



Delft University of Technology

Document Version

Final published version

Licence

CC BY-NC

Citation (APA)

van der Burg, V. (2026). *The Material Gaze: Reflective AI as a Design Practice*. [Dissertation (TU Delft), Delft University of Technology]. <https://doi.org/10.4233/uuid:fb78a0d0-0eb9-46e9-b7af-d54bca8796ca>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

In case the licence states "Dutch Copyright Act (Article 25fa)", this publication was made available Green Open Access via the TU Delft Institutional Repository pursuant to Dutch Copyright Act (Article 25fa, the Taverne amendment). This provision does not affect copyright ownership.
Unless copyright is transferred by contract or statute, it remains with the copyright holder.

Sharing and reuse

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

This work is downloaded from Delft University of Technology.

absurd 0.

The Material Gaze

Reflective AI as a Design Practice

Vera van der Burg

The Material Gaze

Reflective AI as a Design Practice

Dissertation

for the purpose of obtaining the degree of doctor
at Delft University of Technology
by the authority of the Rector Magnificus
Prof.dr.ir. H. Bijl
chair of the Board for Doctorates
to be defended publicly on
Thursday 17th, September, 2026 at 15:00

by

Vera van der BURG

This dissertation has been approved by the (co)promotors.

Composition of the doctoral committee:

Rector Magnificus	chairperson
Prof. dr. P.A. Lloyd	Delft University of Technology, promotor
Dr. R.S.K. Chandrasegaran	Delft University of Technology, copromotor
Dr. A.A. Akdag Salah	Utrecht University, copromotor

Independent members:

Prof. dr. G.W. Kortuem	Delft University of Technology
Dr. O. Kudina	Delft University of Technology
Prof. dr. C. Groth	Aalto University, Finland
Dr. H.K.G. Andersen	Eindhoven University of Technology, the Netherlands
Dr. W. Odom	Simon Fraser University, Canada
Prof. dr. ir. A. Bozzon	Delft University of Technology, reserve member

Printed by: De Kijm
Graphic design: Ward Goes
Copyright: 2026, V. van der Burg

All rights reserved. No part of this publication may be reproduced or transmitted in any form or by any means, electronical or mechanical, including photocopying, recording or by any information storage and retrieval system without permission from the author.

An electronic version of this dissertation is available at:
<http://repository.tudelft.nl>

Part I.....Foundations 10

1.....Introduction 13

1.1.....	The problem with efficiency-driven AI in creative practice	21
1.2	Towards a Reflective AI design practice	23
1.3	Research approach and questions	24
1.4	Thesis structure	32
1.5	Dissemination.....	33

2. Methodological Orientations 35

2.1	Principle 1: Reorienting the purpose from efficiency to reflection	38
2.2.....	Principle 2: AI as material	42
2.3.....	Principle 3: First-person research methods	44

Part II Reflective AI across creative practices..... 48

3.....Ceci n'est pas une chaise 51

.....	<i>An emerging dialogue between an AI and a designer</i>	
3.1.....	Introduction	52
3.2	Background	53
3.3.....	A practice-based inquiry.....	56
3.4.....	Explorations in designer-AI dialogue	57
3.5.....	Seeing-moving-seeing: A model for designer-AI dialogue	73
3.6.....	Conclusion	75

4. Developing the Objective Portrait Method..... 79

4.1	Reading an image	80
4.2.....	Internal design worlds.....	82
4.3.....	The tension of the mirror metaphor	82
4.4.....	A comparative and collaborative research approach	83
4.5.....	Object Detection as a reflective tool	84
4.6.....	The reflective practice with AI: A five-step process.....	85

4.7.....	Insights: The process as the portrait.....	107
4.8.....	Discussion	113
4.9.....	Conclusion	115

5..... The Objective Portrait Workshop..... 123

5.1.....	Reflective AI use in design education.....	125
5.2	Adapting the method for design students	126
5.3	Methodology	128
5.4	The Objective Portrait workshop	132
5.5	Findings: Learning through and about AI	142
5.6	Discussion	154
5.7	Conclusion.....	158

6..... Why do I keep on coming back to this chair? 161

6.1	The backward gaze	163
6.2.....	The story of the chair	165
6.3.....	The role of memory in broader reflective AI practice	190

7..... From Text-to-Clay 195

7.1.....	Introduction	196
7.2.....	Material context	197
7.3.....	The ceramic-AI cycle	199
7.4	The residency	204
7.5.....	Reflections... ..	245
7.6	Conclusion.....	259

Part III..... Synthesis 266

8..... Discussion 269

8.1.....	The character of reflection.....	272
8.2	Learnings for a reflective design practice.....	275
8.3.....	Between particular insights and general method.....	279
8.4.....	Implications for AI technology and discourse.....	281

8.5.....	Limitations and conditions	284
8.6.....	Conclusion.....	285

9..... Acknowledgements..... 289

10..... About the Author 291

11..... Summaries 295

12..... Appendix 311

13..... Bibliography 321

Part I: Foundations

1. Introduction

*Prologue:
A different way of seeing*

The thesis that Artificial Intelligence (AI) can become a tool for reflection, or that it might serve a reflective purpose for designers, did not originate from a critique of technology. It did not stem from an urgent desire to fix the AI industry, nor from a direct reaction against the automation of creative labor. Rather, it emerged from a deep fas-

cination with humans. Specifically, it came from a question that has followed me across the disciplines I was schooled in: how do we see and give meaning to the world around us?

