

Designing Movement Generation Models in Collaboration With Voguing And Dancehall Dancers

Léo Chédin

Laboratoire Interdisciplinaire des Sciences du Numérique
Université Paris-Saclay, CNRS, Laboratoire
Interdisciplinaire des Sciences du Numérique
Orsay, France
leo.chedin@universite-paris-saclay.fr

Baptiste Caramiaux

Sorbonne Université, CNRS, Institut des Systèmes
Intelligents et de Robotique, ISIR
Paris, France
baptiste.caramiaux@sorbonne-universite.fr

Jules Françoise

Université Paris-Saclay, CNRS, Laboratoire
Interdisciplinaire des Sciences du Numérique
Orsay, France
jules.francoise@lisn.fr

Sarah Fdili Alaoui

Creative Computing Institute
University of the Arts London
London, United Kingdom
s.fdilialaoui@arts.ac.uk

Abstract

Recent advances in Artificial Intelligence have enabled powerful generative models, yet few are tailored to dancers' practices. We present a long-term collaboration with a Voguing and Dancehall collective to design movement generation models trained on their repertoire. Our initial study with the dancers revealed that, despite limited physical realism, the generated movements inspired them. Iterative development led to Korai, an interactive tool for monitoring training, visualizing motion data, and prompting generation, which improved output quality. A subsequent structured observation study compared three model variants with high, medium, and low fidelity to the original dataset's style. Results show that dancers favored either highly faithful or highly unfaithful outputs, rejecting medium fidelity as neither authentic to their style nor creatively stimulating. Our findings highlight how direct collaboration with dancers not only informs model design but also deepens understanding of AI's role in supporting creative movement practices.

CCS Concepts

- **Human-centered computing** → **Empirical studies in interaction design**; *Empirical studies in collaborative and social computing*;
- **Applied computing** → *Performing arts*.

Keywords

Generative AI, Dance, Movement Generation, Participatory AI Design

ACM Reference Format:

Léo Chédin, Jules Françoise, Baptiste Caramiaux, and Sarah Fdili Alaoui. 2026. Designing Movement Generation Models in Collaboration With Voguing And Dancehall Dancers. In *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems (CHI '26)*, April 13–17, 2026, Barcelona,

Spain. ACM, New York, NY, USA, 16 pages. <https://doi.org/10.1145/3772318.3791515>

1 Introduction

Since the postmodern era of dance, choreographers have continuously sought new ways to approach and create dance movements. They often experiment with unconventional methods — such as drawing, chance operations, or everyday situations — to develop original motions. From its very beginnings, the computer has been viewed as a potential way of inventing new ways to create movement [12]. In his 1967 article, Michael Noll predicted that: “Different movements might be stored in the computer’s memory and put together at will. [...] Various elements of chance and randomness could be used at the discretion of the choreographer” [38]. Over the years, numerous generative systems have been developed and used in dance. In their article reviewing compositional tools for movement, Fdili Alaoui et al. defined generative tools for dance as creativity support tools (CSTs) that generate movement material either autonomously or manually, through the input of a choreographer [16].

In this context, the technical advances of Artificial Intelligence (AI) have significantly expanded the potential of models that generate dance movements [37]. These models leverage recent neural network architectures to generate dance movements in an increasingly autonomous way [7]. Despite their great capabilities, current research in Human-Computer Interaction (HCI) exhibits a striking deficit in how dance practitioners’ knowledge is included in the design, evaluation, and use of these movement generation models. Dancers build embodied connoisseurship of somatic, sensory, aesthetic, and physical dimensions of movements. Within their communities of practice, they acquire mastery of movements and choreographic thinking, which are situated in the context of specific dance styles and genres. They develop a rich spectrum of movements and repertoires, which are only partially represented, in the current development of dance movement generation models. Indeed, in most cases, the models are trained either on dance improvisations or on the few publicly available dance datasets. Although efforts have been made to produce motion capture datasets of dance across different genres [31], they remain too generic or limited to



This work is licensed under a Creative Commons Attribution 4.0 International License. *CHI '26, Barcelona, Spain*

© 2026 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-2278-3/26/04

<https://doi.org/10.1145/3772318.3791515>

movements from improvisations in the contemporary dance style. Studying the integration of dance connoisseurship from various styles and traditions into the design of AI-based CSTs is thus a research priority. Nevertheless, to date, no study has investigated how generative movement models might foster dancers' knowledge and creativity within a specifically well-defined repertoire.

In this paper, we present a case study of the design of several iterations of a system for dance generation in collaboration with dance practitioners from Voguing and Dancehall dance styles. The study is a close collaboration with two dancers from a dance collective based in Paris that are connoisseurs [44] of these dance styles. The collaboration sought to generate movements grounded in the collective's dance practice in order to support improvisation and choreographic creation within their dance repertoire. A central objective was for the AI to encode the collective's choreographic style faithfully. To that end, we engaged in the iterative design and development of an AI-based dance movement generation model and an interactive system to generate and visualize movement sequences with the model. Our study empirically assessed how dancers perceived and interacted with successive versions of the model, examining its suitability for their specific style and repertoire.

Our contribution is threefold. First, we contribute with the design of an AI-based dance movement generation model, made in collaboration with Voguing and Dancehall dancers and targeting their unique dance practice and repertoires. Second, we developed an interactive system, *Korai*, to facilitate the visualization of motion capture recordings and to allow for direct interaction with the generative model. Third, we contribute with empirical insights gathered throughout multiple studies that help us understand how dancers perceive, experience and interact with various iterations of generative models producing movements with various levels of faithfulness to the training dataset. We discuss these results in light of the different qualities of the generated movements. In addition to discussing the balance between novelty and physical realism in the generated movements, we also underline the importance of stylistic legibility, pointing to a possible uncanny valley in dance generation. We then argue for greater practitioners' participation in the design of AI systems for dance and discuss how software tool such as *Korai* could help mitigate the hardships of designing AI models.

2 Related Works

In this section, we outline the history of motion generation systems in HCI and dance and technology and the recent technical advances leveraging data-driven techniques. We then outline a gap in the development of AI-based systems, which lack grounding in practitioners' work.

2.1 Movement Generation Models for Dance

Over the years, various algorithmic approaches for movement generation have emerged, including genetic algorithms [28], segment concatenation [21], Hidden Markov Models, and Boltzmann Machines [2]. With the latest advances in AI and the increasing accessibility of motion capture technology, motion generation techniques have evolved significantly [1]. Chor-rnn [11] was the first movement generation model based on deep learning techniques to be

developed in 2016 for applications in dance and choreography. They used a recurrent neural network to estimate the parameters of a mixture density network. They trained their model on motion capture data of contemporary dance improvisation that was captured using a Kinect. The model could then create new dance material by generating the continuation of an existing sequence. Pettee et al. investigated the use of variational autoencoders to enable control of the generated movements by manipulating the latent space. They also trained their model using motion capture of their own improvisations in a contemporary dance style [40].

Driven in part by significant progress in movement generation in 3D animation, numerous AI systems for dance generation have since been developed, leveraging various input modalities [37]. Among these modalities, dance generation models often use music-conditioned dance generation [2, 31, 48]. Other models leverage text-to-motion, such as Human MDM [47] or MotionGPT [27], style transfer, as well as spatial or temporal editing (e.g., generating parts of the body or parts of the movements to create the transition between two existing motion sequences [42]).

Training this new generation of AI dance movement generation models requires large training datasets. While several datasets exist for animation [24, 41, 50], datasets of 3D dance movements remain scarce, with only a few publicly available ones. The most popular choice is AIST++, the largest dataset of music-aligned 3D dance movements introduced by Li et al. [31], which they built from the AIST multiview video dataset of dance [49]. It contains 310 minutes of dance movements from several genres, including Hip-Hop, Ballet, Jazz, House, Waack, and more.

While the systems in the literature currently offer impressive technical capabilities, their training often occurs with limited involvement from dance experts. Indeed, they are typically trained on general-purpose dance or animation datasets such as AIST++ and are not specifically tailored to dancers' practices and repertoires. Our work contrasts with these approaches as we involve dance practitioners in the design of our movement generation models, particularly in training the models based on their curated repertoire of Voguing and Dancehall.

2.2 Research on Dance Generation Systems In-the-World

Most research in dance movement generation is currently focused on the technical capabilities of the models themselves, with few studies investigating how practitioners interact with these technologies. Most studies introduce new models, which they typically evaluate based on a combination of quantitative metrics and crowd-sourced platforms such as Prolific, where participants are asked to select their preferred output [3, 22, 31, 48]. While these evaluations are a necessary first step to assess the model's capabilities, there is a strong need to study the use of these systems to support dance creation and improvisation. Previous HCI research has shown the value of deploying and assessing technology for dance and performance *in-the-world*¹ [5, 10, 15].

A few recent works aimed to bridge this gap and assess professional dancers' use of generative models. Wallace et al. [51]

¹We use the term *in the world* instead of *in the wild* as argued by Ssozi-Mugarura et al. [45].

studied how dancers improvise with an abstract dance movement generation system. Their study highlighted the difficulties of using generative systems within established dance practice. In particular, they showed the challenges of providing dancers and choreographers with appropriate interfaces to interact with the generative model without the help of machine learning experts.

More recently, Liu and Sra introduced DanceGen, an interface that allows practitioners to directly interact with a generative model [32]. The model is based on the Motion Diffusion Model [47], fine-tuned on the AIST++ dataset. To evaluate the interface, the authors conducted a study in which dance practitioners used the system to create sequences, followed by semi-structured interviews. The study found that while the system sparked creativity, dancers emphasized the importance of integrating such tools into studio settings.

While these works present valuable insights about how dancers interact with the generative models, their agency often remains limited. The models are not designed or trained with the practitioners, whose input is only gathered through short improvisation sessions once the system is developed. As a result, there might be a mismatch between the dancers' repertoire and the one known to the model, limiting its adoption by dancers in choreographic and improvisation settings. For example, models trained on AIST++ include a mix of styles and genres, while the studies tend to be conducted with dancers practicing contemporary dance styles. This reflects a trend noted by Fdili Alaoui [15], whereby researchers frequently evaluate their systems in improvisational contexts. In such settings, dancers approach the systems without predefined expectations, unlike situations where the systems would directly encode choreographic material or stylistic conventions specific to their practice. Consequently, dancers often engage with and show interest in these systems—not necessarily because of the systems' intrinsic value, but rather due to the dancers' own creative capacities.

