



Fostering AI Literacy with LuminAI through Embodiment and Creativity in Informal Learning Spaces

Chengzhi Zhang
czhang694@gatech.edu
Georgia Institute of Technology
Atlanta, Georgia, USA

Chelsi Cocking
chelsi@gatech.edu
Georgia Institute of Technology
Atlanta, Georgia, USA

Milka Trajkova
mtrajkova3@gatech.edu
Georgia Institute of Technology
Atlanta, Georgia, USA

Zoe Mock
zoe@gatech.edu
Georgia Institute of Technology
Atlanta, Georgia, USA

Gemma Tate
wtate34@gatech.edu
Georgia Institute of Technology
Atlanta, Georgia, USA

Cassandra Naomi Monden
cat.monden@gatech.edu
Georgia Institute of Technology
Atlanta, Georgia, USA

Brian Magerko
magerko@gatech.edu
Georgia Institute of Technology
Atlanta, Georgia, USA

ABSTRACT

LuminAI is an interactive art installation that allows participants to collaborate with an AI dance partner by improvising movements. During the interaction, the participant dances with an AI dance partner who learns from the participant's movements in real-time and remixes them into new, unexpected forms of motion. LuminAI blurs the lines between the participant and AI agent, inviting the participant to improvise and engage with AI in a dance of creativity. Recently, we've redesigned LuminAI with the educational purpose of enhancing the public's AI literacy. We separated LuminAI into three panels with AI-related educational goals for each panel and redesigned the user interaction. In this paper, we present the redesign of LuminAI and educational goals centering on enhancing the public's AI literacy.

CCS CONCEPTS

• Human-centered computing → Interaction design; • Computing methodologies → Artificial intelligence.

KEYWORDS

embodied interaction, AI literacy, AI agent, informal education

ACM Reference Format:

Chengzhi Zhang, Chelsi Cocking, Milka Trajkova, Zoe Mock, Gemma Tate, Cassandra Naomi Monden, and Brian Magerko. 2024. Fostering AI Literacy with LuminAI through Embodiment and Creativity in Informal Learning Spaces. In *Creativity and Cognition (C&C '24)*, June 23–26, 2024, Chicago, IL, USA. ACM, New York, NY, USA, 6 pages. <https://doi.org/10.1145/3635636.3664255>

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).
C&C '24, June 23–26, 2024, Chicago, IL, USA
© 2024 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-0485-7/24/06
<https://doi.org/10.1145/3635636.3664255>

1 INTRODUCTION

LuminAI (previously called Viewpoints AI [5]) is an interactive installation in which participants can improvise movements with an AI dance partner. Due to the rich interaction it supports, we have explored multiple research directions with LuminAI [10], including exploring how to design socially interactive creative systems involving multiple humans and AI agents [8], investigating proper ways of turn-taking while improvising with LuminAI dance agent [25], etc. We varied the design trajectory due to the varying research questions we asked. This includes changing from five-module agent design [21] to three panels design presented in this paper; varying the number of participants in the interaction processes from multiple participants [8] to single participant [25]; varying the scale of the installation, etc. Besides co-improvisation interaction design, we also have a corresponding data visualization panel named Moviz for presenting its learned movements in different colored clusters [7].

With AI becoming more prevalent in our lives, the need to enhance the public's AI literacy becomes essential. Thus, we are applying Duri and Magerko's work in defining AI literacy—in terms of critical learning objectives and design principles – as a framework for redesigning this embodied, creative AI experience for informal learning environments [11]. The AI literacy metrics listed in this paper have served as guidelines for many projects promoting AI literacy [6, 22].

