



Sofia Taipa
taipa.sofia@gmail.com
Visual Artist, Goldsmiths, University of
London, London, UK

Sui Generis exposes how machine learning algorithms reduce human complexity to statistical abstractions. I approach these tools not as neutral agents, but as judgement and political machines. The work materialises this position through an algorithm that infers (or perhaps, fails to infer) how much I like or dislike a given word on a scale from 0 to 100. This algorithm is not rigorous or useful, but an ironic device used to highlight the subjective nature of these technologies. My custom machine learning model is materialised as a dictionary, a symbol of universal knowledge. Each percentage, each word classified by the model, exposes its limits: the impossibility of true personalisation. Displayed within a clinical setting, the audience is invited to examine the book using white gloves and a digital microscope, entering a loop of scrutiny where subjective data is objectified, only to become subjective again through the viewer's judgement. The algorithm ultimately reflects countless unseen contributors embedded within its model. Where do other people's preferences surface in these supposed personal classifications? How does this homogenised entity shape my own sense of identity?

Keywords: Machine Learning,
Homogenisation, Algorithmic Bias,
Individuation, Latent Space, Installation,
Artistic Research.

Concept and Research

Sui Generis emerged from a critical and conceptual investigation into the operational logic of machine learning algorithms, and how we can relate them to human decision-making. These technologies are built to make decisions. They are not objective, sterile, nor neutral instruments. Fundamentally, they are built to categorise, presume, and select: they are judgement machines. By processing vast datasets of human information and interaction, they do not eliminate subjectivity, but rather encode it. Since they are trained on the accumulation of human choices, data that is inherently subjective, they do not merely observe the world; they operationalise its biases. Human data is, by definition, biased. And so is machine learning.

Rooted in this subjective data, the training process for these models operates on mathematical thinking: a vectorisation that maps individuals and their data into a high-dimensional latent space. This reduces every subject to a statistical entity, a mere vector amongst millions. Conceptually, this latent space functions as a closed universe where the model holds all possible answers as predefined coordinates. Within this structure, identity is determined strictly by mathematical proximity. This produces a flattening: once crowded into this dataset, we become indistinguishable from the statistical average (Steyerl 2023). As Hito Steyerl notes regarding captcha authentication, the browser scans our history to compare it against a predefined image it already holds of us (Steyerl 2016). We are only recognisable as “correct data” if we are coherent within this contiguity. Meaning, we can only access our bank account or complete an online purchase if we remain the same as yesterday. If we diverge from this configuration—if we naturally evolve, the system concludes we are not ourselves. “It is an identification game” (Steyerl 2016), where we are identified by our similarity to our past selves, our social graph, or simply those who resemble our clicks online. Ultimately, we are no longer defined by our singularity, but by our proximity to our mathematical neighbours in the vector space.

The premise that we are only identifiable if we are similar to those around us forecloses the process of continuous evolution, what Gilles Deleuze describes as individuation (Deleuze 2001). For Deleuze, to become is to differentiate, to produce the new; yet these models reduce this potential to a crowd of sameness. In human decision-making, evolution is driven by emergence: the capacity to generate ideas and patterns that transcend prior conditions. However, a machine learning model cannot accommodate emergence. The latent space can never be an emergent space: because it is trained on past data, its universe is finite. Every output it generates is merely a retrieval of a pre-existing coordinate within its latent space. It cannot create the new; it can only locate the likely.

By granting these judgement machines the authority to decide, we fundamentally alter the definition of truth. For the algorithm, truth is not derived from the unfolding of lived experience, but from statistical

alignment with its dataset. Truth becomes defined by consistency. Because these algorithms are tested strictly on their ability to reproduce history, they inherently penalise divergence: the future they predict is mathematically bound to equate to the past. They effectively function as a stochastic parrot (Bender et al. 2021), where meaning is not created but statistically recycled. In this closed loop, the future equates to the past; it is automated, embedding all the mistakes and biases, while leaving no space for growth.

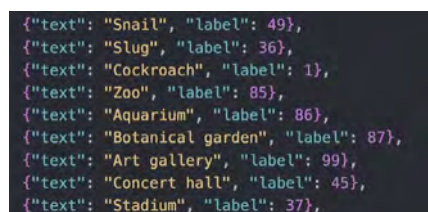
My Judgement Machine

This replacement of “likenesses with likelinesses” (Chun 2021) forms the conceptual basis of *Sui Generis*. The project functions as a critical interrogation of the custom judgement machine I built, constructed to expose the mechanisms of similarity and homogenisation. To this end, I developed a personal classification system designed not for accuracy, but for the amplification of bias. I created a small personal dataset of words, assigning each one of these words a value between 0 and 100 (Fig. 1), based on intuitive preference. The rapidity and subjectivity of this classification task were intentional and crucial. Even with an attempt at a methodical approach, this process remains inherently absurd. Quantifying volatile human attributes into stable data points is inherently impossible, yet it remains the foundational logic upon which these machines operate.