With a focus on Voguing and Dancehall, our goal is to understand whether a generative model can capture the essence of these dance styles and support the practitioners' artistic practices and expectations. To close the loop between AI development and dance practice, we collaborated with a dance collective from the early stage of the research project, following a participatory and collaborative approach to design a system that supports their dance practice and improvisation. Our approach is iterative, leading to generative models that are grounded in and assessed in the real-world practice of the dance collective.

3 Overview of the Method

This paper presents a close collaboration with two dancers from a Voguing and Dancehall dance collective based in Paris, which took place over three years.

3.1 Context of the Project

This research emerged from a larger project called Living Archive initiated in our group that aimed at designing technologies to support the documentation and archiving of non-institutionalized

dances² that have not been the subject of academic research, technological design or institutional archival efforts. In particular, we were interested in the use of generative machine learning as a means to explore, interpret and transmit knowledge in dance. In this context, such models create opportunities to provide dancers with real-time feedback on their practice or serve as improvisational partners that nurture their embodied experiences of these dances.

Among the case studies that were initiated in the project, we were involved in dances that developed in underground and LGBTQ+ milieus, namely “Ballroom dances”. We initiated a collaboration with Ballroom dance practitioners that practice Voguing and Dancehall. We contacted the collective through social media and personal connections. After a first conversation with the collective founder Dancer-M, we initiated a collaboration with him and another dancer from the collective, Dancer-Y. Our collaboration with the collective was motivated by the dancers' interest in movement generation to support their choreographic practice. Additionally, we were interested in assessing the potential of generative models with dances such as Voguing and Dancehall, because these styles are highly codified and follow precise and structured steps. This contrasts with most prior works, in which generative models are trained and evaluated on less codified movements that are usually improvisational [37].

With Dancer-M and Dancer-Y, we embarked on the design of a generative AI model and a series of studies that aimed at investigating how they experience and interact with various iterations of the model to create movements and improvise within their specific styles and repertoires as further described in the following Section 3.2.

3.2 Iterative Design With Connoisseurs

We followed an iterative design process with connoisseurs [44]. We designed a dance movement generation model that aligns with the Voguing and Dancehall repertoires practiced by the dancers with whom we collaborate. Through their practice, Dancer-M and Dancer-Y are connoisseurs of these styles. By collaborating with them, we placed their dance expertise at the heart of our movement generation model design. Our iterative design unfolded over three years, yielding several generations of models and interfaces, each assessed through empirical studies. The three main stages of our research process are presented in Figure 1.

The first stage of the research, described in Section 4, presents the design of the first iteration of the machine learning-based movement generation model. In this stage, Dancer-M and Dancer-Y participated by curating steps from their repertoires, and we recorded both dancers performing them using motion capture. We trained a first iteration of the model, which we evaluated by analyzing the two dancers' experience of improvising with its generated movements.

The results of the improvisation session showed the model's creative potential and technical limitations and the need for an interface to control the generative process. This laid the ground for the second research stage, described in Section 5. In that section, we describe the design of the *Korai* system that helped the research

²Institutionalized dance practices typically include ballet, modern or contemporary dance.

team train a second iteration of the generative model. While the dancers contributed to the design of the generative models, Korai was created primarily to support the researchers' workflow in training the models prior to sharing them with the dancers. After training a new model, we ran a structured walkthrough to gather feedback from the dancers about how they perceive the new model that provides more fluid and physically realistic movement sequences.

The structured walkthrough showed a mixed response to the refined movement generation model, despite its physical realism. This was due to the model's lack of faithfulness to the dance style. This inspired us to assess the dancers' experience according to the generation's level of faithfulness to the dance style. In the third stage of our research (Section 6), we conducted a comparative structured observation study, evaluating the dancer's experience across three variants of the generative model that varied from high, medium, to low fidelity to the original dataset.

4 Stage 1: First Iteration of an AI-based Movement Generation Model

This section describes the development of a first iteration of a generative model tailored to the specific repertoire of the collective by training it exclusively on their data.

4.1 Model Design and Training

We designed an initial generative model that produces sequences of movement based on a training dataset composed exclusively of the collective's motion data. Our generative model uses the *Chor-rnn* architecture [11], considered state-of-the-art at the time. We chose *Chor-rnn* because of the compact size of its architecture, enabling us to train model instances from scratch on limited data without relying on large and generic datasets (as opposed to transformer-based architectures). Besides, *Chor-rnn* also enables performing movement generation without relying on other modalities such as sound or text. By prompting the generation using a specific movement, the model generates movements that are close to the source and that act as their continuation. Since the collective's repertoire consists of two distinct dance styles, we hoped that prompting the model with a movement in a given style would steer it toward that style.

The first step of designing the model involved collecting motion capture of the collective's repertoire. Next, we trained the model and assessed its generated movements qualitatively to ensure they produced inspiring movements.

4.1.1 Recording the Motion Capture Dataset. To create the training dataset, we recorded a set of sequences of Voguing and Dancehall of the collective using motion capture. The sequences of Voguing and Dancehall were performed by both dancers from the collective, Dancer-Y and Dancer-M. They were based on a repertoire curated by the two dancers, composed of foundational steps from the Voguing and Dancehall styles. The complete list of the recorded dance moves is listed in Appendix A. Over two days, we recorded each dancer performing each sequence individually. We used an Optitrack marker-based optical motion capture system composed of 18 Flex 3 cameras connected to the Motive software (version 2.3). The dancers wore full Optitrack motion capture suits with 41

passive retro-reflective markers, as shown in Figure 2. Using these markers, the software reconstructs the dancers as 3D skeletons of 21 joints. The complete motion capture recordings amounted to approximately 35 minutes of data with a frame rate of 100 frames per second, which we exported as BioVision Hierarchical (BVH) files. The BVH format is a data file format that stores and structures 3D movement data in a hierarchical way [35].

4.1.2 Training and Generating Movements With the Model. We reimplemented *Chor-rnn* and trained a version of this model solely on the collective's motion capture data. Technical details of the model's implementation and training are given in Appendix B.2. Models based on this architecture perform *motion continuation*: they are prompted on an existing movement of arbitrary length, which acts as the *source movement*, also called *seed*. The model then generates the continuation of the seed as a time series of motion data in a frame-by-frame fashion. To generate sequences, the model's output is used as its own input for the next frame. Through its architecture, the model retains information from past sequences. The generation process is illustrated in Figure 3.

After training the model, we could generate movements using a command-line interface. The parameters available to control the generation are the extract chosen as the seed (its nature and length) and the desired generation length. Our pipeline then outputs the generated movements as BVH files. We used the software Blender³ to visualize the results. An example of a sequence produced by the model and visualized in Blender is shown in Figure 4.

4.2 Experimenting With the Voguing and Dancehall Generative Model

We ran a movement session with Dancer-Y and Dancer-M, where they were invited to improvise with the generated movements resulting from the model trained on their data. The goal was to probe how they perceive and interact with the generated sequences and to understand the potential and limitations of this iteration of the model. At that point, the model's outputs were glitchy and discontinuous, going from one pose to another every second. By *glitches* in the generation, we refer to physically unrealistic motion artifacts where for example body parts move at unnatural speed, adopting anatomically impossible positions, or intersecting with one another. Previous works in the literature argued that glitches in digital systems hold creative and generative potential for dance [26, 52]. Therefore, despite their glitchy and discontinued nature, we were interested in probing the dancers' response to the generated movements as sources for dance improvisation.

4.2.1 Dance Workshop. The workshop was conducted in a professional dance studio. After warming up, the dancers performed four improvisational duets, each lasting approximately 10 minutes. Generated movement sequences were projected on one of the studio's walls during each improvisation. The dancers were free to look at them, follow them, or ignore them altogether. They were also free to set up the space for themselves and choose any music they preferred to improvise with, in addition to the visuals.

³<https://www.blender.org/>

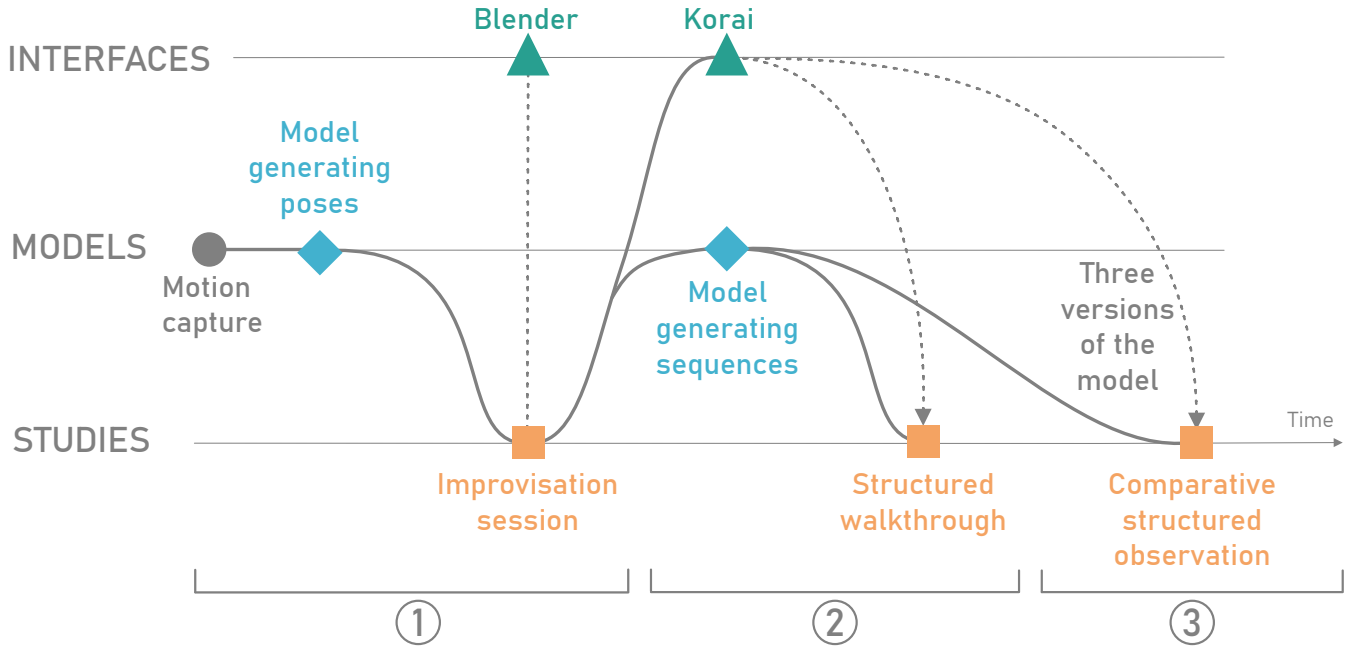


Figure 1: We present our research process in three main stages, detailing the interfaces used or designed, the generative dance models iterations and the studies conducted.

