Authorship Score:

$$\begin{aligned} HAS &= 25 \times (0.25 \cdot 4 + 0.15 \cdot 3 + 0.20 \cdot 4 + 0.20 \cdot 3 + 0.10 \cdot 3 + 0.10 \cdot 4) \\ HAS &= 25 \times (1.00 + 0.45 + 0.80 + 0.60 + 0.30 + 0.40) \\ HAS &= 25 \times 3.55 = \mathbf{88.75}, \text{ rounded to } \mathbf{89}. \end{aligned}$$

Accordingly, the manuscript is classified as **Human-Author, AI-Assisted (HAS $\approx$ 89)**.

8. REFERENCES

- Association for Computing Machinery. (n.d.). Frequently asked questions: ACM policy on authorship. <https://www.acm.org/publications/policies/frequently-asked-questions>
- Committee on Publication Ethics Council. (2023). COPE position: Authorship and AI. <https://doi.org/10.24318/cCVRZBms>
- Geburu, T., Morgenstern, J., Vecchione, B., Vaughan, J. W., Wallach, H., Daumé III, H., & Crawford, K. (2021). Datasheets for datasets. *Communications of the ACM*, 64(12), 86–92. <https://doi.org/10.1145/3458723>
- Gimpel, H., Hall, K., Decker, S., Eymann, T., Gutheil, N., Lämmermann, L., Braig, N., Maedche, A., Röglinger, M., Ruiner, C., Schoch, M., Schoop, M., Urbach, N., & Vandirk, S. (2025). Using

- generative AI in higher education: A guide for instructors. *Journal of Information Systems Education*, 36(3), 237–256. <https://doi.org/10.62273/QLLG7172>
- Gursoy, F., & Kakadiaris, I. A. (2022). System cards for AI-based decision-making for public policy. *arXiv*. <https://doi.org/10.48550/arXiv.2203.04754>
- Institute of Electrical and Electronics Engineers. (2024). Author guidelines for artificial intelligence (AI)-generated text. <https://open.ieee.org/author-guidelines-for-artificial-intelligence-ai-generated-text/>
- Lending, D., Ezell, J. D., Dillon, T. W., & May, J. (2022). A rubric to evaluate and enhance requirements elicitation interviewing skills. *Journal of Information Systems Education*, 33(4), 371–387. <https://doi.org/10.62273/ytnl5269>
- McHenry, W. K. (2022). Teaching tip: Evaluating visualizations with a compact rubric. *Journal of Information Systems Education*, 33(4), 324–337. <https://doi.org/10.62273/pdlh3512>
- Mitchell, M., Wu, S., Zaldivar, A., Barnes, P., Vasserman, L., Hutchinson, B., Spitzer, E., Raji, I. D., & Gebru, T. (2019). Model cards for model reporting. *Proceedings of the Conference on Fairness, Accountability, and Transparency*, 220–229. <https://doi.org/10.1145/3287560.3287596>
- National Information Standards Organization. (2022). ANSI/NISO Z39.104-2022, CRediT, Contributor Roles Taxonomy. <https://www.niso.org/publications/z39104-2022-credit>
- National Institute of Standards and Technology. (2023). Artificial intelligence risk management framework (AI RMF 1.0). <https://doi.org/10.6028/NIST.AI.100-1>
- Parasuraman, R., Sheridan, T. B., & Wickens, C. D. (2000). A model for types and levels of human interaction with automation. *IEEE Transactions on Systems, Man, and Cybernetics – Part A: Systems and Humans*, 30(3), 286–297. <https://doi.org/10.1109/3468.844354>
- Paullet, K., Pinchot, J. L., Kinney, S. T., & Stewart, C. (2025). Detecting AI-Generated Writing Using GPTZero. *Information Systems Education Journal*, 23(6), 44–58. <https://doi.org/10.62273/PZWW7741>
- U.S. Copyright Office. (2023). Copyright registration guidance: Works containing material generated by artificial intelligence. <https://www.copyright.gov/ai/>
- Zhang, X. (2025). Teaching tip: Incorporating AI tools into database classes. *Journal of Information Systems Education*, 36(1), 37–52. <https://doi.org/10.62273/GKZI2477>

APPENDIX A

Detailed HACM Scoring Anchors

The following scoring anchors provide additional guidance for applying the Human–AI Contribution Model (HACM). They are intended to supplement the summary rubric in Table 1 by clarifying how each score from 0 to 4 should be interpreted. These anchors may be used by students, instructors, reviewers, or researchers when assigning axis scores.