2 RELATED WORK

2.1 Embodied interaction for promoting AI literacy

Interacting with LuminAI greatly emphasizes embodied interaction as it requires a full-body interaction. Embodied and tangible interaction has proven to be efficient in STEM education, including Computer Science [3], Mathematics [1], etc. Recent research further demonstrated the advantage of utilizing embodied interaction for promoting AI literacy [26]. As museums are important public spaces for informal learning and the site for deploying our installation,

Panel	Previous Learning Objectives	Revised Learning Objectives
1. Representation	1. Understand that a body can be tracked by a computer.	1. Understand movement data can be represented by AI. 2. Understand movement data can be categorized by AI.
2. Co-creation	1. Computers generate concepts based on the data set it is given. 2. Understand AI can co-create with humans in novel ways using movement. 3. Critically reflect what it means to co-create with AI.	1. Understand AI can learn movement from a person. 2. Understand AI can make a decision based on learned data.
3. Grouping	1. Understand the importance of grouping data and why it is useful in ML. 2. Computers reason about data sets through clustering (to determine similarity). 3. Computers can learn data with or without supervision & the benefit therein.	1. Understand AI can group similar motion data.

Table 1: Learning objectives for LuminAI: each panel breakdown.

Hornecker's book on HCI in the museum has informed our design and evaluation [4]. Furthermore, seventeen design principles for co-creating with AI in public spaces [9] have greatly informed our design practice.

2.2 Projects for promoting AI literacy in informal learning setting

Projects, activities, and exhibits with the goal of promoting AI literacy in the informal learning setting have increasingly emerged. Project-wise, *Creature Features* is a project in which the user can use the training dataset for a feature-based ML bird classification algorithm [14]. *Knowledge Net* is a tangible interface of tiles and arrows to collaboratively build semantic networks [14]. *Doodlebot* is a mobile social robot that promotes grade school (K-12) students' understanding of AI concepts [24]. Projects that require minimal technical requirements like *AI education unplugged* [13] were designed to suit various learning contexts. Activity-wise, Druga et al designed 11 AI literacy activities [2] engaging parent-child partners. Exhibit-wise, the Museum of Science, Boston presented a special exhibit featuring AI [15]. Further, the Misalignment Museum in San Francisco specifically features AI education and learners' critical reflection upon AI [16]. These projects, activities and exhibits directly informed our design and helped us navigate our learning goal selection.

3 LUMINAI REDESIGN

In the redesign of LuminAI, we have divided the holistic interactive experience into three-panel designs. The division helps us better specify the learning objectives in detail, which directs the interaction design that follows. It also helps us modularize the installation of the project and evaluate learning outcomes after participants' interaction. Among the three panels, Panel 2 serves as the place where most dancing and improvising activities happen, and it is physically located in the middle of the exhibit. Panels 1 and 3 help explain the AI's thinking and reasoning processes behind the scene.

3.1 Learning objectives, revised

While discussing the learning objectives together with our museum partner Museum of Science and Industry (MSI), Chicago, we realized our learning objectives were overlaid with scientific language and jargon, which museum designers suggested are not suitable for middle school-aged learners or other learners who do not have any AI-related background. To solve this, we revised our learning objectives by reducing reading levels. For instance, to replace the word "cluster" (which is a Machine Learning term for unsupervised learning) in Panel 3's learning goal, we used "group" instead. We iterated through the learning objectives and decided on the revised learning objectives as listed in table 1. After that, we started the interaction design processes.

3.2 User population aim

The redesign is intended for middle school-aged children (9-14 years old), informed by work indicating the difficulties that younger children have in understanding AI due to their still-developing theory of mind [23]. We choose the middle school age for AI introduction as it is a critical period when students start to consider their future paths [19]. With the broader goal of enhancing AI literacy across the public, special attention is given to groups less represented in AI and computer science—specifically aiming to include middle schoolers with no computer science background, girls in this age range, and students from Title 1 schools¹. Our project seeks to broaden AI understanding among museum visitors, not requiring or asking about their past AI interactions.

3.3 LuminAI panel redesign

3.3.1 Panel 1. The finalized learning objective of panel 1 is to *understand how raw movement data can be represented by AI* (See table

¹Title I schools are schools that receive federal funding to help low-income students achieve academically.