4.2.2 Data Collection and Analysis. The improvisation sessions were video-recorded. After each improvisation, we engaged in group discussions with the dancers, which were also recorded. The discussions were transcribed by the first author. Additionally, the first and last authors watched the videos of the improvisations and annotated them according to the most salient movement features performed. The last author is a certified Laban Movement Analyst. She used her movement observation skills to produce the movement annotations from the videos [anonymous reference]. Following a thematic analysis approach, the first, second, and last authors worked independently to familiarize themselves with and generate codes from both the annotation and the transcript. Through collaborative discussion, these codes were then merged, compared, refined and grouped, resulting in the three themes described in the next findings section.

4.2.3 Findings. We describe the three themes that reveal the potential of this first iteration of the AI-based movement generation model for dance improvisation in the context of Voguing and Dancehall, as well as its limitations.

Incorporating Postures in Fluid Movement Improvisations. Throughout the improvisations, the dancers interacted with the model by embodying the postures that were in front of them, either partially or entirely. Since the model's outputs were displayed in a pose-by-pose fashion, it allowed the dancers to analyze and incorporate the poses into their improvisations. The slow pace of the display made the generated poses more understandable and easier to perceive which enabled the dancers to integrate them on the fly into their improvisation.

We observed that when a specific gesture in a posture was clear to the dancers, they imitated the part of the body that performed the gesture and improvised with the rest of their bodies. For example, they imitated the specific gesture of an arm or a leg while improvising with the rest of their bodies. They performed this type of imitation and integrated gestures into their improvisations fluidly and seamlessly. Occasionally, when the poses were stylistically evocative of Voguing and Dancehall, they mirrored the model's entire poses, blending them into their ongoing improvisation. Note that because of the glitches present, most of the time the poses were only partially stylistically legible as the dancers recognized one hand's gesture or one leg's angles.

Borrowing from the Poses to Start, Interpolate or Transition Between Sequences. Even if the poses were not physically realistic or glitchy, the dancers drew inspiration from them and incorporated them in their improvisations. The model acted as a source of inspiration, providing what we call *posture cards*. These posture cards were used by the dancers to draw ideas from to start new sequences of movements. They also used them to integrate new sequences in their improvisations. While improvising, they interpolated between the sequences that they were performing and the ones they borrowed from the model. This interpolation was performed while maintaining fluidity and continuity in their improvisation. They also used the postures to transition between the two dance styles of Voguing and Dancehall, especially when these postures were ambiguous, borrowing from both or merging Voguing and Dancehall styles.



Figure 2: One of the dancers recording a Voguing sequence in the motion capture studio.

Limited Creative Possibilities over time. We observed that over time, the dancers paid less attention to the projection of the generated movement and favored improvising with each other and with the music. Indeed, once they had interacted with the model over a few sessions, they stopped drawing inspiration from it. We believe that this is due to the model’s technical limitations, which did not allow for the generation of enough dance material to sustain a long interest in the dancers. Once they exhausted the creative possibilities of the model, we observed how they shifted the source from which they drew their motivation for improvising and focused more on interacting with one another and on responding to music.

4.3 Lessons Learned

We articulate several lessons learned from our experiments. First, regarding the model’s quality: although the generated poses were stylistically evocative of postures and gestures present in the collective’s repertoire, the experience seemed to reveal that higher-quality generation with physically realistic sequences would be needed to enable dancers to interact with complete dance sequences, moving beyond isolated and glitchy poses.

Second, the generation workflow lacked interactivity: Working with the terminal with offline rendering and importing each file into Blender was tedious when repeated several times. Moreover,

generating sequences offline made it difficult to assess movement generation quality. Conventional metrics, such as validation losses, are insufficient for evaluating the creative potential of generated poses. This highlights the need for real-time rendering tools that can allow visual assessment of the generation.

Lastly, several challenges arose when attempting to visualize the files effectively, as BVH skeleton data is not designed to be viewed directly in 3D animation software. Instead, it should be paired with an avatar, which requires expert knowledge in “attaching” (a process called rigging) the animation data to a 3D avatar. This process is both complicated and time-consuming, which reveals that future workflows should prioritize accessible tools and streamlined processes to enable customizable renderings.

5 Stage 2: Designing an Interactive Workflow to Improve the Generative Model

This section describes the second iteration of the design of the generative model for Dancehall and Voguing movements based on the lessons learned in the previous iteration. We motivate and describe the design of *Korai*, an interface that allows the research team to directly interact with the generation processes to track improvements over the model easily. We discuss preprocessing and hyperparameter choices for the model, which constitutes a core part of the model’s design. Finally, we present the results of a structured walkthrough with Dancer-Y.

5.1 Designing *Korai*, an Interface for Interactive Movement Generation and Visualization

Our initial experiments with Voguing and Dancehall movement generation underlined several challenges related to the workflow for training, generating, visualizing and assessing movement sequences. In particular, we lacked an appropriate tool for visualizing the results with the dancers and, more critically, for the research team to monitor and evaluate the quality of the generated dance movements throughout the model development cycle. To address these challenges, we designed *Korai*, a web application to visualize 3D movements and interact with movement generation models.

5.1.1 Design Goals. To support the researchers’ workflow, we derived a set of design goals for *Korai*:

- (1) **Facilitate Data Exploration and Interactive Seeding.** The generative model performs movement continuation, requiring a *seed* – a segment of source movement – to generate a new sequence. The system should enable interactive selection of generation seeds by navigating the motion capture recordings. Movement generation should be performed with low latency so that the exploration process remains interactive, fostering navigation between recorded and generated content.
- (2) **Enable Interactive Visualization of Source and Generated Movements.** Instead of a tedious process requiring the manual rendering of generated motion capture sequences with Blender, *Korai* should implement real-time 3D rendering of moving bodies. The rendering should be interactive, so that users can manipulate the 3D scene, and preferably web-based so that a variety of users can access the movement

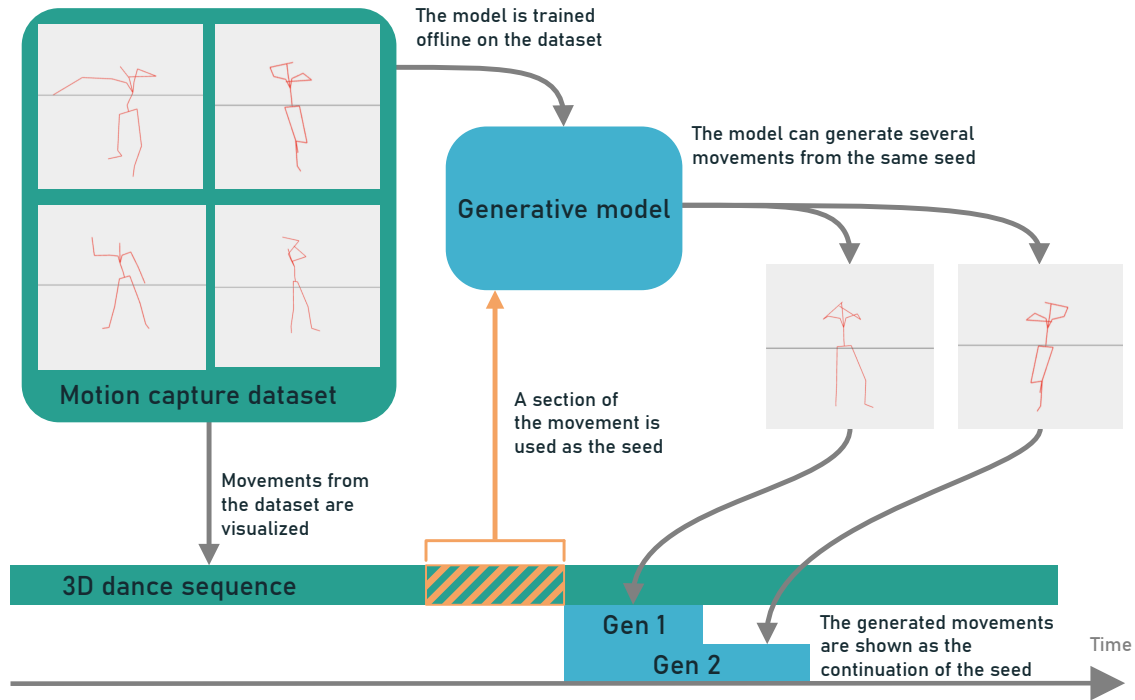


Figure 3: Diagram illustrating the workflow of the generative model. The motion capture dataset is used to train the model, which can then generate new movements. Both the original dataset and the model's outputs can be visualized for comparison.

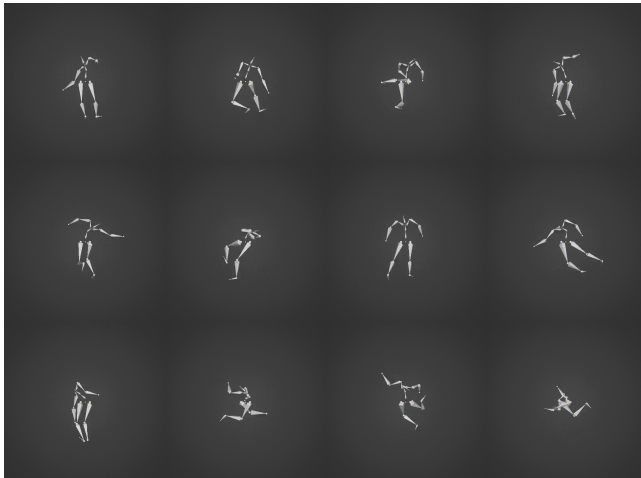


Figure 4: Examples of poses generated by the model and visualized using Blender.