Years ago, I entered a master program in a design school not with any practical experience in art or design, but with a background in neuroscience. From this position, I found myself confused and often frustrated. While my peers possessed an intuitive fluency with materials, thinking through their hands, building things with ease, I felt like the odd one out. I became a design student who could analyze theoretically and talk forever, but struggled with developing a making practice. Yet, where I somewhat lacked the tacit knowledge of craft,

I developed an obsessive curiosity about perception and how designers see and work, especially experiencing from up close how the people in this school worked with materials and tools. I wanted to understand why we as people, but specifically designers, can see and feel a material and be attracted to it without knowing why. I wanted to know why a specific composition of an inanimate object can make us feel something, or remind us of something. I was touched by the fact that an artwork or design can strike us with emotions before we have cognitively processed it.

It was this fascination with perception, understanding the why behind our intuitive visual preferences, that led me to Artificial

Intelligence. In the mechanisms of computer vision, and more specifically object detection models, I saw a technological parallel to my own questions. These tools were designed to do something like seeing, to interpret and act upon things in the human visual world. I wanted to understand how that ‘seeing’ of those machines worked, and by extension, what that machine vision could reveal about how people see, or maybe even more interesting, how we not see.

As I began to train my own object detection models in order to understand them better, I quickly realized that these models possessed no independent vision. Their ‘sight’ was entirely dependent on its training data, collections of images that existed simply because someone,

somewhere, decided to place them in a folder and give them a name, becoming a simple digital structure from which those models learned.

This realization formed my research trajectory during my PhD. I saw that AI as a reflection of human decisions, decisions of the one(s) that trained and mastered these tools. It was somewhat like a mirror. While this metaphor is typically used to critique the technology, highlighting how it reflects the worst of humans in the form of societal biases and harmful prejudices that get embedded in these tools, I saw an opportunity in that reflective characteristic. Instead of deleting subjectivity in order to create a ‘correct’ tool, I started to train an object detection model explicitly on my own subjective

perspective. I used that subjective stance as a starting point, in order to *learn* from it, turning the act of training AI into a way of looking at my own looking. This approach to AI formed the very first basis of what I started to call a *Reflective AI design practice*, the practice that I will describe in this thesis: a practice with AI models with the goal to reflect on your own ways of seeing as a maker and designer.

As I developed this work over the course of this PhD, starting in 2020 and ending in 2026, the context in which I did so, changed. In the creative industry and beyond, working with AI tools transformed from a niche practice into a dominant cultural force, evolving from experimental models to the ubiquitous ‘Fast AI’ of commercial text and

image generators. Large models like ChatGPT, Gemini and Midjourney, among many others, entered the realm of creative work.

As the commercial AI industry pushed for efficiency, promising to write texts and create images for us through big models obscured by easy-access interfaces, my reflective approach transformed into a form of resistance. I realized that the drive for automation of creative processes undermined the very thing I found most valuable in my own work: using AI to create a space for reflection. Therefore, I felt it necessary to position Reflective AI as a critical alternative approach to this bigger force, even though my personal starting point came from a fascination for that technology instead of a critique itself.

Besides being a personally meaningful practice for me, the Reflective AI design practice also questions the drive to automate important parts of the creative process, like making an image or writing a text, and argues instead that the value of AI technologies lies in what it can reveal about its creator, especially when their creator is a designer. This dissertation is the result of staying with that realization while the world around it changed.

A note on terminology

Throughout this thesis, I use the term ‘Artificial Intelligence’ or ‘AI’ as an umbrella term encompassing the specific technologies I work with: object detection systems, image generation models, and the machine learning pipelines that underlie them. I find the term very imprecise, and sometimes misleading. It carries connotations of general intelligence, autonomy, and even consciousness that do not reflect the narrow, pattern-matching systems I actually engage with. Yet, ‘AI’ remains the term through which these technologies are discussed in design practice, popular discourse, and increasingly in daily life. Using more technically precise terms like ‘convolutional neural networks’ or ‘diffusion models’ throughout would create unnecessary

distance from the design contexts this research addresses. I therefore use ‘AI’ pragmatically, while asking readers to understand it as shorthand for the specific, trainable, image-based systems documented in the chapters that follow.

1.1 *The problem with efficiency-driven AI in creative practice*

The integration of artificial intelligence technologies as tools into design practice is increasingly influencing creative processes (Hwang, 2022; Wu et al., 2021). A logo can now be generated from a handful of keywords; a rough sketch transformed into a polished product in moments. These tools, predominantly understood as Generative AI, are typically presented as fixed, seamless interfaces optimized for efficient workflows. The interaction paradigm is deceptively simple: input a text-based prompt describing what you need, and receive immediate output. While this approach might align with productivity-driven goals and can be convenient at times when time is short, it also imposes significant constraints on how designers engage with AI in their creative work.

What is lost when speed is gained? If these tools promise to make creative work faster and easier, we should ask what disappears from creative practice when efficiency is prioritized. Donald Schön’s concept of *reflection-in-action* offers a way to think about this (Schön, 1988, Schön, 1983). Schön distinguishes

between knowing-in-action, the tacit knowledge that enables skillful practice, and reflection-in-action, the thinking that occurs while doing when a situation presents as uncertain or surprising. He argues that skillful practice involves a dynamic dialogue between practitioner and materials, where understanding emerges through iterative cycles of action and reflection. Designers learn by constantly adjusting their approach based on how materials respond to their interventions. The friction, resistance, and unexpected responses that materials provide are essential components of the learning process that enables creative work. Creative practice involves navigating situations of uncertainty and instability (Gaver et al., 2003; Cross, 2006), the uncomfortable space of not knowing, of failure. Good design practice thrives on this messiness. And it involves more than generating a thing: it is just as much thinking about the why and how behind the creation of a thing, which can be driven by specific fascinations, worldviews, and material preferences

of its maker. The question becomes whether efficiency-driven AI tools allow for this kind of engagement.