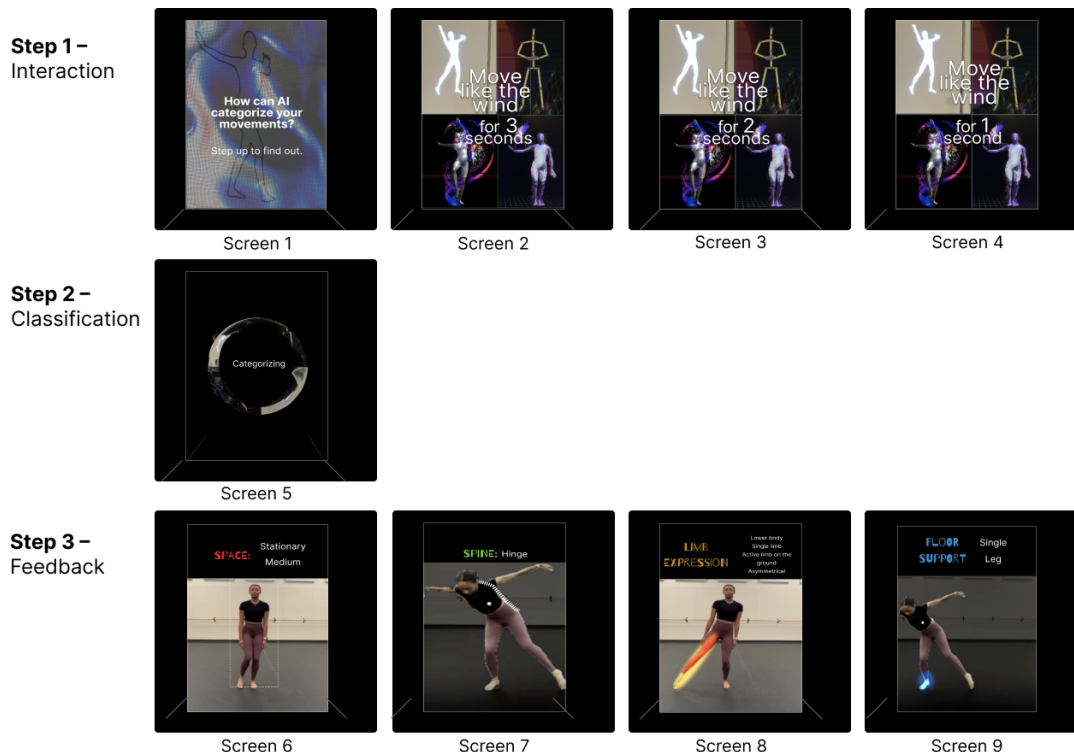


Figure 1: LuminAI Panel 1 design(from left to right, top to bottom): Step 1-Interaction: Participants are greeted by a question (Screen 1). A countdown with a dance prompt appears (Screen 2-4). Step 2-Classification (Screen 5): Feedback on system status appears, informing participants their movement is being categorized. Step 3-Feedback: The classification results appear with the first category of Space (Screen 6). The second category of spine (Screen 7). The third category of limb expression (Screen 8). The final category of floor support (Screen 9). All the categories have a text-based output and visual feedback on the avatar.

1). The foundational underpinning resides in helping learners ² understand how we can use computer vision to track their movements and how we can use AI, specifically supervised machine learning, to classify participants' movements. The user experience of this panel as shown in Fig. 1, begins with step 1- the interaction mode. Learners are met with a question (Screen 1) to entice middle schoolers - "How does AI categorize your movements? Step up to find out." We specifically chose to use the terminology of "categorize" rather than "classify" to eliminate the use of jargon for the comprehension level of middle schoolers. After the learner steps closer to the screen, the next screen appears that details a short dance prompt, e.g., "Move like the wind" (Screen 2-4). A series of back-end prompts will give structured guidance on how to move rather than relying on open-ended improvisation [21]. Learner will dance for three seconds which subsequently will greet them with feedback of system status indicating that their movements are being categorized (Step 2, Screen 5). The system will categorize learners' movements according to the Body Action Taxonomy (BAT) that we designed internally [21]. This taxonomy provides a structured vocabulary designed to more objectively describe observable body actions by using a standardized language. BAT has four overarching categories used to describe movements: 1. Space - representing which vertical