Figure 5: Dancers improvise with the generated poses projected onto the studio's wall. Starting with the second session, we used two camera perspectives within Blender at the dancers' request, allowing them to draw from different parts of the avatars.

data without platform or installation constraints. To facilitate quality assessment, the seed and its generated continuation should be visualized together consistently, enabling users to compare the input and output visually.

- (3) **Present Comparisons of Movement Alternatives.** As the model is stochastic, it produces different results every time,

even when prompted with the same seed. Ideally, the generation model should capture variability within the training data, generating new but realistic variations when repeatedly prompted. The system should provide a way to generate and visualize several sequences using the same seed to compare them. Korai should also allow users to switch between different versions of the model to assess the differences in generated sequences in similar contexts.

5.1.2 System Description. Korai is a web application, available as open-source software⁴. A screenshot of Korai's interface is presented in Figure 6. Korai is configured with a dataset of motion capture recordings. Users can navigate this corpus using a drop-down menu to select animation files. When choosing a file, the system loads the BVH data in a motion capture player for real-time interactive visualization. Movements are rendered in a 3D viewer, allowing users to move, rotate or zoom into the scene. Users can play and pause the animations using buttons and they can navigate in the animation using a progress bar, and adjust the playback rate – that can be negative, in which case the animations play in reverse. Regarding the visual aspect of the avatars performing the dance movements, we kept a skeleton-like visual with thick bone segments. This choice was a compromise between simple stick figures composed of points and lines without thickness, which can be too rigid for dancers, and more elaborate avatars that impose specific body representations. Our skeleton-like visualization results in a representation close to a human body but abstract enough to help dancers focus on the movements.

The generative movement model is directly integrated into Korai, enabling users to request new generated sequences interactively. Users can directly select the *seed* used for generation using the progress bar of the current *source animation*. They can choose the seed duration (how many seconds of the movement are used by the model) and the duration of the generated sequence, which is a continuation of the seed in the dance style on which the model was trained. To give the model enough context, we constrained the source movement to be at least 1 second long, but the users can generate sequences as long as they wish.

Thanks to the lightweight nature of the model, inference can be achieved in quasi-real-time with any modern computer, even without a dedicated GPU.⁵ A settings drawer gives access to a few visualization parameters, such as the version of the generation model used, enabling them to choose different trained models with varying parameters for easy comparison.

When a movement is generated, it is displayed in the same 3D viewer as the source animation. Generated movements play synchronously with the source movement from which they were generated. Each time a sequence is generated, a new timeline is added at the bottom of the player, where the generated sequence is highlighted as a colored segment. Users can toggle the visibility of the various sequences with checkboxes, making it possible to visualize several generated sequences together, with or without the source animation. For each generated output, a collapsible panel is created to display information about the model used for generation and the training parameters, allowing users to delete it.

Rather than enabling users to compose dance sequences by concatenating movements generated from various source animations, we chose to attach the model output to their source animation. The application's state is automatically stored: generated movements are removed from the view when the user changes the source animation, and are restored if the user reloads the same source animation. Users can delete the generated movements and recreate

as many as they wish. Technical implementation of Korai is detailed in Appendix C.

Following the study in Stage 1, our initial goal for the generative model was to produce complete, fluid dance sequences. Over time, we refined a workflow facilitating the model development work, which required many experiments on data preprocessing, movement representation, data augmentation, and hyperparameter tuning. Training of the models was done on a remote cluster, and we used Tensorboard⁶ for monitoring the training process, using training and validation losses to ensure proper training of the model. Once a new model was trained, it was evaluated through direct interaction in Korai and visual assessment. For such motion generation, loss metrics are often not informative enough and do not capture aspects related to anatomical realism or the creative potential of movement. Korai was therefore instrumental in visualizing and evaluating how changes in data representation or hyperparameters affected the model's capacity to produce realistic output and to generalize with new sequences.

5.2 Structured Walkthrough

To review and discuss the improvements made to the model, we ran a structured walkthrough with Dancer-Y, which is an effective method to obtain feedback about a system [55]. The model used in the walkthrough was trained with the goal of generating movement sequences that were fluid and physically realistic. They also needed to be creatively inspiring to the dancers. To do so, we trained a model that produces movements that are neither too close nor too far from the original dataset. These movements are both recognizable but also provide novelty for the dancers. This choice was inspired by the guidelines from Wallace et al. that argue for the creative value of AI models that respect a good balance between novelty and realism [53].

5.2.1 Method. During the structured walkthrough, the first author served as the presenter. They prepared a variety of movement outcomes in advance generating them using seeds from both Dancehall and Voguing movements. The purpose of the walkthrough was to walk Dancer-Y through the generation of these movement outcomes visualized directly in the Korai system and identify with them as many problems and opportunities as possible. During the walkthrough, we made space for critiques, questions and open discussions to emerge. The walkthrough lasted approximately 40 minutes. The first author captured the walkthrough via audio recording then transcribed the data. The research team separately familiarized themselves with this transcript and, through collective discussion, identified the two main limitations highlighted by Dancer-Y.

5.2.2 Findings.

Difficulty Recognizing Dance Styles. The generative model works by creating a possible continuation of an existing movement. The second iteration of the model that we used in the walkthrough produced more continuous and fluid sequences of movements. While these movement sequences were deemed physically realistic, they left Dancer-Y surprised and unconvinced by the results of the generation. He found that the movements generated were not legible

⁴<https://gitlab.lisn.upsaclay.fr/lched/korai>

⁵As a reference, generating 30 seconds of movements takes approximately 10 seconds on the laptop we used, which had an 11th-generation i5 processor and no GPU, including exporting the files.

⁶<https://www.tensorflow.org/tensorboard>

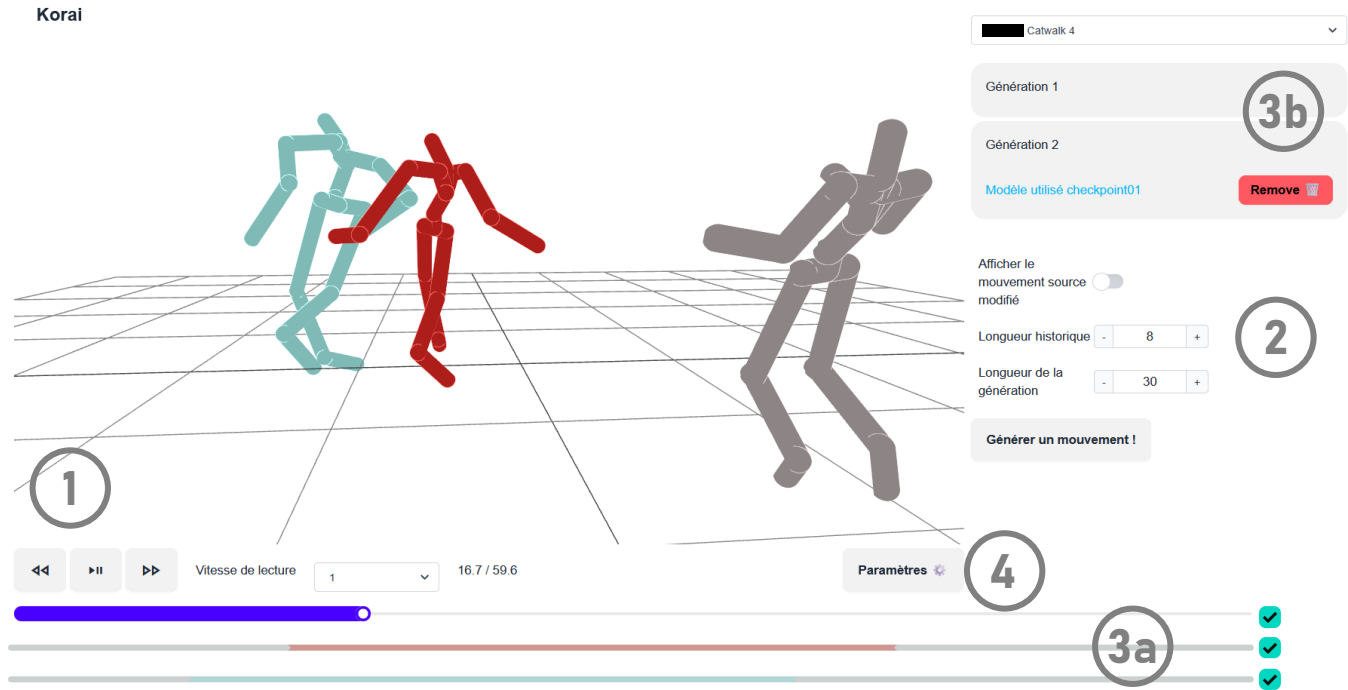


Figure 6: The *Korai* application. The main parts of the application are (1) the 3D viewers to visualize movements, along with navigation controls, and (2) the panel on the left that allows users to generate new movements from the currently displayed source movement. Users can choose how long the seed should be, how long the generated movement should be and if they wish to keep the seed in the resulting animation. (3a) The progress bars display the position and duration of each generated movement played alongside its seed. (3b) Panel displays the models used to generate the movements. They allow users to delete generated movements individually. (4) Finally, a settings drawer lets users change the generative model used.

and recognizable, knowing the dance steps in Voguing and Dancehall. He recognized the source movements that were movement sequences from their repertoire captured using motion capture but did not recognize the generated outputs following them. An explanation for this is that the generative model is not explicitly conditioned to a dance style during training or inference. Because our training repertoire consisted of both Voguing and Dancehall, the resulting generated movements appeared to be in between these two dance styles and did not clearly belong to one or the other. This created difficulty associating the generated dance sequences with Voguing and Dancehall styles.

Difficulty Recognizing Dance Structures. According to Dancer-Y, despite the improved quality of the generation, the sequences did not have consistency in rhythm. He also argued that the movement sequences did not follow the structures of the dances. He did not understand the logic of the sequence and why one movement was followed by another: “I don’t see how it [the model] can take what I was doing below [the training movement] and then do that [the generated movement].” An explanation for this is that the generative model was seen by the dancer as generating movement that partially resembled Voguing and Dancehall but they didn’t perceive the overall sequences as corresponding to Voguing or Dancehall steps. This led him to form higher expectations about the structural and

stylistic legibility of the sequences. Indeed, the dancers tended to compare the generated sequences to the sequences from the training dataset, arguing that the generation was less faithful to the styles and structures of Voguing and Dancehall than the recorded motion capture data and is thus less usable for dances that are highly codified, such as Voguing and Dancehall.

