



Arthur Kuhn
arthur.escabasse@gmail.com
Lab. Paragraphe, ED CLI,
Université Paris 8, Paris, France

Munda M'ennuie is a live noise performance investigating the co-direction between human memory and machine learning in sound creation. The project confronts the opposition between a human inability to imagine beyond remembered sonicities and the unconvincing imitations produced by AI when attempting to implement the same sonicities from its description. Rather than relying on existing models, this work explores a personally curated approach through training a custom model on the soundtrack of the game *Legacy of Kain – Soul Reaver* (1999), sourced from gamerips—a type of fan recording. The performance explores technostalgia and foreverism, embodying a shift from archival preservation to a poetics of remix where glitches and emulation markers constitute an aesthetic vocabulary. By deliberately misaligning the dataset with its intended technical use, the work generates noise textures that evoke foggy remembrance mixed with abrupt sonic shifts, creating drone-like soundscapes. This piece aims to un-cover machine learning as an inhuman meeting point, exploring the contingencies' archipelago of the model's latent space. Rather than pursuing faithful reproduction or AI promises, it focuses on computational dissolving and the aesthetic potential of in-between spaces.

Introduction

Munda M'ennuie is a live noise performance project, built around the real-time use of an audio convolution Machine Learning (ML) model, trained on a videogame soundtrack. It emerged from a tension regarding the infusion of past aesthetical experiences in new materials, when trying to establish a collaboration between an Artificial Intelligence (AI) technology and a human composer.

In the context of this performance, the AI is considered a co-director. An agent to meet on its own terms, offering and conditioning a specific mode of sonic exploration. Though a human performer ultimately enables the performance, they must do so while dealing with the technical specification and *ontology* (Smith 2019) of ML. Such a setting was already experimented in my previous piece *melodia atomizacji – asemic dissolution ritual*.¹ In this previous performance, the model used was provided already trained, and my sonic work was mostly informed by an interest in the *Harsh Noise Wall* (HNW) movement,² as well as defined by an over-

Keywords: Noise Performance,
Sound Art, Technological Nostalgia,
Artificial Intelligence,
Computational Aesthetics.

DOI [10.34626/2026_xcoax_040](https://doi.org/10.34626/2026_xcoax_040)

1. Kuhn, Arthur. s. d. *melodia atomizacji – Asemic Dissolution Ritual*. 24:32.

<https://arthurkuhn.bandcamp.com/album/melodia-atomizacji-asemic-dissolution-ritual>

2. The *Harsh Noise Wall* movement is a sub-genre of harsh noise music, consisting of “a literal consistent, unflinching and enveloping wall of monolithic noise”. For more informations, see: Williams, Russell. 2014. “LIVE REPORT: Harsh Noise Wall Festival III.” *The Quietus*, May 22.

<https://thequietus.com/news/live-report-hnw-fest-iii/>

arching fictional work designed to explore ML. Continuing the investigation on a human-AI co-conducted practice, *Munda M'ennuie* allows me to go further, taking charge of the dataset curation and training phases, developing a less linear structure of live composition, and exploring a more personal theme of nostalgia.

Technical Context

Munda M'ennuie uses a custom RAVE model through its dedicated decoder for real-time use: the [nn~] Max/MSP object.³ Both of these technologies were developed by the ACIDS research team at IRCAM, Paris.⁴ It is important to state that this idea of human-AI real-time co-creation is far from an exclusive inquiry. Several other artists and researchers, amongst them the members of the ACIDS team themselves, have experimented in this direction, as can be seen in the video recordings of their *ACID Workshop beta* event from 2022.⁵ My choice of a RAVE model and its wrapper object for this project was informed by previous personal creations, but also because of RAVE's ease of use and capability to be trained on relatively small amount of data for a great audio quality, allowing me to pursue my goal of *Do-It-Yourself ML* (Kuhn 2024). Because these characteristics exist thanks to Chemla-Romeu-Santos, Caillon and Eisling's work, I feel the need to address this lineage. The originality of my project, I hope, lies in a creative misappropriation of their technologies, and the themes it allows me to explore.

Fig. 1. *munda m'ennuie* (live) illustration visual.



3. Caillon, Antoine, and Philippe Esling. 2022. *Streamable Neural Audio Synthesis With Non-Causal Convolutions*. Version 1. <https://doi.org/10.48550/ARXIV.2204.07064>

4. For a presentation of the research team and their different projects, see the ACIDS-IRCAM website: <https://acids-ircam.github.io/>

5. See the official event page on the IRCAM website: <https://forum.ircam.fr/agenda/acids-workshop-beta/detail/>

Memory and Archive

After performing *melodia atomizacji* several times, I sought to try different methods for AI-assisted sound creation. Trying out several tools, from Suno, the dominant platform for prompt-based musical generation,⁶ to ChatGPT by asking it to generate sound clips after a discussion of their qualities. In both cases, I was confronted to my own incapacity to go beyond the evocation of past “sonicities” (Ernst 2016); fixating on the technical apparatus of a sound—like the hollowness or blunt quality of a square sinewave when shifting its duty cycle—moreso than imagining a composition. The AIs, in return, proved either too close to any given reference (Suno), or straight-up bad at conveying what it talked about eloquently into sound (ChatGPT). Hence, exploring deeper RAVE models and neural audio convolution seemed more interesting.