level a movement lies; 2. Location of Spine - representing the position and alignment of the spine; 3. Limb Expression - representing the intentional movements created by upper and lower limbs; and 4. Floor Support - representing the primary segment responsible for bearing weight. Generally, what body part is touching the floor. Each category has different levels (description of these is outside the scope of the paper). The learner will see their three-second movement broken down into these four categories. The final step is feedback. The next screens (6-9) will present the feedback associated with each category and code in the taxonomy. Feedback will be represented both in texts and also visually embodied on the human. By using multimodal feedback we can ensure greater appendability as well as accessibility [20]. Space will be visually represented by a dotted box around three levels of the body. Spine visual feedback will be focused on the spine itself, demonstrating the curve. Limb expression is visually highlighted by the body parts used. Finally, the last category of floor support will show a visual highlight of the body parts touching the floor. The prototype can be viewed here.

3.3.2 Panel 2. The latest version of LuminAI is an immersive holographic public art installation (inspired by the concept of a Pepper's ghost) that allows audience members and dancers to engage in this co-creative AI dance experience [21]. However, an initial evaluation

²We used the term "learners" to encompass museum-goers, participants, and users.

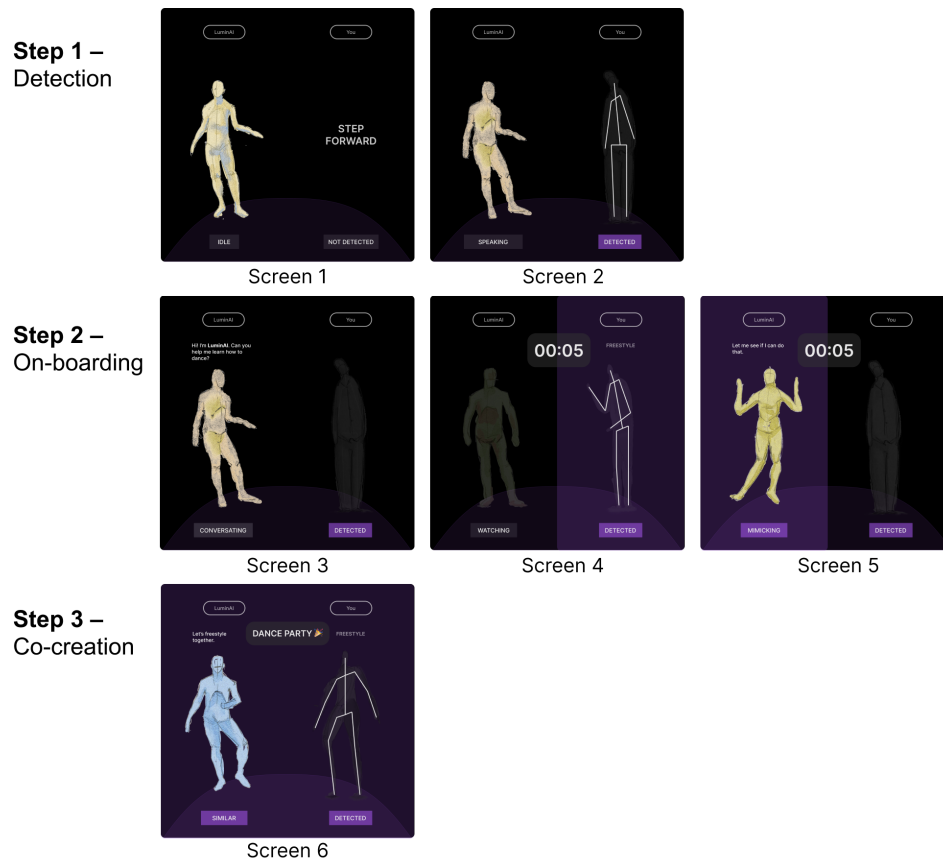


Figure 2: LuminAI Panel 2 design (from left to right, top to bottom): Step 1 - Detection: the user is prompted to step forward into the interaction zone to be detected by the Kinect sensor (Screen 1-2). Step 2 - Onboarding: the user goes through two rounds of a 'turn-taking' dance experience to learn how to collaborate with LuminAI (Screen 3-5). Step 3 - Co-creation : the user is free to collaboratively and concurrently freestyle and co-create with LuminAI through improvisational dance (Screen 6)

of the installation indicated a lack of affordance to cue public audience participation, and the audience typically incorrectly perceives the AI agent's capabilities [21].