5.3 Reflections

In the experimentation with the first iteration of the model, despite the glitchy, physically unrealistic, and discontinuous nature of the generated poses, the dancers were able to use them creatively in their improvisation. Surprisingly, their responses to our structured walkthrough indicated that they were critical of the second iteration of the model that provided continuous, fluid and more physically realistic movement sequences. The second model did not seem novel enough to spark their creative interest, nor stylistically legible enough to be useful in their practice. These results led us to explore whether their perception and experience of the model depend on the stylistic legibility of the movements generated by the model.

6 Stage 3: Studying Dancers' Experience of Movement Generation According to Levels of Stylistic Legibility

To investigate the effect of stylistic legibility on the dancer's experience, we designed a comparative structured observation study [34] that compares the dancers' perception, experience and response to three different movement generation models used as variants. The three models varied from high, medium, or low fidelity movements generated depending on the degree to which the model fits the training data, which is obtained by varying the number of training steps of the model.

6.1 Experimental Setup

The structured observation was conducted with Dancer-Y in a professional dance studio that was one of his usual rehearsal studios. The first author brought their laptop to the studio and ran Korai to visualize the movements generated from the three generative models using the same setup. They then projected their screen onto a wall to allow Dancer-Y to view the movements while dancing freely in the studio space.

6.2 Procedure

During the study, Dancer-Y engaged in three improvisations with movements generated by three different models visible in Figure 7, in the following order:

- (1) The first model was called "medium-fidelity" and was the same model used to showcase the movements during the structured walkthrough. It provided movements that were not too close nor too far from the original dataset. This model was obtained with the training procedure detailed in appendix B.2.
- (2) The second model was called "high-fidelity". It provided movements from a highly overfitted version of the training data. These movements were close to the original dataset both in terms of physical realism and stylistic faithfulness. This model was obtained by overfitting on the collective's training data. Compared to the previous training procedure, we did not use a validation set but the whole dataset to train the model. Therefore, we did not use early stopping and instead trained the model for a fixed 230 epochs.
- (3) The last model was called "low-fidelity". It produced movements that were far from the original dataset. These movements did not resemble the dance styles and occasionally produced glitches and physically unrealistic movements. This model was obtained using the checkpoint saved after the 10th epoch of this same training dataset.

Each improvisation with a given model lasted approximately 20 minutes. The dancer could move freely and choose the music he preferred. After each improvisation, an interview was conducted with the dancer.

6.3 Data Collection and Analysis

The improvisations were video-recorded using a fixed camera in the studio that recorded a wide-angle view, including most of the rehearsal space and the wall with the projected Korai interface.



(a) Movement sequence generated with the medium-fidelity model. The movements were physically realistic and presented partial stylistic faithfulness.



(b) Movement sequence generated with the high-fidelity model. The movements generated were close to the original dataset and were both physically realistic and stylistically faithful.



(c) Movement sequence generated with the low-fidelity model. The movements generated were far from those in the original dataset, occasionally producing glitches and physically unrealistic poses.

Figure 7: Examples of generated movements for the three sessions. The sequences all last 10 seconds, and the postures shown here are extracted every 1 second.

A research team member also took close-up videos during the improvisations. The interviews after each improvisation session were audio-recorded. Logs of the generated movements with their metadata were also collected during the improvisations.

We conducted a reflexive thematic analysis [8] on all the data collected. We used a mix of inductive and deductive thematic analysis. Guided by the structure of our comparative structured observation, we followed a deductive approach that focused on comparing the responses of the dancer to the three models' generated movements. We also followed an inductive approach where we actively defined and named themes from the codes generated to describe the data.

Precisely, the first author transcribed the interview data. Then, along with the last author, who is a Laban Movement Analyst, they watched the videos of the improvisations and annotated them according to the most salient movement features performed. Following that, the first, second, and last authors read both the transcripts of the interviews and the video annotations. Then, they highlighted quotes from the interviews that they identified as relevant to understand the responses of the dancer to the different levels of faithfulness of the models. They generated codes independently from both the annotations and the highlighted quotes in the transcript. In order to reach inter-coder reliability, differences among researchers' codes were discussed and resolved through consensus. Through collaborative discussion, these codes were then merged, compared, refined, and grouped, resulting in the three themes described in the next findings section.

6.4 Findings

We extracted three themes related to how the dancer perceived the models, how he interacted with the system, and the improvisation strategies he deployed in response to the three models.



Figure 8: The dancer improvising with generated movements visualized in Korai.

6.4.1 Stylistic Legibility Affects the Dancer’s Perception and Appreciation of the Generated Movements. The dancer’s perception of the generated movements varied significantly from one model to another based on both criteria of stylistic legibility, and physical realism of the generated movements.

Movements generated with the medium-fidelity model were perceived as difficult to interpret. The dancer did not consider them as a source of inspiration and they were seen as hindering the improvisation. The dancer described the outputs of the model as “not legible”, adding, “It’s hard to compose with something that’s very abstract.” The difficulty in interpreting the medium-fidelity model’s outputs was due to the fact that they were seen as lacking both spatial and temporal structure present in Voguing and Dancehall.

In contrast, after the improvisation with the overfitted model providing high-fidelity movements, the dancer had many positive comments on the generated movements. He described them as “legible”, “clear” and providing “movement, structure”: “I found the [generated movements] much more legible, much easier to bounce off, much more pleasant.” The “simplicity” and fluidity of the movements generated with the high-fidelity model made the dancer perceive them as closer to what he described as *a danced movement*, stating “I find [the high-fidelity model generation] much more lively”. Dancer-Y reported that he was more inspired by these movements.

Dancer-Y found the movements generated with the low-fidelity model difficult to interpret. He reported that glitches in the generated movements were a disadvantage and led to physically unrealistic movements that lacked fluidity. Despite the fact that the low-fidelity model generated the most glitched movements, the dancer was able to recognize elements of their Voguing repertoire, more so in this model than in the medium-fidelity model. He reported finding the low-fidelity movements more inspiring than the ones generated by the medium-fidelity model. He described the movements as “strange, fast” yet related them more to the original vocabulary and found inspiration in them. This indicates that physical realism alone does not fully explain the appeal of the model.

In summary, he saw the movements provided by the model that was closest to the original dataset, the high-fidelity model, as fluid, clear and legible. He considered the movements provided by the model that was furthest to the original dataset, the low-fidelity model, as physically unrealistic but novel and inspiring. He judged the movements from the medium-fidelity model as stylistically illegible and creatively uninspiring and described them as “abstract,” and hard to interpret.

6.4.2 High-fidelity Movement Generation Facilitates Fluid Improvisation. Overall, the dancer preferred the high-fidelity model, which allowed for a more natural improvisation. He said that “it allowed me to be much more lively and to make more creative choices compared to the other time [referring to the medium-fidelity model].” Therefore, fidelity to the training data, stylistic legibility and interpretability of the movements encouraged improvisation: “Seeing [the avatar] move in a much more natural and much cleaner way, it made me want to move more.” The recognizable movements of the high-fidelity model were easy for the dancer to pick up and to interact with. They encouraged him to borrow from the movements displayed on the screen and integrate them continuously into their improvisation.

With the low-fidelity model, the dancer stated that he struggled to “understand” the glitchy movements and “pick up something interesting” to use in his improvisation. This forced him to stop the improvisation and think more consciously about how to incorporate elements of the generation into the dance. He indicated that “too many things are happening” in the generated movements, often forcing him to pause: “I’d pause, try to understand what was happening on the screen, and see how I could react to that.” In this scenario, the dancer described his interaction with the generated movement as a “confrontation”: “There were moments when I imagined that it was really a confrontation, so for example when the AI was moving towards me, well when the model was moving towards me, I would move backward.” Since the model was not providing movement that the dancer recognized clearly as Voguing or Dancehall, he wondered how the unfamiliar movements could contribute to his improvisation beyond mere imitation, asking “what can I do with it”.

While the dancer reported verbally that the medium-fidelity model didn’t inspire him, we observed that it changed his way of dancing. When we asked him about it during the interview following the improvisation with the medium-fidelity model, he explained that while the interaction was not easy, he still drew elements from the projected movements, such as introducing ruptures in his improvisation: “Even if it’s very abstract, seeing it raises questions, intrigues me, it makes me want to go and look, makes me want to stop, freeze the moment. But in any case, even if it’s very abstract, I respond to it. And that also helps to break down the structure I am creating, so that’s what I find interesting.” Thus, the response of the dancer to the medium-fidelity model providing movement with no particular novelty appeal was to reflect on he structure of his improvisation. In his own words: “it allowed me to structure [my improvisation], but by doing something unstructured.”

6.4.3 The Level of Stylistic Legibility Implies Diverse Improvisational Strategies with the Generated Movements. At the beginning of the study, the dancer mentioned that his dance style had evolved since

recording the motion capture, partly because he had begun learning and practicing contemporary dance. This had an effect on how he utilized different dance styles and improvisation strategies in the different improvisations with the three models, something we observed from the video recordings and that he acknowledged during the interviews. These styles and improvisation strategies varied according to the different models' levels of stylistic legibility.

When improvising with the medium-fidelity model, because of the difficulty of interpreting the generated movements, the dancer preferred to distance himself from the generated movements: "it could have been something like "I see Dancehall on the screen, so I respond by doing Dancehall" or "I see Voguing on the screen and I respond by doing Voguing" and so on, but it's not necessarily the case." His improvisation strategy did not consist of imitating the generated movements or reproducing known steps of Voguing or Dancehall. Instead, we observed that he alternated between using the rhythm and using poses of the generated movements, but he rarely incorporated entire movements. In this configuration, his improvisation style was closer to a contemporary dance style or free improvisation.