Commercial AI tools do not give designers the opportunity to actually explore AI as a design material (Benjamin et al., 2021; Steinbrück and Stankowski, 2023), partially due to how their actual material components are abstracted away. The polished interfaces—a text prompt box, a few parameter sliders, perhaps a style selector—conceal the complex technical pipeline underlying them. Critical stages such as data collection, annotation, and model training are rendered invisible, tucked away into backend processes that designers never see, much less are expected to engage with. Designers are left to interact only with the surface-level outcomes: the prompt, the generated images, the text completions, the predictions. The processes that produce these outputs—the thousands of training images, the human labor of annotation, the algorithmic decisions about what becomes a pattern—become irrelevant.

This abstraction, combined with the speed in which a perfect image or text is generated, has consequences for design practice. It creates what researchers identify as ‘design fixation’ (Jansson and Smith 1991; Hwang 2022), where the speed and completeness of output generation causes designers to become prematurely attached to initial results, constraining the broad exploration essential to a design process. It produces ‘creativity exhaustion’ (Li et al., 2024), as designers struggle to keep pace with the AI’s rapid generation, shifting from

creators of output to curators of machine output. Most critically for this research, it prevents designers from developing what Schön would recognize as a conversation with materials (Schön, 1983). There is little material to converse with when the most important components of the material itself are hidden. The cultural response shows this disconnect in practice, where the rise of the term ‘AI slop’ (Crawford, 2024), referring to some of the low-quality, generic, or mass-produced content generated by AI tools, signals a moment where some AI-generated content is increasingly seen not as progressive innovation but as degradation of creative standards.

By hiding complexity, these tools promise a so-called ‘democratized’ access to AI capabilities as it becomes easy and intuitive for people to use it (Caramiaux et al., 2025), which follows a pattern common to emergent digital tools where the actual material of technology (like the underlying code, for example) is obscured for the sake of usability. However, this promise comes at a cost as the opportunities for nuance, iteration, and reflection become constrained. There is little room for exploring alternative workflows situated in a designer’s specific practice, little opportunity to actually question the AI’s mechanisms from within the technology, and limited capacity to cultivate agency over the co-creative process. This leaves the creative practitioner with a tool that offers immense generative power but little room for the reflective, material engagement that gives their work meaning.

1.2 Towards a Reflective AI design practice

As I described in the prologue, my relationship with AI has always been different. I began by training object detection models on my own subjective perspective, and discovered that this process itself became a site for reflection. Therefore, I wondered if AI could become a reflection tool rather than one that would automate something from my own practice.

But what exactly is this reflection I am describing? The term is used widely, and many carry an intuitive sense of what it means, but to articulate what I mean more precisely, I turn to Phoebe Sengers’ work on Reflective Design. Sengers argues that reflection should be a core outcome of technology design, and defines it as:

“...bringing unconscious aspects of experience to conscious awareness, thereby making them available for conscious choice.” (Sengers et al., 2005, p. 50)

This definition positions reflection as a mechanism for surfacing what usually remains hidden. Sengers adds that this process operates beyond rational thought alone:

“Additionally, reflection is not a purely cognitive activity, but is

folded into all our ways of seeing and experiencing the world. Unconsciously held assumptions are not things we rationally know; they are part of our very identity and the ways we experience the world. Similarly, critical reflection does not just provide new facts; it opens opportunities to experience the world and oneself in a fundamentally different way.” (Sengers et al., 2005, p. 50)

Reflection, then, is not merely an intellectual exercise; it shapes how we perceive and engage with the world. Sengers articulates why this capacity matters:

“This critical reflection is crucial to both individual freedom and our quality of life in society as a whole, since without it, we unthinkingly adopt attitudes, practices, values, and identities we might not consciously espouse.” (Sengers et al., 2005, p. 50)

This is precisely what the efficiency-driven paradigm obscures. By hiding the training process, commercial AI tools encourage the unthinking adoption of the model’s defaults, values, and outputs. But when designers engage with the training process

directly, an opportunity opens up: by reflecting on creative practice through the process of training AI models, we gain the chance to understand ourselves as designers more deeply, create an understanding of the workings of the technology at hand, and thereby to ‘*see the world around us in a fundamentally different way*’. This is the core goal of this thesis, and what I will call a *Reflective AI design practice*.

Reflection in design has a rich theoretical lineage in design discourse, most notably Schön’s reflection-in-action (as earlier mentioned). Sengers builds on this tradition, but focuses specifically on reflection in relation to technology. Chapter 2 of this thesis develops how Schön informs the methodology of this research; here, Sengers serves as the anchor because her definition captures both the mechanism and the stakes of reflection when working with AI systems.