Though, that preliminary work proved useful, as it gave the performance its major themes. First, “technostalgia” (Habib 2024) designates a tendency to remember and cherish the memory of a technological support more than of a cultural content: being nostalgic of the 8mm color grading and noisy image, not of a specific movie; of an audio chipset sound, not a particular music. Then, this project also reflects on “foreverism” (Tanner 2023): the transformation of our collective archival fascination—preservation-oriented, aiming for a mimical quality towards the original—into a poietic of perpetual remix and re-activation. Together, these two questions outlined a framework, the exploration of a re-processed “noise mapping” (Parikka 2011) of my obsessive technicalities.

Dataset Curation

Equipped with this framework, rose the question of the dataset. Foreverists objects in themselves, datasets stand on the line between meticulous preservation, well-ordained archival tagging, and memorial proceduralization, remixing normalization. I stipulated the following criteria for choosing a dataset:

1. It must come from a personal memory;
2. It must be chosen as much, if not more, for its technicalities than its aesthetic;
3. It must require work in terms of collecting and standardization;
4. It must be close, but not quite adapted to the model’s training it will be used for.

The last criterion is here for both aesthetic reasons—I am always more attracted to the sound of things broken—and poietic ones: in accordance with my theoretical questions, I was aiming for an un-naturalistic, anti-archival sort of curation. A bricolage of artistic investigation with the goal to explore the in-between space birthed from computational dissolving and remapping, not an exercise in faithful recreation. Fulfilling

6. Suno’s official website, allowing for prompt-based musical generation: <https://suno.com>

all of these criteria, I chose to work from the soundtrack of the PlayStation's videogame *Legacy of Kain – Soul Reaver* (1999).⁷

1. This game is a highly personal memory. A game which I began over and over. It encapsulates very significant moments of my own artistic construction;

2. Yet I hardly remembered its soundtrack. I remembered its mood, its digital-sounding quality (lots of acidic, digital-sounding reverb), but no specific melody;

3. Particularly adequate, *Soul Reaver's* soundtrack isn't commercially available.⁸ The only way for fans to get it was through gamerips: people playing the game and recording its sound. A practice helped by the fact that *Soul Reaver* is one of the very first games where one could go through almost all of it without loading breaks. Hence, the material I trained my model on is actually the creation of a highly dedicated player following a very specific protocol: starting in the game's first area, crossing it all over while killing all the enemies, going back at the beginning, and revisiting the whole environment; while regularly staying in place to let the music of each area loop long enough (Figure 2).

Fig. 2. *Legacy of Kain – Soul Reaver's* Gamerip cover art. Gamerip by Raina, the v8 indicates that it is their eighth version of it.



4. The soundtrack carries a similar mood, even sharing a common beat, for a large part. It comprises a fairly small number of timbres and textures, with each of the game areas' score coming in two versions; one

⁷ For more information, see the game's wikipedia page:
https://en.wikipedia.org/wiki/Legacy_of_Kain:_Soul_Reaver

⁸ The wikipedia page provided above mentions that "Music from both *Soul Reaver* and *Soul Reaver 2* was released on a promotional soundtrack in 2001". This actually only refers to a bonus feature on the *Soul Reaver 2* PlayStation 2 DVD that let you listen to a few selected tracks. No soundtrack CD were ever officially commercialized.

being drenched in the same hazy reverb each time. Still, it is not a clean series of samples of the same instrument, proving to be on the edge of untrainability for a RAVE model. Offering a fine, broken enough, quality.

Training took place in Vancouver during the Artificial Muse residency,⁹ co-piloted by the French Canadian Institute and the Simon Fraser University, through the Metacreation Lab and its director Philippe Pasquier.¹⁰ Several training and tryouts were needed, while I ended up following Philippe Pasquier's advice of programmatically cutting the sound files in overlapping 2seconds samples, mixed with approximately 60% of the total runtime of silence. The goal being to, following Philippe's advice, "teach my model to be silent, sometimes".

Fig. 3. Preview from *munda m'ennuie* performance visuals.