Similarly, with the low-fidelity model, the generated movements' lack of physical realism led to reduced legibility for the improviser. We observed that the dancer faced various challenges in keeping up with the generated movements. He often took pauses to understand and respond to the less realistic outputs. This had the effect of pushing him away from the original repertoire's style (Voguing or Dancehall) and towards more abstract movements that can be identified as contemporary dance or free improvisation. When the dancer picked up a movement that he wasn't able to recognize or to integrate in his ongoing style, we saw him shift his improvisation style to free improvisation. This was confirmed by the dancer: "what I did during this improvisation session... it was just an exploration of movement qualities and techniques that are completely... In fact, I'd say there's not even any technique, because it's just feeling, just lines, whereas Dancehall has steps and Voguing has steps."

In contrast, with the high-fidelity model that produced more physically realistic and stylistically legible movements, the dancer was imitating the model more closely and borrowing more from what was generated, which he interpreted as Voguing or Dancehall. The improvisation, therefore, consisted of creating variations of the generated movements. The more realistic and legible the movements, the more fluid the improvisation with movements closely aligned with the original dance styles of Dancehall and Voguing.

The evolution of the dancer's practice had opened up new improvisational possibilities where he could go beyond imitation of Voguing and Dancehall, utilizing free-form improvisation when the generated movements projected were unfamiliar to him. This revealed a contrasting strategy to respond to the legibility and novelty of the projected movements, either by directly imitating them when they are the clearest or by transforming them and using them as prompts for free improvisation when they are the least familiar

7 Discussion

This paper presents a long-term collaboration with a Voguing and Dancehall collective, with whom we designed and assessed movement generation models leveraging their knowledge of the repertoire.

7.1 Navigating Stylistic Legibility, Realism and Novelty: An Uncanny Valley of Movement Generation?

Compared with earlier studies that focus on balancing novelty and realism [53], we found that physical realism by itself fails to reflect how dancers react to generated movements, particularly within highly structured styles such as Voguing and Dancehall. Our studies elicited responses regarding both physical realism and stylistic legibility of the generated movements. Our findings for the various models used in the studies are summarized in Table 1. Despite a technically limited model with glitchy poses, our first movement session revealed that the dancers found the generated poses stimulating and generative of movement ideas. These movements were not physically realistic and presented limited stylistic legibility, but they were novel and unfamiliar. In the structured observation study, the low-fidelity model, which provided movements that were unfamiliar, sparked dancers' creativity and inspired new ways of moving and improvising.

These findings align with prior CST research showing that glitches and generation of unfamiliar movements can act as catalyst for choreography Wallace et al. [54] by disrupting dancers' habitual movement patterns [33] and allowing them to make creative choices during improvisation [20, 39]

On the other hand, the dancer responded critically to the model that provided more physically realistic and fluid movement sequences presented in the structured walkthrough and as the medium-fidelity model in the structured observation, even though the quality of the sequences had dramatically improved. The model balanced physical realism with novelty, providing movements that did not overfit the original dataset and only partially incorporated features of Voguing and Dancehall. Yet, the dancer perceived it as uninspiring and lacking stylistic legibility.

In our comparative structured observation, we introduced an even more physically realistic and stylistically legible model, resulting from clearly overfitting the original dataset. To our surprise, the dancer expressed a clear preference for this high-fidelity model over the medium-fidelity one, for the clarity and legibility of the dance structures and steps displayed. Because of the highly codified nature of Voguing and Dancehall, dancers expected the movement generation model to produce sequences of movements that reproduced the steps and represented the structures inherent to how the styles are defined.

Our results showed that when both physical realism and stylistic legibility were present, the model inspired the dancers to imitate the movements and incorporate them into his Voguing and Dancehall improvisations. When these qualities were absent (in the low fidelity and the glitchy models), the generated movements were seen as unfamiliar and novel which led the dancers to set aside their expectations and use their creativity to draw inspiration from what

Model	Type of motion	Physical realism	Stylistic legibility
Study 1	Poses	Not physically realistic	Some poses were evocative of the style
Study 2*	Sequences	Physically realistic	Difficulty recognizing the steps' style and structure
Study 3 - Low-Fidelity	Sequences	Physically unrealistic, glitches and lack of fluidity	Some poses were evocative of the repertoire, but the movements were mostly novel and unfamiliar
Study 3 - Medium-Fidelity*	Sequences	Physically realistic	Difficulty recognizing the steps' style and structure
Study 3 - High Fidelity	Sequences	Good physical realism	Stylistically legible

Table 1: Overview of the different models used in this study.**These two models are the same.*

was in front of them. However, with the tradeoff of the medium-fidelity model, the generated movements were physically realistic, but that was not enough. Because they lacked stylistic legibility, the dancer perceived them as neither faithful enough to feed their Voguing or Dancehall improvisation nor novel and unfamiliar enough to spark his creativity. This led him to express harder critiques towards such a model, considering it both "unclear" and uninspiring.

A possible interpretation of these results could be made using the concept of an *Uncanny Valley* of AI-generated movement for dance, by analogy to the term coined by Mori [36] for animation and robotics. Similarly to Mori's theory, the results from our structured observations seem to suggest an uncanny valley where the response to the model becomes negative as the model becomes more—but not exactly—like the original dance styles. As the model gets closer to the original dataset, the response becomes positive again. This concept was recently explored for dance movement generation in a context where participants passively observed the generated dance sequences [30]. This "uncanny valley" effect is likely centered around the capacity of the model to encode stylistic legibility rather than physical realism alone. Thus, our findings complement previous work by Wallace et al., where AI models were trained on non-codified movements from improvisation and generated sequences that did not need to fulfill stylistic legibility [53]. Our results therefore open up an interesting avenue of research in the design of generative models for dance practices, suggesting the need to carefully negotiate stylistic legibility – going beyond mere physical realism – and novelty, tailoring this balance to the intended use of the generation and the codified nature of the repertoire at stake.

7.2 Towards More Participation in AI Technology Design for Dance

An ambition of this research project was to develop AI-based systems in close collaboration with dancers, grounded in their practice and repertoire. We aimed to integrate their expert embodied knowledge in the design of the generative model, so that it serves their own creative endeavors.

Our iterative design of the movement generation models illustrates a process including the dancers at multiple stages of the design and assessment of the models. We positioned the dancers as

"subjects of knowledge" from which we sought to understand how they apply dance knowledge in their practice [23]. As "subjects of knowledge", the dancers provided a curated set of movements according to their collective practice of Voguing and Dancehall that we recorded and used to train the generative model. Their involvement in assessing the models through various studies and discussions offered important insights into how they embody and experience the generated movements which motivated, inspired and led the various iterations of the design process. Thus, drawing from Fdili Alaoui et al.'s taxonomy, our study can be seen as adopting a second-person perspective, centering the development of technology around the unique perspectives of the dancers we collaborated with [17]. Yet, the level of involvement of the dancers in our process is partial when compared to participatory approaches in the HCI literature where dancers are involved in defining what interactive technologies do and how they do it [14, 46]. Indeed, in our case, iterations were long, involving researchers to mediate between the dancers and the models, as in a typical approach to machine learning development [4].

Our method thus presents a first attempt to develop AI models centered around close collaborations with dance connoisseurs whose embodied, and subjective dance practice shaped (although partially) the training and development of the movement generation models. This contrasts with the current literature of dance movement generation models that are neither designed nor trained with the practitioners, whose input is only gathered through short improvisation sessions once the system is developed. This also contrasts with the dominant trends in the development of AI models, which are now discarding years of research in HCI that demonstrated the benefit of participation in technological design [43]. While high-level approaches towards "participatory AI" have been discussed in the literature [6, 29], there remains a critical need for collaborative tools that bridge AI development with the embodied knowledge and expertise of dance practitioners. Thus, we argue that researchers in HCI should continue to aim towards more participation and engagement with communities of practice, both when designing, developing and training generative models and when studying them in situ.

7.3 Korai: Attenuating the Hardships of Designing AI Models

The model development process was particularly tedious, as it required a tremendous amount of experiments (e.g., with movement representation, data preprocessing and augmentation, stopping criteria, etc.). In practice, working with machine learning models still requires expert engineering skills, and the rigid workflows of available software tools restrict the opportunities for domain experts to directly act upon the training process. Indeed, machine learning models are not originally developed to accommodate a method centered on users where the technology can be adapted to their needs, knowledge or practice. Another main challenge is that generative models lack an objective measure of the model's quality, which forces developers to rely on subjective assessments to guide the model-building process. Such subjective assessments are hard to define and set up systematically. While approaches such as Interactive Machine Learning (IML) have been proposed to involve users in a tighter interaction loop with the training process [13, 18], architectures such as Chor-rnn require a large amount of data, computing infrastructures and are time-consuming, preventing such interactive workflows.

We designed Korai to attenuate some of these issues. Korai implements an interactive workflow that reduces the time-consuming processing and representation of data. In addition, interaction with model inference enables the model to be explored interactively and fluidly, making subjective evaluation easier. This interaction seems to facilitate an understanding of the behavior of the generative model through experience. Using experience to assess AI output has also been reported as the main way by which visual artists work and assess AI models' creative outputs [9, 25]. However, to the best of our knowledge, this has never been reported in studies designing generative models in embodied artistic practices such as dance. While Korai was primarily designed to be used by the researchers to facilitate the development of the model, we used it in the studio to generate and visualize movements during our structured walkthrough and structured observation studies. This allowed us to provide a tool for both the dancers and the researchers to visualize, reflect on, and experientially assess the generative models' outputs and the training datasets simultaneously.

7.4 Limitations and Future Works

In this paper, our design method with the dancers from the collective allowed for an in-depth collaboration over a significant period of time, enabling us to gain a deep understanding of the potential that movement generation models hold to support their practice depending on the level of realism, stylistic legibility and novelty provided. This opens up interesting avenues for future research on movement generation in dance. A complementary study could involve a larger number of dancers and dance styles to investigate more systematically how physical realism and stylistic legibility affect dancers' experiences across practices and repertoires.

As we highlighted in our discussion, our study presented a partial participation of the dancers. We see potential in extending participation by deploying movement-generation models within choreographic creation. This would allow practitioners to engage more fully in the design of the generative models, for instance, by

training and testing them over long periods of time dedicated to creation, typically during dance residencies. This could reveal how the generated movements engage the dancers creatively over an extended period, a limitation hinted at in our first study 4.2.3.

Our study builds on the Chor-rnn algorithm, which was state-of-the-art when our collaboration started. While our work shows compelling results in the movements generated, the emergence of new architectures in recent years could help improve the model's representation of the dance styles. Yet, developing models that can both maintain physical realism and faithfully capture the stylistic essence of dance from limited data remains a major challenge.