By relating to AI models in this way, I effectively use them as mirrors. This metaphor is well established in critical AI discourse, typically highlighting how systems reflect the biases, values,

and power structures embedded in their training data, revealing problematic patterns that need to be addressed (Vallor, 2024; Crawford, 2021). This critical perspective treats what the AI-mirror shows as evidence of flaws that need to be corrected. For this research, however, I redirect that mirroring capacity toward a different purpose. Building on the understanding that AI systems trained on human-curated data inevitably carry the subjective judgments of those who construct them (Dignum, 2021), I create deliberately personalized mirrors to understand my own creative practice. A Reflective AI design practice becomes valuable because the resulting system functions as this kind of mirror; a mirror that reveals not what needs fixing in society, but what I as a practitioner value, notice, and judge in my work. The mirror metaphor, then, offers a way to think about how AI’s capacity to capture the implicit, to bring ‘unconscious aspects of experience to conscious awareness’, might serve individual creative inquiry alongside its critical societal function.

1.3 *Research approach and questions*

The chapters that follow document the development of a Reflective AI design practice through five projects that test and refine the approach across different contexts and creative situations, each

focusing on a different aspect of creative practice. The research follows *Practice-Based Research*, where questions arise from the process of practice and the answers are directed toward

enlightening and enhancing that practice (Candy and Edmonds, 2018). This means the framework presented here was not predetermined and then tested. Rather, it emerged through the act of making: encountering AI in different contexts and reflecting on what those encounters revealed about the designers involved. The principles, tactics, and insights documented here were discovered by doing.

The investigations span solo explorations by myself, a collaborative project with professional designers, a workshop with design students, auto-ethnographic inquiry into my own artistic past, and a three-month art residency in the context of ceramic practice. Through this long-term engagement of approximately six years, across various contexts and forms of practice, three methodological principles crystallized that characterize this Reflective AI design practice and gave shape to the investigations that follow: 1) reorienting purpose from efficiency toward reflection, 2) treating AI as malleable material, and 3) relying on first-person research methods. Chapter 2 develops these principles in depth before the projects unfold in Part II.

This thesis focuses specifically on image-based AI models: object detection and text-to-image generation systems. This focus reflects the visual and material nature of the design practices explored here, where curating, annotating, and training on images creates direct dialogue with the designer’s ways of seeing. While text-based models and other AI systems could potentially support reflective practice, image-based systems offer particular

affordances for engaging with the visual fascinations, aesthetic judgments, and material concerns central to the creative work documented here.

Through this practice, one overarching research question frames the investigation:

What specific design practices enable reflective engagement with AI?

This overarching inquiry is investigated through two interconnected sub-questions. First:

What characterizes reflection when designers engage with AI training processes?

This question examines how technical processes like data curation, annotation, and inference¹ can become sites for reflection, and what forms this reflection takes within creative practice. Second:

What design learnings emerge from this practice that can transfer to other design research and practice contexts?

This question asks what principles and tactics emerge across the five projects that other researchers and practitioners might adapt. The aim is to understand what can be shared across these variations, and what forms the method takes in different contexts.

¹ Inference refers to the moment when a trained model is applied to new inputs to produce outputs.

What
specific
design
practices

Overarching research question

enable
reflective
engagement
with AI?

What
characterizes
reflection
when

28

First sub-question

designers
engage with
AI training
processes?

29

What design
learnings
emerge
from this
practice that

30

Second sub-question

can transfer
to other
design research
and practice
contexts?

31

1.4

Thesis structure

This thesis is organized into three parts that trace the development, demonstration, and implications of Reflective AI as a design practice.

- *Part I: Foundations establishes the theoretical and methodological groundwork.*

Chapter 1: Introduction (this chapter) situates the research within debates about AI in creative practice, establishes the Reflective AI framework through the mirror metaphor, defines core concepts, presents the research questions, and outlines contributions.

Chapter 2: Methodological Orientations articulates the three principles that characterize Reflective AI practice and explains the practice-based research approaches employed across the projects. This chapter shows how the methodological principles from this work emerged through practice and how they guide the investigations documented in subsequent chapters.

- *Part II: Reflective AI Across Creative Practices documents five practice-based investigations arranged chronologically to show the evolution of the research and the iterative development of the practice.*

Chapter 3: An Emerging Dialogue between an AI and a Designer establishes a foundational model

for designer-AI dialogue through object detection, demonstrating how AI's 'surprising' perspective can participate in Schönian 'framing conversations' that help reframe visual content.

Chapter 4: Developing the Objective Portrait Method explores how three designers (myself and two collaborators) use self-trained AI to reflect on our individual 'design worlds' and fascinations. This chapter introduces the Objective Portrait method.

Chapter 5: The Objective Portrait Workshop investigates whether the Objective Portrait method developed with expert practitioners can translate to design education. A three-day workshop with eleven design students shows the development of personal annotation strategies combined with hands-on knowledge about AI models, and the adaptability of Reflective AI practice to novice designers.

Chapter 6: "Why Do I Keep Coming Back to This Chair?" turns the method toward auto-ethnographic inquiry, using 'AI-as-Memory', an evolution of the mirror metaphor, to excavate artistic meaning in past work. This chapter documents a deeply personal investigation into a foam chair created eight years prior, demonstrating how AI trained on personal archives can reveal latent patterns and values.

Chapter 7: From Text-to-Clay brings Reflective AI into dialogue with ceramic practice through a three-month residency at the European Ceramic Work Centre. This chapter explores what happens when the digital processes of AI training encounter the slow, physical materiality of ceramics, creating an iterative loop where physical making generates new training data.