A.1 Source Origination (SO): Where core ideas come from

Score 0: The AI system originates the topic, angle, structure, or primary claims with little or no meaningful human direction. The human accepts the direction supplied by the AI.

Score 1: The human selects a broad topic or domain, but the AI determines the angle, framing, organization, or argument.

Score 2: The human provides both a topic and a general direction, such as a comparison, position, or intended focus, but does not provide a detailed structure or claim sequence.

Score 3: The human supplies a partial outline, key claims, examples, or an argument plan that meaningfully frames the work.

Score 4: The human independently originates the full conceptual structure, including the topic, thesis or purpose, argument sequence, evidence frame, and intended contribution before involving AI.

A.2 Directive Specificity (DS): Precision of instructions provided to AI

Score 0: Prompts are vague or open-ended, giving AI broad discretion over content, organization, and emphasis.

Score 1: The human provides only basic parameters, such as approximate length, tone, or general topic.

Score 2: Prompts include moderate scaffolding, such as named sections, required points, examples to include, or limited formatting expectations.

Score 3: The human provides detailed instructions that guide organization, rhetorical goals, evidence use, style, or definitional boundaries.

Score 4: The human provides a fully structured brief or assignment-style prompt with explicit constraints, section-by-section expectations, required content, and clear boundaries on what AI may and may not change.

A.3 Iterative Steering (IS): Human revision and control across iterations

Score 0: The AI output is accepted after a single prompt with no meaningful review, correction, or further direction.

Score 1: The human makes minor reactive edits, usually at the sentence or formatting level, but provides little conceptual guidance.

Score 2: The human gives targeted feedback across one or two iterations, correcting discrete issues or adjusting emphasis.

Score 3: The human engages in multiple revision cycles, directing structural changes, improving claims, clarifying reasoning, or correcting AI misinterpretations.

Score 4: The human actively leads sustained iterative development by challenging AI assumptions, restructuring arguments, rejecting unsuitable suggestions, refining reasoning, and ensuring alignment with the intended purpose.

A.4 Human Text Inclusion (HT): Proportion of human-authored text in the final artifact

Score 0: The final artifact contains no meaningful human-authored prose; it is entirely or almost entirely AI-generated.

Score 1: Less than 25% of the final text is human-authored, often limited to short insertions, minor edits, or isolated phrases.

Score 2: Between 25% and 50% of the final text is human-authored. Human and AI contributions are substantially intertwined.

Score 3: Between 50% and 75% of the final text is human-authored, with AI used for selected expansions, rewrites, or refinements.

Score 4: More than 75% of the final text is human-authored. AI is used only in a limited editorial, proofreading, formatting, or support role.

A.5 Transform Depth (TD): Degree of AI rewriting or alteration

Score 0: AI performs deep structural rewriting that substantially changes the meaning, organization, logic, or argumentation of the human's original input.

Score 1: AI performs extensive rewriting or restructuring while preserving only the broad idea or general intent of the human's input.

Score 2: AI performs moderate rewriting, replacing or reorganizing some sentences or paragraphs while leaving the underlying reasoning partly intact.

Score 3: AI performs light revision, such as smoothing phrasing, improving clarity, adjusting tone, or correcting awkward wording without materially changing the argument.

Score 4: AI changes are strictly surface-level, such as grammar, spelling, mechanics, formatting, or proofreading, with no substantive rewriting.

A.6 Attribution & Judgment (AJ): Human responsibility for factual and interpretive integrity

Score 0: The human accepts AI-generated facts, citations, claims, interpretations, or recommendations without checking their accuracy.

Score 1: The human performs limited spot-checking, but most facts, references, examples, or claims remain unverified.

Score 2: The human verifies selected claims or citations but does not systematically review the entire artifact.

Score 3: The human reviews most factual claims, corrects inaccurate or hallucinated references, and checks the logic and evidentiary support of the artifact.

Score 4: The human fully verifies factual claims, citations, data, reasoning, and interpretations and assumes complete responsibility for the final artifact.

APPENDIX B

HACM Student Reflection Template

This template may be used by instructors who want students to document and reflect on their use of generative AI in an assignment. Students should complete the template after finishing their work and submit it with the final artifact when required by the instructor.

Assignment Information

Student Name: _____

Course: _____

Assignment: _____

Date: _____

AI tool(s) used, if any: _____

Description of AI Use

Briefly describe how AI was used in this assignment. Include the purpose of use, the stage of the assignment when AI was used, and whether AI contributed to brainstorming, outlining, drafting, revising, proofreading, citation checking, coding, analysis, or another task.