We think that the system that we developed, Korai, has the potential to facilitate these three possible future studies. The tool could be adapted to be used directly by dancers, which would enable scaling up the research's scope to accommodate more dancers and other genres or more longitudinal studies in the context of residencies. Moreover, it could be extended to support other movement generation models and investigate their use in choreographic settings.

8 Conclusion

In this paper, we presented a collaboration with a Vaguing and Dancehall collective to design and assess a movement generation model leveraging their own repertoire. To alleviate the difficulties related to developing and evaluating generative models, we designed Korai, an interactive system that allows us to visualize and monitor the generation process. Throughout our design process, we included the dancers in a series of studies to investigate how they perceived and experienced various iterations of the movement generation model. The results of our studies showed that dancers found inspiration in movements that were either strange, glitchy and novel, or in movements that were familiar and that closely aligned with their repertoire. Generated movements that fell somewhere in between, physically realistic but lacking stylistic legibility, were found neither inspiring nor faithful to the structure of the dance styles. We suggest a parallel with the concept of uncanny valley, applied to dance movement generation. It is through our long-term engagement with the collective that we were able to gain such a deep understanding of the potential that movement generation models hold to support their practice depending on the level of fidelity to the original style and novelty provided. We conclude by arguing for the need to engage with communities of practice, both when creating generative models and when studying them in situ.

Acknowledgments

We thank the collective members for their collaboration on this project.

This work benefited from State aid managed by the Agence nationale de la recherche under the France 2030 Program, under the references ANR-23-PEIC-0001 and ANR-20-CE33-0006- "LivingArchive: Interactive Documentation of Dance".

References

- [1] Omid Alemi. 2021. *Expressive Movement Generation with Machine Learning*. Ph.D. Dissertation. Simon Fraser University.
- [2] Omid Alemi, Jules Françoise, and Philippe Pasquier. 2017. GrooveNet: Real-time Music-Driven Dance Movement Generation Using Artificial Neural Networks.

- networks 8, 17 (2017), 26.
- [3] Simon Alexanderson, Rajmund Nagy, Jonas Beskow, and Gustav Eje Henter. 2023. Listen, Denoise, Action! Audio-Driven Motion Synthesis with Diffusion Models. *ACM Trans. Graph.* 42, 4 (July 2023), 44:1–44:20. doi:10.1145/3592458
 - [4] Saleema Amershi, Maya Cakmak, William Bradley Knox, and Todd Kulesza. 2014. Power to the People: The Role of Humans in Interactive Machine Learning. *AI magazine* 35, 4 (2014), 105–120. doi:10.1609/aimag.v35i4.2513
 - [5] Steve Benford, Chris Greenhalgh, Andy Crabtree, Martin Flintham, Brendan Walker, Joe Marshall, Boriana Koleva, Stefan Rennick Egglestone, Gabriella Gianachi, Matt Adams, Nick Tandavanitj, and Ju Row Farr. 2013. Performance-Led Research in the Wild. *ACM Transactions on Computer-Human Interaction* 20, 3 (July 2013), 1–22. doi:10.1145/2491500.2491502
 - [6] Abeba Birhane, William Isaac, Vinodkumar Prabhakaran, Mark Diaz, Madeleine Clare Elish, Jason Gabriel, and Shakir Mohamed. 2022. Power to the People? Opportunities and Challenges for Participatory AI. In *Proceedings of the 2nd ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization (EAAMO '22)*. Association for Computing Machinery, New York, NY, USA, 1–8. doi:10.1145/3551624.3555290
 - [7] Daniel Bisig. 2022. Generative Dance - a Taxonomy and Survey. In *Proceedings of the 8th International Conference on Movement and Computing*. ACM, Chicago IL USA, 1–10. doi:10.1145/3537972.3537978
 - [8] Virginia Braun and Victoria Clarke. 2006. Using Thematic Analysis in Psychology. *Qualitative Research in Psychology* 3, 2 (Jan. 2006), 77–101. doi:10.1191/1478088706qp0630a
 - [9] Baptiste Caramiaux and Sarah Fdili Alaoui. 2022. "Explorers of Unknown Planets": Practices and Politics of Artificial Intelligence in Visual Arts. *Proceedings of the ACM on Human-Computer Interaction* 6, CSCW2 (Nov. 2022), 1–24. doi:10.1145/3555578
 - [10] Marianela Ciolfi Felice, Sarah Fdili Alaoui, and Wendy E. Mackay. 2021. Studying Choreographic Collaboration in the Wild. In *Proceedings of the 2021 ACM Designing Interactive Systems Conference (DIS '21)*. Association for Computing Machinery, New York, NY, USA, 2039–2051. doi:10.1145/3461778.3462063
 - [11] Luka Crnkovic-Friis and Louise Crnkovic-Friis. 2016. Generative Choreography Using Deep Learning. In *Proceedings of the Seventh International Conference on Computational Creativity*. Sony CSL Paris, France, Paris France. doi:10.48550/arXiv.1605.06921
 - [12] Scott Delahunta and Norah Zuniga Shaw. 2006. Constructing Memories: Creation of the Choreographic Resource. *Performance Research* 11, 4 (Dec. 2006), 53–62. doi:10.1080/13528160701363408
 - [13] John J. Dudley and Per Ola Kristensson. 2018. A Review of User Interface Design for Interactive Machine Learning. *ACM Trans. Interact. Intell. Syst.* 8, 2 (June 2018), 8:1–8:37. doi:10.1145/3185517
 - [14] Sarah Fdili Alaoui. 2019. Making an Interactive Dance Piece: Tensions in Integrating Technology in Art. In *ACM Designing Interactive Systems*. ACM, San Diego, United States, 1195–1208. doi:10.1145/3322276.3322289
 - [15] Sarah Fdili Alaoui. 2023. *Dance-Led Research*. Habilitation à Diriger Des Recherches. Université Paris-Saclay.
 - [16] Sarah Fdili Alaoui, Kristin Carlson, and Thecla Schiphorst. 2014. Choreography as Mediated through Compositional Tools for Movement: Constructing A Historical Perspective. In *Proceedings of the 2014 International Workshop on Movement and Computing*. ACM, Paris France, 1–6. doi:10.1145/2617995.2617996
 - [17] Sarah Fdili Alaoui, Thecla Schiphorst, Shannon Cuykendall, Kristin Carlson, Karen Studd, and Karen Bradley. 2015. Strategies for Embodied Design: The Value and Challenges of Observing Movement. In *Proceedings of the 2015 ACM SIGCHI Conference on Creativity and Cognition (C&C '15)*. Association for Computing Machinery, New York, NY, USA, 121–130. doi:10.1145/2757226.2757238
 - [18] Rebecca Fiebrink, Perry R. Cook, and Dan Trueman. 2011. Human Model Evaluation in Interactive Supervised Learning. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. ACM, Vancouver BC Canada, 147–156. doi:10.1145/1978942.1978965
 - [19] Jules Françoise, Baptiste Caramiaux, and Téo Sanchez. 2021. Marcelle: Composing Interactive Machine Learning Workflows and Interfaces. In *The 34th Annual ACM Symposium on User Interface Software and Technology*. ACM, Virtual Event USA, 39–53. doi:10.1145/3472749.3474734
 - [20] Jules Françoise, Sarah Fdili Alaoui, and Yves Candau. 2022. CO/DA: Live-Coding Movement-Sound Interactions for Dance Improvisation. In *CHI Conference on Human Factors in Computing Systems*. ACM, New Orleans LA USA, 1–13. doi:10.1145/3491102.3501916
 - [21] Satoru Fukayama and Masataka Goto. 2015. Music Content Driven Automated Choreography with Beat-Wise Motion Connectivity Constraints. *Proceedings of SMC* (2015), 177–183.
 - [22] Kehong Gong, Dongze Lian, Heng Chang, Chuan Guo, Zihang Jiang, Xinxin Zuo, Michael Bi Mi, and Xinchao Wang. 2023. TM2D: Bimodality Driven 3D Dance Generation via Music-Text Integration. In *2023 IEEE/CVF International Conference on Computer Vision (ICCV)*. 9908–9918. doi:10.1109/ICCV51070.2023.00912
 - [23] Tove Grimstad Bang, Sarah Fdili Alaoui, and Elisabeth Schwartz. 2023. Designing in Conversation With Dance Practice. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. ACM, Hamburg Germany, 1–16. doi:10.1145/3544548.3581543
 - [24] Chuan Guo, Shihao Zou, Xinxin Zuo, Sen Wang, Wei Ji, Xingyu Li, and Li Cheng. 2022. Generating Diverse and Natural 3D Human Motions from Text. In *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*. 5142–5151. doi:10.1109/CVPR52688.2022.00509
 - [25] Drew Hemment, Dave Murray-Rust, Vaishak Belle, Ruth Aylett, Matjaz Vidmar, and Frank Broz. 2024. Experiential AI: Between Arts and Explainable AI. *Leonardo* 57, 3 (June 2024), 298–306. doi:10.1162/leon_a_02524
 - [26] Beverley Hood. 2012. Choreographing the Glitch. In *4th International Workshop on Physicality*.
 - [27] Biao Jiang, Xin Chen, Wen Liu, Jingyi Yu, Gang Yu, and Tao Chen. 2023. MotionGPT: Human Motion as a Foreign Language. In *Proceedings of the 37th International Conference on Neural Information Processing Systems (NIPS '23)*. Curran Associates Inc., Red Hook, NY, USA, 20067–20079.
 - [28] François-Joseph Lapointe and Martine Époque. 2005. The Dancing Genome Project: Generation of a Human-Computer Choreography Using a Genetic Algorithm. In *Proceedings of the 13th Annual ACM International Conference on Multimedia*. ACM, Hilton Singapore, 555–558. doi:10.1145/1101149.1101276
 - [29] Min Kyung Lee, Daniel Kusbit, Anson Kahng, Ji Tae Kim, Xinran Yuan, Alissa Chan, Daniel See, Ritesh Noothigattu, Siheon Lee, Alexandros Psomas, and Ariel D. Procaccia. 2019. WeBuildAI: Participatory Framework for Algorithmic Governance. *Proceedings of the ACM on Human-Computer Interaction* 3, CSCW (Nov. 2019), 1–35. doi:10.1145/3359283
 - [30] Amalaswitha Leh, Dominik Endres, and Mathias Hegele. 2025. Dancing through the Uncanny Valley: On the Likeability of Model-Generated Dance Movements. *Psychology of Aesthetics, Creativity, and the Arts* (2025). doi:10.1037/aca0000726
 - [31] Ruilong Li, Shan Yang, David A. Ross, and Angjoo Kanazawa. 2021. AI Choreographer: Music Conditioned 3D Dance Generation With AIST++. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 13401–13412.
 - [32] Yimeng Liu and Misha Sra. 2024. DanceGen: Supporting Choreography Ideation and Prototyping with Generative AI. In *Proceedings of the 2024 ACM Designing Interactive Systems Conference (DIS '24)*. Association for Computing Machinery, New York, NY, USA, 920–938. doi:10.1145/3643834.3661594
 - [33] Lian Loke and Toni Robertson. 2013. Moving and Making Strange: An Embodied Approach to Movement-Based Interaction Design. *ACM Trans. Comput.-Hum. Interact.* 20, 1 (April 2013), 7:1–7:25. doi:10.1145/2442106.2442113
 - [34] Wendy E. Mackay and Joanna McGrenere. 2025. Comparative Structured Observation. *ACM Trans. Comput.-Hum. Interact.* 32, 2 (April 2025), 14:1–14:27. doi:10.1145/3711838
 - [35] Maddock Meredith and Steve Maddock. 2001. Motion Capture File Formats Explained. *Department of Computer Science, University of Sheffield* 211 (2001), 241–244.
 - [36] Masahiro Mori. 1970. "The Uncanny Valley". *Energy* 7, 4 (1970), 33–35.
 - [37] Maria Rita Nogueira, Paulo Menezes, and José Maçães de Carvalho. 2024. Exploring the Impact of Machine Learning on Dance Performance: A Systematic Review. *International Journal of Performance Arts and Digital Media* 0, 0 (2024), 1–50. doi:10.1080/14794713.2024.2338927
 - [38] A. Michael Noll. 1967. Choreography and Computers. *Dance Magazine* XXXXI, 1 (Jan. 1967), 43–45.
 - [39] Victor Paredes, Jules Françoise, and Frederic Bevilacqua. 2022. Entangling Practice with Artistic and Educational Aims: Interviews on Technology-based Movement-Sound Interactions. In *International Conference on New Interfaces for Musical Expression*. doi:10.21428/92fbeb44.5b9ac5ba
 - [40] Mariel Pettee, Chase Shimmmin, Douglas Duhaime, and Ilya Vidrin. 2019. Beyond Imitation: Generative and Variational Choreography via Machine Learning. arXiv:1907.05297 [cs, stat] doi:10.48550/arXiv.1907.05297
 - [41] Matthias Plappert, Christian Mandery, and Tamim Asfour. 2016. The KIT Motion-Language Dataset. *Big Data* 4, 4 (Dec. 2016), 236–252. doi:10.1089/big.2016.0028
 - [42] Sigal Raab, Inbal Leibovitch, Guy Tevet, Moab Arar, Amit H. Bermano, and Daniel Cohen-Or. 2024. Single Motion Diffusion.
 - [43] Yvonne Rogers. 2011. Interaction Design Gone Wild: Striving for Wild Theory. *Interactions* 18, 4 (July 2011), 58–62. doi:10.1145/1978822.1978834
 - [44] Thecla Schiphorst. 2011. Self-Evidence: Applying Somatic Connoisseurship to Experience Design. In *CHI '11 Extended Abstracts on Human Factors in Computing Systems (CHI EA '11)*. Association for Computing Machinery, New York, NY, USA, 145–160. doi:10.1145/1979742.1979640
 - [45] Fiona Ssozi-Mugarura, Thomas Reitmaier, Anja Venter, and Edwin Blake. 2016. Enough with 'In-The-Wild'. In *Proceedings of the First African Conference on Human Computer Interaction (AfriCHI '16)*. Association for Computing Machinery, New York, NY, USA, 182–186. doi:10.1145/2998581.2998601
 - [46] John D. Sullivan, Sarah Fdili Alaoui, Pierre Godard, and Liz Santoro. 2023. Embracing the Messy and Situated Practice of Dance Technology Design. In *Proceedings of the 2023 ACM Designing Interactive Systems Conference*. ACM, Pittsburgh PA USA, 1383–1397. doi:10.1145/3563657.3596078
 - [47] Guy Tevet, Sigal Raab, Brian Gordon, Yonatan Shafir, Daniel Cohen-Or, and Amit H. Bermano. 2022. Human Motion Diffusion Model. arXiv:2209.14916 [cs] doi:10.48550/arXiv.2209.14916