- *Part III: Synthesis and Implications*

Chapter 8: Discussion synthesizes patterns across the five investigations, identifying five types of reflection that emerged through different configurations of the practice. The chapter formalizes two design learnings for practitioners: a disentangled workflow and engagement with productive friction. It concludes by articulating what can be transferred beyond the specific contexts studied and reflecting on the contributions and limitations of this research.

1.5

Dissemination

Throughout this research, I have shared the work across academic, professional, and public contexts. Beyond design research conferences like *Designing Interactive Systems* (DIS) and *Design Research Society* (DRS), I presented at *Dutch Design Week* in Eindhoven and *Salone del Mobile* in Milan, and discussed the research at *FOAM Photography Museum* in Amsterdam and *Nieuwe Instituut* in Rotterdam. Two projects, *Objective Portrait* (Chapter 4) and *From Text-to-Clay* (Chapter 7), also became installations. Koskinen et al. (2011) describe this 'showroom' mode of design research as one grounded in art and design traditions, where exhibitions function as sites for communication and debate. The goal of these installations was to make

the research process visible: showing datasets, annotations, training outputs, and physical artifacts together, so visitors could trace the path from personal data to AI output to reflection. The installations themselves carry the research argument, inviting audiences to consider how AI might serve purposes beyond efficiency. This thesis documents the research underlying these exhibitions rather than the development of the installations themselves, though images of the exhibited work appear throughout the chapters. Exhibiting in art and design contexts brought conversations with audiences outside academia, helping me understand how differently people think and feel about AI beyond the research communities I usually work within.

2. Methodological Orientations

The previous chapter established the problem: commercial AI tools prioritize efficiency in ways that obscure reflection. This chapter develops the methodological foundation for an alternative approach.

The main goal of this thesis is to develop a reflective design practice with AI through practice-based research. Practice-based research is an approach common in design and the arts where the researcher

is also a practitioner, and the practice itself becomes the primary method of inquiry. Within this research method, knowledge emerges from making itself. As Candy and Edmonds (2018, p. 68) write, *“questions arise from the process of practice”* and *“the research goals and opportunities were only discovered out of practice.”* This describes how I worked: over the course of this PhD, I responded to opportunities as they arose, which means the five research projects spanning Chapters 3–7 employ different methodological approaches and occurred in different contexts. Chapter 3 documents a solo exploration. Chapter 4 follows a collaboration with two designers curious about similar questions. Chapter 5 emerged from an opportunity

to teach a workshop at Design Academy Eindhoven. Chapters 6 and 7 return to solo practice through an invited research stay at Sorbonne University and an artist residency at the European Ceramic Work Centre. Table 2.1 provides an overview.

Yet across this methodological diversity, three shared principles emerged that characterize Reflective AI practice as a distinct approach on its own. These principles are not prescriptive methods I thought of in advance, but patterns I recognized by reflecting on all the practices together. The first principle reorients the purpose of AI engagement from efficiency toward reflection. The second treats AI as a malleable material rather than a fixed tool. The third

relies on first-person research methods that position the practitioner’s direct experience as primary knowledge. This chapter establishes these three principles, showing how they provide methodological coherence to the projects documented in Part II.

2.1

Principle 1: Reorienting the purpose from efficiency to reflection

The first principle of a Reflective AI design practice is the reorientation of the goal from efficiency to reflection. In the projects presented here, the primary objective was not necessarily the generation of a specific design or tangible outcome. Instead, the aim was to cultivate reflection on the design practice itself.

Chapter 1 established this reorientation and stated that commercial AI tools, by prioritizing speed and abstracting technical processes, prevent the Schönian ‘conversation with materials’ that makes design practice meaningful. But how does one actually design for/with reflection rather than efficiency with this type of technology? To answer this, I turn

to the philosophy of Slow Technology, which provides a way of thinking when designing interactions with technology with reflective goals.

Hallnäs and Redström (2001) established Slow Technology as “a design agenda for technology aimed at reflection and moments of mental rest rather than efficiency in performance.” Their core insight challenges a fundamental assumption about technology: that faster interaction represents progress. They argue that the distinction between fast and slow technology is not about actual time perception, but about time presence:

“When we use a thing as an efficient tool, time disappears, i.e., we get

Table 2.1
Overview of practice-based research projects across Chapters 3–7, showing the methodological diversity and contexts through which Reflective AI practice emerged.

Chapter	Title	Context	Participants	Research Method	AI Technology	Reflected On	Output Created
3	An Emerging Dialogue between an AI and a Designer	TU Delft	Solo practice	Practice based exploration	Object Detection (YOLO)	Images of paintings and designed objects	Blueprint for a reflective designer AI interaction, paper
4	Developing the Objective Portrait Method	Personal practices	Two designers and myself (Gijs de Boer, Maurik Stomps)	Collaborative project	Object Detection (YOLO)	Design Fascinations	3 Object Portraits, exhibition, pictorial
5	The Objective Portrait Workshop	Design Academy Eindhoven	Design students	Workshop	Object Detection (YOLO)	Research Interests	11 Object Portraits, paper, framework
6	Why do I keep coming back to this chair?	ISIR lab, Sorbonne University, Paris	Solo practice	Autoethnography	Image Generation	Artistic Pasts	Video Essay
7	From Text-to-Clay	European Ceramic Work Centre	Solo Practice	Art residency	Image Generation	Ceramic Practice	Sculptures, exhibition

things done. Accepting an invitation for reflection inherent in the design means on the other hand that time will appear, i.e., we open up for time presence.” (Hallnäs and Redström, 2001, p. 203)

This distinction is crucial for understanding how I developed Reflective AI. Fast technology makes time disappear. The text prompt generates an image instantly, the task is complete, we move on. Slow Technology makes time appear by creating what they call ‘reflective space’: temporal and conceptual room to examine the technology itself, to understand its logic, and to question one’s own relationship to it.