Based on those empirical findings, the goal of the Panel 2 redesign is to 1. solve existing usability pitfalls and 2. serve the learning objectives outlined in Table 1 based on the co-creative AI for public spaces design principles, outlined by Long et al. [9]— especially the goals of apprehendability, learnability, creativity, and engagement.

In the latest version of the installation [21], the high-level two-step user flow is: having the user enters a curtained-off makeshift black box room, which serves as a stage/interaction zone, and immediately upon entrance, the user begins interacting with LuminAI—the AI agent, in concurrent collaborative freestyle dance. While this design aimed to achieve an intuitive interaction to start the experience, it inadvertently discouraged public participation. The makeshift spotlighted black-box stage deterred novice users from publicly engaging in dance-based, embodied interactions and separated the interacting participant from the wider public audience, causing a lack of social engagement around the experience [21]. Furthermore, upon beginning the collaborative dance interaction, the user needed help understanding how the AI agent responded

to or reacted to them. Many users wrongly thought the agent was responding with pre-authored movements. In fact, they were interacting with an AI agent that actively improvises dance moves onstage live [21].

The redesign of Panel 2 utilizes a 'turn-taking' tutorial-like experience in the middle of the user flow before launching the user into the co-creative AI dance collaboration experience. This serves as an 'on-boarding' phase. (see Fig. 2) Adding this step to the overall experience helps intuitively onboard the participants into the experience of interacting with LuminAI. By first interactively educating them on how LuminAI learns dance moves from them, responds to them with new dance moves, and collaborates with them, we are creating more intuitive interaction as suggested in Design Principle 11[9].

The new user flow includes: A panel screen installed in a public setting has an ambient visual of LuminAI, the AI agent, idly swaying and marking movements with a prompt "step forward" on the opposite side to attract participants. Once a participant steps into the dedicated interaction zone—a zone in an open space rather than an enclosed room, monitored by a Kinect sensor—their abstract skeletal representation appears on the opposite side of the