- [48] Jonathan Tseng, Rodrigo Castellon, and Karen Liu. 2023. Edge: Editable Dance Generation from Music. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 448–458.
- [49] Shuhei Tsuchida, Satoru Fukayama, Masahiro Hamasaki, and Masataka Goto. 2019. AIST Dance Video Database: Multi-Genre, Multi-Dancer, and Multi-Camera Database for Dance Information Processing. In *ISMIR*, Vol. 1. 6. doi:10.57765/2003396
- [50] Carnegie Mellon University. [n. d.]. CMU Graphics Lab - Motion Capture Library.
- [51] Benedikte Wallace. 2023. *AI-generated Dance and The Subjectivity Challenge*. Ph. D. Dissertation. University of Oslo.
- [52] Benedikte Wallace and Charles P. Martin. 2022. Embodying the Glitch: Perspectives on Generative AI in Dance Practice. arXiv:2210.09291 [cs]
- [53] Benedikte Wallace, Charles P. Martin, Jim Tørresen, and Kristian Nymoen. 2021. Exploring the Effect of Sampling Strategy on Movement Generation with Generative Neural Networks. In *Artificial Intelligence in Music, Sound, Art and Design*, Juan Romero, Tiago Martins, and Nereida Rodríguez-Fernández (Eds.). Vol. 12693. Springer International Publishing, Cham, 344–359. doi:10.1007/978-3-030-72914-1_23
- [54] Benedikte Wallace, Kristian Nymoen, Jim Tørresen, and Charles Patrick Martin. 2024. Breaking from Realism: Exploring the Potential of Glitch in AI-generated Dance. *Digital Creativity* 35, 2 (April 2024), 125–142. doi:10.1080/14626268.2024.2327006
- [55] Edward Yourdon. 1989. *Structured Walkthroughs: 4th Edition*. Yourdon Press, USA.

A Voguing and Dancehall Movements Recorded with Motion Capture

Name of the step	Dance Style	Takes per dancer
Bogle Step	Dancehall	1
Catwalk	Voguing	4
DuckWalk	Voguing	4
Gun Steps	Dancehall	1
Gyal Step	Dancehall	1
Hands Performance	Voguing	3
Improvisation	Dancehall	1
Nuh Behavior	Dancehall	1
Nuh Linga	Dancehall	1
Scandal	Dancehall	2
Wait Yuh Turn	Dancehall	1

Table 2: Dance moves chosen by the collective and recorded during the motion capture session. Each dancer of the collective recorded these steps, resulting in 40 motion capture recordings.

B Technical Details on the Model

B.1 Implementation of the Model

We implemented a movement generation model based on the *Chor-rnn* architecture [11], a Mixture Density Recurrent Neural Network (MDRNN) that has been extensively used in the literature. The model’s architecture consists of a recurrent neural network using long short-term memory layers that estimates the parameters of a mixture density network (MDN). The MDN used 4 Gaussian kernels. For the recurrent part of the model, we used 3 long-short term memory layers (LSTMs) with hidden sizes 1024, 512 and 256, which are the values used in [53]. They are then connected to 3 parallel separate linear layers estimating the parameters of the MDN, namely the mixture coefficients, the means and the variances. The model generates movement by estimating a frame from the previous ones. Because it uses an MDN, it is stochastic: the

model estimates the parameters of the MDN, and to actually get the generated frame, the resulting distribution is sampled. For motion representation, the model encodes movement as joint angles using an Euler representation.

B.2 Training the Model

We resampled the animation data from 100 to 30 frames per second to train the model. We used an Adam optimizer with a learning rate of 3.10^{-4} and a batch size of 256. We used the negative likelihood loss. The optimizer was given batches of data from 10-second-long sequences (300 frames). The loss was computed for each of the 299 frames’ distribution estimated from the previous ones and averaged over the batches. The training data comprises all of the motion capture data we recorded with the collective. The sequences are cut into 10-second overlapping sequences (over one frame only). We kept five files from the motion capture data as a validation set, which was used as a reference to check the progress of the training loss against the same loss computed on that test set and perform early stopping when the loss stopped improving for five epochs. We trained our models on a single V100 GPU.

A checkpoint of the model’s training was saved every 10 epochs, and intermediate generated dance movements were saved using Korai every 5 epochs to monitor and understand the training process. These were always generated from the same seeds, which were randomly chosen from the validation set at the start of the training. Training stopped when the validation loss did not improve after 10 epochs. This resulted in 47 epochs of training.

C Implementation of Korai

Korai is a web application built using Marcelle [19], an open-source toolkit for programming interactive machine learning applications. The main interface is a web page built with the Svelte framework. It leverages the Three.js⁷ and Threlte⁸ libraries to render the 3D models in real-time in the browser. Each source animation is loaded as a BVH file and rendered as a Three.js object.

For movement generation, the web client communicates with a Python script for inference using WebSocket communication. This process is mediated by the backend system of Marcelle, a Node.js server. When the user requests a new sequence, the generation parameters are passed to the Python script (seed time and duration, sequence length). The script retrieves the source animation, extracts the seed and performs the inference. The resulting motion sequences are stored as a BVH file. The generated movements are also stored in the backend. Once this process is complete, the client is notified and the interface is updated in real-time to render the new output. The application’s state is stored in a MongoDB database managed by the backend, ensuring that refreshing or reopening the application after closing it will restore the same state.

In practice, we used an 11th-generation i5 CPU-based personal computer to run the application. This hardware setup is enough to bring low-latency generation and keep the exploration interactive, with about 10 seconds of inference for a 30-second generated sequence.

⁷<https://threejs.org/>

⁸<https://threlte.xyz/>