Hallnäs and Redström identify five ways technology can be slow, distinguishing intentional slowness (designed for reflection) from unintentional slowness (poor design or complexity) (Hallnäs and Redström, 2001, p. 203). Technology is slow when it takes time to:

learn how it works



*understand why it works
the way it works*



apply it



see what it is



*find out the consequences
of using it*

Crucially, they argue: *“Using such an object should not be time consuming but time productive; we should get time for new reflective activities. It is not technology for compressing time to do given tasks, but technology supplying time for doing new things”* (Hallnäs and Redström, 2001, p. 203). This reframes slowness not as inefficiency, but as a resource for reflection.

Traditionally, the Slow Technology philosophy brings forward discrete artifacts embodying this interpretation of slowness in their interaction. Odom et al. (2019)’s Olly exemplifies this: a music player that only plays songs from the same day in previous years, transforming music consumption from background accompaniment into a contemplative encounter with personal memory. By restricting choice and introducing temporal constraint, Olly creates a reflective pause where users can consider their relationship to time, memory, and media consumption rather than simply selecting the next song.

Some work has explored slow technology principles in AI-based contexts, though with limited scope. Park et al. (2019) showed that intentionally slowing algorithmic response times in high stakes decision making contexts, such as crime prediction systems, prompted users to reflect on the processes behind algorithmic decisions. When the system delayed its output, users reported questioning how the algorithm arrived at its conclusions, what data it used, and whether they should trust it, reflections that disappeared when the

same algorithm responded instantly. Similarly, Cremaschi et al. (2024) introduced temporal constraints in generative AI through a modified typewriter interface for ‘slow tweets.’ The mechanical typewriter required physical effort for each character, deliberately slowing the act of text generation to critique the accelerated, frictionless content creation of platforms like Twitter. The physical labor of typing became a moment to reconsider what was worth saying.

These approaches assume AI technology functions as a unified tool, or artifact adapted at the interface level. They slow down interaction with a finished system, a completed model that users encounter as a black box. In both cases, the slowness occurs around the AI, not within it: Park’s crime prediction algorithm still processes the same data through the same model, only the display of results is delayed; Cremaschi’s typewriter still feeds text into the same generative system, only the input is slowed. The reflective space is created at the edges of the system.

AI tools, however, are not only unified tools but a composite collection of interconnected components: data collection, annotation, training, and inference. Recognizing AI as this collection rather than a single artifact shows different possibilities. Rather than slowing the interface to a finished system, Slow Technology principles can be applied to engagement with these constituent components themselves. This embeds slowness in the mechanism: the practitioner spends time selecting which images to

include, deliberates over how to label them, waits while the model trains on these choices, and interprets outputs that reflect their decisions. Each component becomes a site where time can appear, where the practitioner encounters the technology’s logic and their own assumptions.

This brings Slow Technology into dialogue with Schön’s reflection-in-action. For Schön, reflection in action is fundamentally different from routine problem solving, which he calls knowing in action. He sees this distinction as central to design practice. In routine situations, practitioners apply established techniques to familiar problems, but design situations are characterized by uncertainty, uniqueness, and value conflicts (Schön, 1983). In these situations, practitioners must engage in what Schön describes as a conversational process: they frame the problem in a particular way, make a move based on that framing, observe what the situation ‘says back,’ and often must reframe their understanding based on unexpected responses. This iterative cycle of seeing-moving-seeing constitutes the reflective conversation through which design understanding develops. Crucially, this conversation requires temporal space; time to observe material responses, interpret what they reveal, and adjust one’s approach.

When approached through Slow Technology principles, AI training processes possess characteristics particularly suited for this Schönian conversation. The act of curating a dataset forces externalization of implicit decisions: *Which images represent my*

interest? *Why these and not others?* The annotation process demands articulation of tacit judgments: *How do I actually define ‘beauty’ or ‘chaos’ when I must label it consistently in images?* Training the model creates temporal extension, hours or days rather than seconds, during which the practitioner anticipates how their training choices will manifest. Finally, the model’s outputs function as Schön’s material backtalk: revealing how the system interpreted curatorial choices, often in unexpected ways that challenge initial assumptions.

This goal shift shows the first principle that shapes a Reflective AI practice. It uses Slow Technology’s philosophy to structure a Schön-inspired reflective conversation with AI. Where Schön observed architects in conversation with sketches, Reflective AI proposes that designers can enter conversation with AI systems through the deliberate, slow process of training them on personally meaningful data, creating conditions where time appears, reflection becomes possible, and understanding develops through making.

2.2

Principle 2: AI as material

Principle 1 established that Reflective AI practice requires slowing down engagement with AI’s constituent processes to create a reflective space. This component-based engagement is possible when we shift how we view the technology itself. Rather than treating AI as a fixed tool or black-boxed service, I approach AI as a material with distinct properties that can be shaped, explored, and engaged with through practice.