Signals and Synths

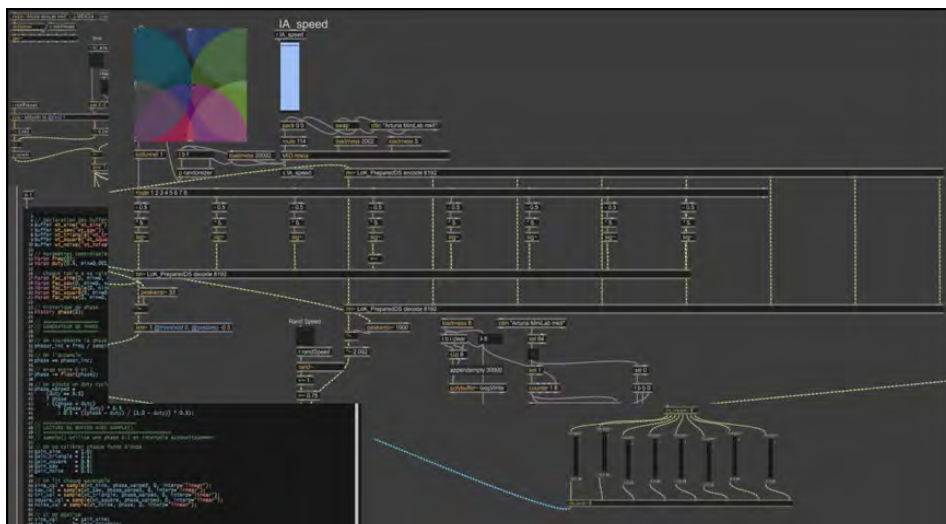
With the model satisfyingly trained, it moved to Max/MSP and the co-creation phase. To better explain my following choices, I will now offer a brief rundown of the audio convolution [nn~] performs; but I want to underline that I do so with a self-taught, empiric form of knowledge. During the model's training, the audio provided is progressively down-sampled and its evolution in the time-frequency domain is converted in unnamed latent dimensions, which are discovered by the model. These dimensions are notoriously hard to figure in human terms, but would sometimes loosely correlate to the audio's dynamics/energy shifts, spectral brightness, periodical or noisy nature, and other specifications fundamentally impossible to parametrize. Each of these dimensions serve to map out a continuous latent space, in which each set of coordinates correspond to a probability distribution of frequencies, allowing for the reconstruction of a sound. More importantly, this latent space is specific to each model, depending on the dataset it was trained on. Its number of dimensions, what they represent, and what type of sound it will be able to recreate is unique. But, for any RAVE model, the convolution will

9. Artificial Muse residency, a program by the French Canadian Institute: <https://www.institutfrancais.ca/en/nos-residences/artificial-muse-vancouver-surrey>

10. See the Metacreation Lab website: <https://www.metacreation.net/>

occur from a signal given to the model that will endure the same down-sampling to find a corresponding set of coordinates, and then recreate the closest resembling result according to its abilities. In short, it will try to create the version of what it knows that resembles the most what it is given. When used within the [nn~] object, this operation happens in real-time, although with an adjustable but unavoidable latency. It is also important to note that [nn~] can either be used as a simple translator: sound → object → sound, or as a more flexible pair of objects: sound → object:encoder → signals → object:decoder → sound.

Fig. 3. Edited screenshot of *munda m'ennuie*'s main patch. On the right are the [nn~] signals and treatment.



The signals added in-between the encoder and decoder correspond to an automatically selected portion of the latent dimensions, allowing them to be directly modified. This second option also allows you to multiply the number of both encoders or decoders. From that, I envisioned three ways of working with [nn~]:

1. Sending it a sound that possesses qualities I wish to see in its recreation, without certainty of how the model will react, or even if these qualities were indeed preserved in the model's dimensions.
2. Same, but also adding scalars to the dimensions as the encoded sound is about to be decoded. These scalars being constant, fluctuant or random; added, multiplied or subtracted; applied on one dimension, several, or all of them.
3. Directly sending sets of coordinates to explore the content of the model latent space in the decoder.

I decided to combine solutions 2 and 3.

In order to produce a sound that could serve as a basis, I coded a wavetable synth, equipped with an ADSR envelope filter, a waveshaping distortion and a state-variable filter.¹¹ This instrument allowed me to produce sound I deemed close enough to the basic material of *Soul Reaver*'s soundtrack, while offering enough variations to push the

11. For more information on these different techniques, see the Max/MSP documentation, respectively: *Wavetable Oscillator*: https://docs.cycling74.com/learn/articles/05_mspbasicchapter03/ *adsr~*: <https://docs.cycling74.com/reference/adsr~/> *Waveshaping*: https://docs.cycling74.com/learn/articles/07_samplingchapter05/ *svf~*: <https://docs.cycling74.com/reference/svf~/>

model into unexpected reactions. I also considered it an interesting opposition to have an AI technology, based on unsupervised training and inhuman design, dialog with a highly parametric construct, based entirely on simulated equations. It was as if the expert systems of the eighties were exchanging with contemporary AI.