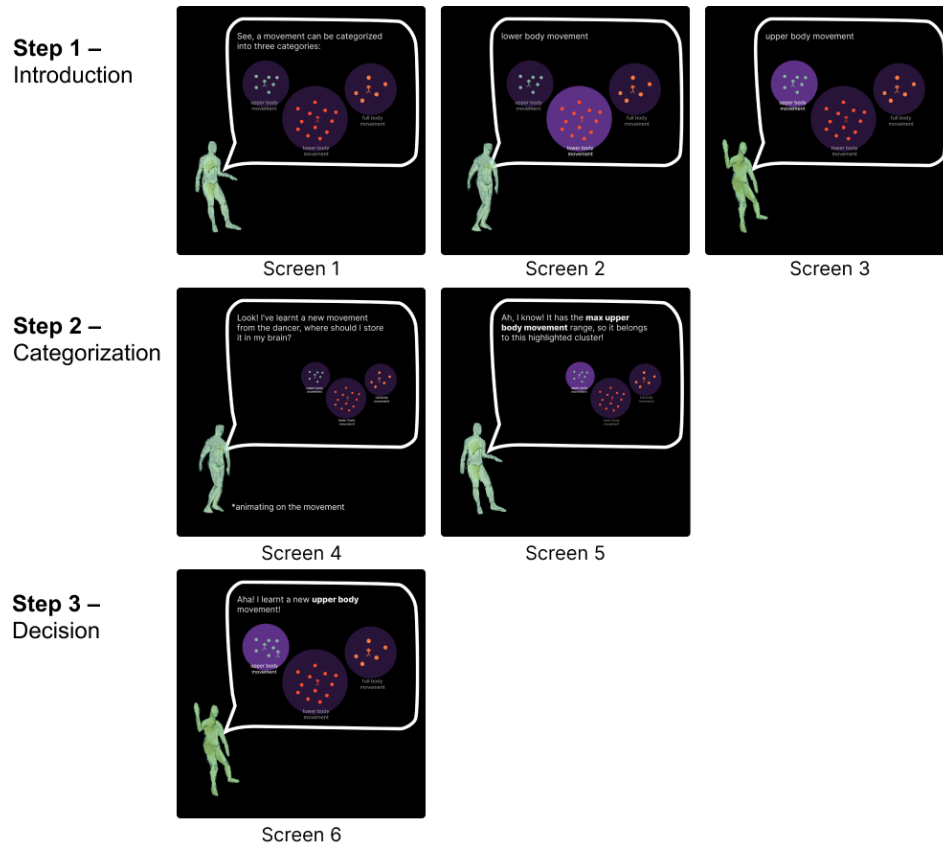


Figure 3: LuminAI Panel 3 design (from left to right, top to bottom): Step 1-Introduction: the agent introduces three clusters by highlighting each cluster and its name (Screen 1-3). Step 2-Categorization: the agent categorizes one movement and gives the reason (Screen 4-5). Step 3-Decision: the agent decides where to store a movement (Screen 6).

AI agent, showing the position of their body in space. LuminAI then introduces itself and prompts the user to help it learn how to dance by teaching it a dance move: "Hi! I am LuminAI. Can you help me learn how to dance? ... Teach me your best dance move! Ready?" The system then presents the user with a timed five-second experience encouraging them to demonstrate a dance move. LuminAI is animated to make it seem like it is actively watching the user dance during the timed experience. At the end of the 5-second timer, LuminAI responds: "Let me see if I can do that", and then executes the dance the participant just demonstrated by mimicking their exact movements. The system then repeats this tutorial experience once more, with LuminAI prompting the participant to try a more challenging, more elaborate dance movement: "Let us go again! Give me some harder dance moves. Ready?" After another timed round of teaching by the participant, LuminAI responds this time by saying: "I will dance something different," and using its AI capabilities to respond with dance movements that contrast with what the participant just taught. After this second, slightly more advanced tutorial, the participant is dropped into the original collaborative mode, which is now step 3 of the experience, where they can co-creatively collaborate with the AI agent through constant ongoing improvisational freestyle dance (see Fig. 2).

3.3.3 Panel 3. Panel 3 has a significant subsequent correlation with Panel 2 (See Fig. 3). After the participant dances with the AI agent, the AI agent will gain new movements, and the new movements will appear in the Moviz movements visualization panel. In the design, we stressed the movement **categorization** processes to meet the learning goal: *Understand that AI can group similar motion data*. We had an AI agent at the bottom left corner of the screen, which explained the whole categorization process of new data points (movement).

Previously, in Moviz [7], we used keyframes to cluster movements in K categories. However, these clusters do not bear human-understandable characteristics. Specifically, there is no clear semantic meaning to each cluster. To solve this issue, we decided on having three clusters labeled "upper body movement", "lower body movement" and "whole body movement". The name of the cluster is self-explanatory. For instance, in the "upper body movement" cluster, the movement has the most noticeable movement in the upper body. In this way, the agent can not only reason about their selection but also **explain** that to the participants. To reduce the visual complexity, we only show the center cluster movement, while the other movements in the cluster are shown in dots.