The craft tradition provides a foundation for understanding how practitioners develop expertise through long-term material engagement. Richard Sennett (2008) describes how craftspeople learn by working directly with materials that possess their own agency and resistance. He illustrates this with the example of a carpenter

planing a piece of wood. When the plane moves with the grain, the wood yields smoothly, producing thin, even shavings. But when the carpenter works against the grain, the wood tears and splinters, the blade catches, and the surface becomes rough. The grain’s direction and density push back against the craftsperson’s intentions, asking for adjustment — change the angle of working, reverse direction, alter the pressure. Through repeated encounters, the carpenter learns to read subtle variations in color and texture that indicate grain direction before the plane ever touches the surface. The material reveals its properties through how it responds to the maker’s actions.

This dialogue operates through iterative cycles. The practitioner acts on the material, observes the response,

adjusts their approach, and acts again. Through repetition across many encounters, craftspeople develop what Sennett calls ‘tacit knowledge’: embodied understanding of material behavior that guides practice below the level of conscious articulation. The woodworker learns to read grain patterns and adjust tool pressure without explicitly reasoning through each decision. Crucially, this learning emerges directly from the act of doing, from hands-on engagement with the material’s specific resistances.

The framing of machine learning (ML)², a more specific term for AI, as a design material has been explored before. Dove et al. (2017) framed ML as a new and specifically difficult design material for UX designers, while Holmquist (2017) drew parallels with traditional design practice, arguing that just as product designers must understand the physical characteristics of materials like plastic, wood, and metal, designers working with AI need similar hands-on engagement to understand what it can do. Both identify a challenge: developing this understanding is hard for designers to do with software, but it’s particularly hard with machine learning due to its technical complexity and opacity. Unlike wood, which shows its properties through direct physical manipulation, ML systems are opaque, with inner workings hidden even from those who build them. They operate through probabilistic inference from data patterns, making their material qualities less immediately apparent.

What, then, are ML’s material properties? Benjamin et al. (2021)

advanced this question by positioning this uncertainty as a defining material attribute to work with, instead of against. Rather than executing fixed commands like traditional software, ML systems exhibit behaviors shaped by training data, model architecture, and inference context—behaviors that remain partially uncertain and variable. This uncertainty is an intrinsic property of how these systems function.

Outside of these academic framings, artists working with ML have been engaging with these material qualities through practice. Ploin et al. (2022) explored how several artists integrate ML into creative practice, finding that their workflows became iterative, alternating between research, generation, and curation phases. Curation emerged as the primary site for artistic intention, both in assembling training datasets and in selecting from the model’s outputs. Artists like Anna Ridler build custom datasets of self-photographed tulips to train image-generation models, using the dataset itself as an artistic statement about value, scarcity, and the economics of both flowers and AI-generated art (Ridler, 2018). Others like Jake Elwes fine-tune pre-trained models to achieve particular visual effects in drag performance, shaping the model’s understanding of gender presentation and facial features (Elwes, 2019). This curatorial work which entails deciding what goes into the training data, how it is labeled, what gets selected from outputs, shapes what the model learns and how it

² a subset of AI focused on systems that learn from data rather than explicit programming,

generates. The practice creates iterative cycles between human curation and machine generation that resemble how a carpenter works with grain: acting, observing response, adjusting approach.

These examples highlight AI's key material property for Reflective AI practice: trainability. AI systems can be shaped through the data they are trained on and the way that data is structured. This trainability is what makes them materials rather than tools. In this thesis, the primary creative act is not prompting a pre-trained model, but crafting the model itself through data curation and annotation as training decisions, guiding how to deal with the resulting output.

This material property opens up a reflective dialogue described in

Principle 1. The model's probabilistic behavior creates a specific kind of friction. A model trained on 'beauty' might consistently detect features the trainer, or practitioner never consciously recognized as contributing to that judgment, or label images in ways that reveal patterns in the curating choices they might not have articulated to themselves. The gap between intention and interpretation forces reflection. The practitioner can ask: *Did I actually annotate this way without realizing? What does this reveal about my implicit assumptions?* This is how AI's trainability creates conditions for the Schönian seeing-moving-seeing cycle within the Slow Technology framework established in Principle 1.

2.3

Principle 3: First-person research methods

The material engagement described in Principle 2 presents a methodological challenge. Practice-based research, as established in this chapter's introduction, generates knowledge through making. But when the subject of inquiry is reflection itself, understanding emerges through experiencing the reflective process as it happens. Sennett's woodworker develops tacit knowledge through direct physical encounter, learning to read grain, adjust pressure, feel resistance. This knowledge requires

working the wood oneself; it emerges through one's own hands. But when the subject of inquiry is reflection, something that happens inside the practitioner, you also need methods that treat that inner experience as valid data. First-person research methods do this by positioning the researcher's direct experience as a primary source of knowledge.

First-person research methods have gained traction within HCI and design research specifically for studying subjective, experiential,

and reflective aspects of human technology interaction. As Desjardins et al. explain, "*First-person research embraces subjectivity, and the subjectivity of the researcher. It focuses attention on self experience and the senses*" (2021, p. 5). At the same time, such methods maintain methodological rigor in how attention is directed to experience and how it is described, aiming for relevance and rich descriptions rather than repeatability and objectivity (Hornecker et al., 2017). This approach deliberately departs from research traditions that prioritize objectivity and third person observation, recognizing that understanding certain phenomena within technology requires intimate, long term engagement through direct, subjective experience. These methods range from somaesthetic practices that use bodily awareness techniques to inform design, to autoethnographic studies documenting personal experiences with specific types of technology, to contemplative approaches using meditation practices for self observation (Desjardins et al., 2021, p. 2).