This subjective dataset was then used to fine-tune a pre-trained language model (DistilBERT). Crucially, this model is not a blank slate, it has already learned the relationships between words, their semantic weights, and conceptual associations from a massive corpus of text representing millions of users online. In this process, I did not teach the machine what words are or how they relate, that knowledge is inherited from the crowd of its original dataset. My fine-tuning process, associating words with my personal preference through numerical classification, merely rearranges these pre-existing connections to fit my personal rules. Consequently, the resulting “personal” machine can never be truly individual. Because its foundational logic is built upon the aggregated preferences and biases of millions of other voices, the model remains irrevocably tethered to the external world, making true isolation, and therefore true personalisation, impossible.

The central artefact, the dictionary, acts as a physical critique of this false personalisation. As a symbol of universal knowledge and definition, the dictionary form reinforces the project’s central thesis: that I am generating yet another instance of standardised universality, rather than true individuality. Inside, instead of definitions, words are paired with their predicted classification scores (Fig. 2). Each of these numbers exposes the limitations of the model, freezing fluid preferences into predefined statistical certainties.

The installation environment reframes this object not as a volume to be read, but as evidence to be examined. Visitors are invited to re-



```
{"text": "Snail", "label": 49},
{"text": "Slug", "label": 36},
{"text": "Cockroach", "label": 1},
{"text": "Zoo", "label": 85},
{"text": "Aquarium", "label": 86},
{"text": "Botanical garden", "label": 87},
{"text": "Art gallery", "label": 99},
{"text": "Concert hall", "label": 45},
{"text": "Stadium", "label": 37},
```

Fig. 1. Section of my personal dataset (words and their classifications).

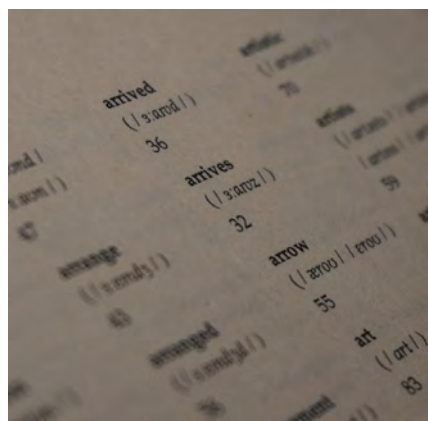


Fig. 2. Close up of the dictionary.

view the book through a digital microscope in a setting that mirrors a clinical laboratory, using white gloves to enforce a deliberate, tactile engagement. These choices disrupt the usual affordances of reading, promoting a posture of scrutiny. Under the microscope, the text becomes a specimen. This creates a continuous cycle of judgement, where subjective data is objectified by the machine, only to become subjective again through the audience's interrogation.

Fig. 3. Installation with book, digital microscope and white gloves.



Conclusion

Ultimately, *Sui Generis* does not seek to fix the bias of the machine, but to embrace the dictionary as a tangible iteration of the stochastic parrot. The project demonstrates that the aggregation generated during the training process is not an aggregation of people or their data, but an abstraction of it. As the model reduces us to mathematical probabilities, it reveals a critical flaw in the pursuit of total knowledge: the data provided by these models can never fully capture reality, for they operate by adding a layer of calculation to an existing experience we already understand.

The model I fine-tuned inherits not only the biases of its previous training, but also the imprints of countless voices across the internet: every individual who has ever contributed to its datasets. This raises urgent questions: What does the statistical output given by the model tell me about other people's preferences? In what ways are they embedded in these supposed personalised classifications? And crucially, how is my own identity being influenced by the homogenised entity built by the model? We are naturally connected by our similarities, but these similarities should not define us. Unfortunately, as this work aims to highlight, these datasets mine our data not to identify who we are, but to reduce us to our relation to others.

We cannot uninvent these technologies, nor can we simply reject them. Instead, we must learn to use them as a mirror. As observed in *Sui Generis*, the model often hallucinated, amplifying my preferences into extreme values that defied the given scale. This exaggeration suggests that these systems do not merely replicate; they amplify their biases. Rather than using them to generate static outputs to predict a future that is merely a recycled past, we can leverage them to understand our present. Their predictions provide a lens through which we can better understand where our biases lie. We need to lay these predictions against the underlying data that only we can understand, shifting the focus from automated categorisation to a critical reflection on ourselves.

References

Bender, Emily M., Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell.

2021. "On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? 🦜." In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–23. New York: Association for Computing Machinery. <https://doi.org/10.1145/3442188.3445922>

Chun, Wendy.

2021. *Discriminating Data: Correlation, Neighborhoods, and the New Politics of Recognition*. The MIT Press.

Deleuze, Gilles.

2001. *Difference and Repetition*. Athlone Contemporary European Thinkers. Continuum.

Steyerl, Hito.

2023. "Mean Images." *New Left Review*, nos. 140/141 (March–June): 82–97. <https://doi.org/10.64590/uhm>

Steyerl, Hito.

2016. "Talk: Hito Steyerl's Why Games? Can an Art Professional Think? (June 2016)." GAMESCENES (blog). August 06, 2016. Accessed 31 May 2024. <https://www.gamescenes.org/2016/08/talk-hito-steyerls-why-games-can-an-art-professional-think-june-2016.html>