As shown in screens 1-3, in a default state, the agent first explains, "a movement can be categorized into three categories" and goes through each cluster by their labels. When a movement is learned in Panel 2, the agent says "Look! I've learned a new movement from the dance; where should I store it in my brain?" The pondering moment is followed by an "aha" moment when the agent explains that due to the "max body movement" range it sees, it stores the movement in the upper body cluster, and it learns a new movement. Through active explanations, we aim to convey the learning goal that *Understand AI can group similar motion data*.

4 EVALUATION PLANS AND FUTURE WORK

Researchers use personal meaning maps [18] and learning talk [17] data collection with learners to evaluate the learning outcomes and interaction design. However, given the fact that we want to assess a learner's knowledge gain for a specific learning objective, we choose to adopt a questionnaire for evaluating the learner's subjective knowledge of AI before and after their interaction. Secondly, we plan to adopt a questionnaire for evaluating content knowledge acquisition (CKA) that mainly consists of open-ended questions, as well as a Likert scale for evaluating their future learning interests in AI. Lastly, we plan to adopt the APEX [12] framework for on-site evaluation. The evaluation goal is to both inform our future design iteration and also to evaluate the knowledge gained after interacting with each panel.

REFERENCES

- [1] Arthur Bakker and Dor Abrahamson. 2016. Making sense of movement in embodied design for mathematics learning. *Cognitive Research: Principles and Implications* 1, 1 (2016), 0–0.
- [2] Stefania Druga, Fee Lia Christoph, and Amy J Ko. 2022. Family as a Third Space for AI Literacies: How do children and parents learn about AI together?. In *Proceedings of the 2022 CHI conference on human factors in computing systems*. 1–17.
- [3] Michael S Horn, Erin Treacy Solovey, R Jordan Crouser, and Robert JK Jacob. 2009. Comparing the use of tangible and graphical programming languages for informal science education. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. 975–984.
- [4] Eva Hornecker and Luigina Ciolfi. 2022. *Human-computer interactions in museums*. Springer Nature.
- [5] Mikhail Jacob and Brian Magerko. 2015. Viewpoints ai. In *Proceedings of the 2015 ACM SIGCHI Conference on Creativity and Cognition*. 361–362.
- [6] Phoebe Lin and Jessica Van Brummelen. 2021. Engaging teachers to co-design integrated AI curriculum for K-12 classrooms. In *Proceedings of the 2021 CHI conference on human factors in computing systems*. 1–12.
- [7] Lucas Liu, Duri Long, and Brian Magerko. 2020. Moviz: A visualization tool for comparing motion capture data clustering algorithms. In *Proceedings of the 7th International Conference on Movement and Computing*. 1–8.
- [8] Duri Long, Mikhail Jacob, Nicholas Davis, and Brian Magerko. 2017. Designing for socially interactive systems. In *Proceedings of the 2017 ACM SIGCHI Conference on Creativity and Cognition*. 39–50.
- [9] Duri Long, Mikhail Jacob, and Brian Magerko. 2019. Designing co-creative AI for public spaces. In *Proceedings of the 2019 Conference on Creativity and Cognition*. 271–284.
- [10] Duri Long, Lucas Liu, Swar Gujrana, Cassandra Naomi, and Brian Magerko. 2020. Visualizing improvisation in luminai, an ai partner for co-creative dance. In *Proceedings of the 7th International Conference on Movement and Computing*. 1–2.
- [11] Duri Long and Brian Magerko. 2020. What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI conference on human factors in computing systems*. 1–16.