Description:

HACM Axis Scores

For each axis, assign a score from 0 to 4 using the HACM rubric and scoring anchors. Then provide a brief justification for the score.

Axis	Score 0–4	Brief Justification
Source Origination (SO)	_____	_____
Directive Specificity (DS)	_____	_____
Iterative Steering (IS)	_____	_____
Human Text Inclusion (HT)	_____	_____
Transform Depth (TD)	_____	_____
Attribution & Judgment (AJ)	_____	_____

Human Authorship Score Calculation

Use the HACM formula below to calculate the Human Authorship Score.

$$\text{HAS} = 25 \times (0.25\text{SO} + 0.15\text{DS} + 0.20\text{IS} + 0.20\text{HT} + 0.10\text{TD} + 0.10\text{AJ})$$

Substitute your scores into the formula:

$$\text{HAS} = 25 \times (0.25 \times \underline{\hspace{1cm}} + 0.15 \times \underline{\hspace{1cm}} + 0.20 \times \underline{\hspace{1cm}} + 0.20 \times \underline{\hspace{1cm}} + 0.10 \times \underline{\hspace{1cm}} + 0.10 \times \underline{\hspace{1cm}})$$

Calculated HAS: _____

HACM Classification

Use the Human Authorship Score to identify the appropriate HACM classification.

HAS Range Classification

0–24	AI-Created, Human-Curated
25–49	AI-Led Co-Creation
50–74	Human-Led Co-Creation
75–100	Human-Author, AI-Assisted

My HACM classification: _____

Student Disclosure Statement

Write a short disclosure statement explaining how AI was used in this assignment. Your statement should be specific enough that the instructor can understand the role AI played in your work.

Disclosure Statement:

Student Certification

By signing below, I certify that this reflection accurately represents my use of AI in completing this assignment and that I remain responsible for the final submitted work.

Student Signature: _____

Date: _____

APPENDIX C

Sample AI Use Disclosure Statements

The following sample disclosure statements illustrate how writers might describe AI use under each HACM classification. These statements are intended as examples and should be adapted to fit the assignment, course policy, publication venue, or institutional requirements.

C.1 AI-Created, Human-Curated

Sample Disclosure:

Generative AI was used to generate most or all of the content in this submission. My role was limited primarily to selecting the topic, reviewing the output, and making minor edits. I did not substantially originate the argument, structure, or wording of the final artifact. This submission is classified as AI-Created, Human-Curated under HACM.

Appropriate Use Contexts:

This level of AI use may be appropriate for brainstorming, informal exploration, low-stakes practice, or demonstrations of AI capabilities. It is generally not appropriate for assignments or publications requiring original human authorship unless explicitly permitted.

C.2 AI-Led Co-Creation

Sample Disclosure:

Generative AI was used to produce a substantial portion of the draft, including organization, phrasing, and content development. I provided the topic, general direction, and selected materials, and I made limited revisions to the AI-generated output. This submission is classified as AI-Led Co-Creation under HACM.

Appropriate Use Contexts:

This level of AI use may be appropriate for early-stage drafts, formative exercises, background exploration, or assignments where AI-assisted drafting is explicitly allowed. It may not be appropriate for final submissions intended to assess independent reasoning, writing ability, or mastery of course concepts.

C.3 Human-Led Co-Creation

Sample Disclosure:

Generative AI was used to assist with drafting, organization, and revision. I developed the central ideas, argument structure, examples, and overall direction of the work, and I reviewed and revised the AI-assisted content before submission. This submission is classified as Human-Led Co-Creation under HACM.

Appropriate Use Contexts:

This level of AI use may be appropriate when instructors or publication venues allow AI-assisted drafting and revision, provided the human author retains responsibility for the ideas, organization, accuracy, and final presentation of the work.

C.4 Human-Author, AI-Assisted

Sample Disclosure:

Generative AI was used for limited assistance such as proofreading, grammar correction, clarity improvements, formatting, or minor wording suggestions. I originated the ideas, structure, analysis, and substantive wording of the work, reviewed all AI suggestions, and remain fully responsible for the final submission. This submission is classified as Human-Author, AI-Assisted under HACM.

Appropriate Use Contexts:

This level of AI use is generally consistent with human authorship when AI functions as an editorial or support tool. It may be appropriate for formal assignments, manuscripts, reports, or professional writing when permitted by the relevant policy.