Both the synth's original sound and [nn~] interpretation of it are then combined through real-time Fast Fourier Transforms, allowing me to quite literally fuse their sonic structures, phases and amplitudes, while controlling which one is more prevalent at every second; as if balancing between my version of *Soul Reaver* sound's recreation, and the AI's, while keeping both intrinsically linked. Parallely, an object is continuously generating 8 axis coordinates, slowly moving from one random point to the next in a literal exploration of the latent space. This sound, although continuously produced, is enveloped by another ADSR, mirroring the synth's one settings, with varying amount of delay. The goal being to reintroduce harsher, sudden and sharp elements that plays with the somewhat elongated perception of time I mobilize during the performance.

Fig. 5. Lecture performance on musical creation and AI, for the *Master in Digital Media's* students. *Centre for Digital Media*, Vancouver, 2025



Both the synth's original sound and [nn~] interpretation of it are then combined through real-time Fast Fourier Transforms, allowing me to quite literally fuse their sonic structures, phases and amplitudes, while controlling which one is more prevalent at every second;¹² as if balancing between my version of *Soul Reaver* sound's recreation, and the AI's, while keeping both intrinsically linked. Parallely, an object is continuously generating 8 axis coordinates, slowly moving from one

12. Respectively using the [fft~] and [gen~] functions of Max/MSP. See the software documentation for more details: https://docs.cycling74.com/userguide/frequency_domain/

random point to the next in a literal exploration of the latent space. This sound, although continuously produced, is enveloped by another ADSR, mirroring the synth's one settings, with varying amount of delay. The goal being to reintroduce harsher, sudden and sharp elements that plays with the somewhat elongated perception of time I mobilize during the performance.

Fig. 6. Photographs from performances during the *des objets étranges* exhibition. Ateliers Alain le Bras, Nantes. September and October 2023.



Conclusion

Expanding on that, I will briefly conclude by some words on the performance's aesthetic in itself. I purposefully avoided describing its composition while giving—I hope—enough elements to point in the direction of the kind of moment I am interested in. *Munda M'ennuie* is a mostly improvisational performance and cannot be otherwise, since the other, machinic half of the project tends to never give out the same result twice. Furthermore, I do not think that its subject matter should invoke a place of tranquility, but rather of turmoil. Hence, what I hope to achieve is a drone noise project that is equally foggy, melancholic and abrasive. As the title implies, cleanliness is not what I am after.

Audio recording—Preview medley:
<https://arthurkuhn.bandcamp.com/track/munda-mennuie-preview-medley>

References

**Abuzurairq, Ahmed M.,
and Pasquier, Philippe.**

2025. « Explainability-in-Action: Enabling Expressive Manipulation and Tacit Understanding by Bending Diffusion Models in ComfyUI ». Version 1. Prepublication, arXiv. <https://doi.org/10.48550/ARXIV.2508.07183>

Bogost, Ian.

2014. "Inhuman." In *Inhuman Nature*, version 1, by Jeffrey Jerome (Ed) Cohen. Punctum books. PDF, 166pp. <https://doi.org/10.21983/P3.0078.1.00>

Bonnet, François J.

2022. *Les mots et les sons: un archipel sonore*. L'éclat-poche 50. Éditions de l'Éclat.

Caillon, Antoine, and Philippe Esling.

2022. *Streamable Neural Audio Synthesis With Non-Causal Convolutions*. Version 1. <https://doi.org/10.48550/ARXIV.2204.07064>

Ernst, Wolfgang.

2016. *Sonic Time Machines: Explicit Sound, Sirenic Voices, and Implicit Sonicity*. 1re éd. Amsterdam University Press. <https://doi.org/10.1017/9789048528479>

Habib, André.

2024. *Technostalgie et mélancolie de la technique : actualités de l'obsolète*. EUR ArTeC. 1:10:45. <https://vimeo.com/1053755276?share=copy>

Huyghe, Pierre-Damien.

2023. *Poussées techniques, conduites de découverte: À quoi tient le design*. De l'incidence éditeur.

Kuhn, Arthur.

2024. "Do It (as) Your(other)self. Bricolage and Heteronyms, Artistic Strategies Regarding Artificial Intelligence." Paper presented at *3ai24. Art in the Age of AI*, Corfou, Greece, May 27. <https://3ai-24.sciencesconf.org/541784/document>

Parikka, Jussi.

2011. "Mapping Noise: Techniques and Tactics of Irregularities, Interception, and Disturbance." In *Media Archaeology: Approaches, Applications, and Implications*, edited by Jussi Parikka and Erkki Huhtamo. University of California Press, 256-277.

Smith, Brian Cantwell.

2019. *The Promise of Artificial Intelligence: Reckoning and Judgment*. The MIT Press.

Tanner, Grafton.

2024. *Foreverism: quand le monde devient un jour sans fin*. Façonnage Éditions.