- [12] Duri Long, Tom McKlin, Nylah Akua Adjei Boone, Dezarae Dean, Mirina Garoufalidis, and Brian Magerko. 2022. Active Prolonged Engagement EXpanded (APEX): A Toolkit for Supporting Evidence-Based Iterative Design Decisions for Collaborative, Embodied Museum Exhibits. *Proceedings of the ACM on Human-Computer Interaction* 6, CSCW1 (2022), 1–33.
- [13] Duri Long, Jonathan Moon, and Brian Magerko. 2021. Unplugged assignments for K-12 AI education. *AI Matters* 7, 1 (2021), 10–12.
- [14] Duri Long, Sophie Rollins, Jasmin Ali-Diaz, Katherine Hancock, Samnang Nuonsinoen, Jessica Roberts, and Brian Magerko. 2023. Fostering AI Literacy with Embodiment & Creativity: From Activity Boxes to Museum Exhibits. In *Proceedings of the 22nd Annual ACM Interaction Design and Children Conference*. 727–731.
- [15] Boston Museum of Science. 2022. Museum of Science, Boston Reveals How Artificial Intelligence is Everywhere in New Exhibition, Exploring AI: Making the Invisible Visible Museum of Science. <https://www.mos.org/press/press-releases/Exploring-AI-Making-the-Invisible-Visible-Exhibit-Announcement> Accessed on April, 2024.
- [16] noauthor. 2022. The Misalignment Museum 501(c)3 is a place to learn about Artificial Intelligence and reflect on the possibilities of technology through thought-provoking art pieces and events. <https://misalignmentmuseum.com/> Accessed on April, 2024.
- [17] Jessica Roberts and Leilah Lyons. 2017. The value of learning talk: applying a novel dialogue scoring method to inform interaction design in an open-ended, embodied museum exhibit. *International Journal of Computer-Supported Collaborative Learning* 12 (2017), 343–376.
- [18] Won-Joo Suh. 2010. Personal meaning mapping (PMM): A qualitative research method for museum education. *Journal of Museum Education* 4 (2010), 61–82.
- [19] Rivka Taub, Michal Armoni, and Mordechai Ben-Ari. 2012. CS unplugged and middle-school students' views, attitudes, and intentions regarding CS. *ACM Transactions on Computing Education (TOCE)* 12, 2 (2012), 1–29.
- [20] Milka Trajkova and Francesco Cafaro. 2018. Takes Tutu to ballet: designing visual and verbal feedback for augmented mirrors. *Proceedings of the ACM on Interactive, Mobile, Wearable and Ubiquitous Technologies* 2, 1 (2018), 1–30.
- [21] Milka Trajkova, Manoj Deshpande, Andrea Knowlton, Cassandra Monden, Duri Long, and Brian Magerko. 2023. AI Meets Holographic Pepper's Ghost: A Co-Creative Public Dance Experience. In *Companion Publication of the 2023 ACM Designing Interactive Systems Conference*. 274–278.
- [22] Christiane Gresse Von Wangenheim, Livia S Marques, and Jean CR Hauck. 2020. Machine Learning for All—Introducing Machine Learning in K-12. (2020).
- [23] Henry M Wellman, David Cross, and Julianne Watson. 2001. Meta-analysis of theory-of-mind development: The truth about false belief. *Child development* 72, 3 (2001), 655–684.
- [24] Randi Williams, Safinah Ali, Raúl Alcantara, Tasneem Burghle, Sharifa Al-ghowinem, and Cynthia Breazeal. 2024. Doodlebot: An Educational Robot for Creativity and AI Literacy. In *Proceedings of the 2024 ACM/IEEE International Conference on Human-Robot Interaction*. 772–780.
- [25] Lauren Winston and Brian Magerko. 2017. Turn-taking with improvisational co-creative agents. In *Proceedings of the AAAI Conference on Artificial Intelligence and Interactive Digital Entertainment*, Vol. 13. 129–135.
- [26] Weipeng Yang, Xinyun Hu, Ibrahim H Yeter, Jiahong Su, Yuqin Yang, and John Chi-Kin Lee. 2023. Artificial intelligence education for young children: A case study of technology-enhanced embodied learning. *Journal of Computer Assisted Learning* (2023).