Within this methodological landscape of first-person research, I borrow and adapt key principles from autobiographical design, a form of first-person research methods tied specifically to self-use of a new technological system. Neustaedter and Sengers define autobiographical design as "*design research drawing on extensive, genuine usage by those creating or building the system,*" where "*a researcher or designer's own experiences are embodied in the design of a system and its exploration.*"

That is, as the researcher(s) build the system, they use it themselves, learn about the design space, and evaluate and iterate the design based on their own experiences" (2012, p. 514). This choice reflects the nature of the inquiry: to understand how AI training can become reflective practice requires not merely studying users of my systems, but studying first-person engagement with the training process itself. While autobiographical design does not prove generalizability, it can reveal what Neustaedter and Sengers term 'big effects': insights that could significantly impact understanding of system capabilities, where surprises during genuine usage provide crucial understanding that might be missed in controlled studies. Autobiographical design particularly supports early innovation for exploratory systems that occupy new design niches, where there may be no existing system or established culture of use to guide investigation, which is the situation for the design practice of Reflective AI I aim to develop. Writing about and documenting these experiences becomes a form of reflective practice that generates insights about both the technology and the practitioner's own assumptions and creative processes.

Rigor in first-person, practice-based research is measured by the trustworthiness and transparency of the process itself. Neustaedter and Sengers emphasize that autobiographical design requires "*extensive data collection including tracking system usage, documenting changes to the code, and communal blogs*"

shared among project participants” to maintain an audit trail of developing experiences (2012, p. 520). Every project in this thesis was documented accordingly, recording AI-training decisions, process artifacts, conversations, and all resulting datasets and outputs. This documentation was then translated into analysis through writing and conversation. The goal of this process is to generate what

Neustaedter & Sengers describe as a *“detailed, nuanced, and experiential understanding of a design space”* (2012, p. 521), while examining how my own assumptions and biases shaped both the AI systems and their interpretation. This shows that while the methodology is subjective, it is also systematic and accountable, providing the foundation for the practice-based inquiries that follow.

Part II: Reflective AI across creative practices

3.

Ceci n'est pas une chaise

*An emerging dialogue between
an AI and a designer*

This chapter establishes the foundational model for designer-AI dialogue that subsequent chapters build upon. What happens when a designer treats an AI's interpretations not as answers to accept or errors to dismiss, but as alternative perspectives to engage with? Through three explorations with object detection systems, I investigate how an AI's 'surprising'

readings can provoke the reframing central to design practice.

3.1

Introduction

Well before the recent rise of public interest in generative tools, I saw AI-based technologies already making their way into many aspects of creative work. From automatic design suggestions in presentation software to games where the plot and characters are prepared with AI tools, these changes were taking hold. However, at this point, it leaves me with the question of how designers are influenced by this change. For instance, in the context of design practice, AI has been used in analyzing creative work (Maher and Fisher, 2012), in exploring the design space of possible forms for a given product (Burnap et al., 2016), and in generative design (Kazi et al., 2017; Matejka et al., 2018). Then, at a more cognitive level, AI-based systems have been successful in generating non-obvious solutions that match and sometimes surpass a level of human ingenuity (Serra & Miralles, 2021). It is thus logical to expect AI to at least have the capability to support human designers in exploring non-obvious problem and solution spaces, which are key aspects of a ‘good’ design practice (Cross, 2006).

Zooming in the process itself, designing is an act of abductive

reasoning: a new design solution is offered as suiting the problem at hand, and both the design as well as the constraints of the problem are reexamined, often with the result that the understanding of the problem as well as the understanding of the solution changes (Kolko, 2010). Simply put, this change in the perception of the problem and the solution constitutes *framing* in design. Dorst (2015) writes:

“In questioning the established patterns of relationships in a problem situation, design abduction creates both a new way of looking at the problem situation and a new way of acting within it” (p. 53).

Framing is thus an essential part of designing, and novice designers are taught various methods as a way to challenge their assumptions and explore the non-obvious in problem and solution spaces.

Yet, methods create challenges primarily through making design work visible and social (Lloyd, 2019): they depend on dialogue, critique, and external perspectives. Their effectiveness also remains dependent on the engagement, experience and

skill of the designer (Daalhuizen and Cash, 2021). This raises a question about what other sources of challenge might exist for provoking the reframing central to design.

This chapter, as the first in a chronological series of practice-based inquiries, addresses this challenge by positioning an ‘off-the-shelf AI model’³ as an active partner in the fundamental design act of framing. Before the subsequent chapters can go into more personal or material aspects of creative practice—such as reflecting on design fascinations or engaging with ceramics—it is crucial to first establish a foundational interaction model for a designer-AI dialogue. I investigate how an AI’s ‘surprising’ interpretations can provoke framing and reframing,

creating a dialogic model that the rest of this thesis builds upon.

In sections that follow, I report on my practice-based interactions with AI in processes of design. I present three explorations where I engage with object recognition algorithms. The first two explorations are purely interpretive. The AI looks at creative content, and I treat its interpretation as an alternative perspective on that content. The final exploration takes the form of a dialogue between the designer and the AI model in a cycle of creation, modification and interpretation of a designed object, while appreciating the value of the unexpected and surprising interpretations offered by the AI model as an underlying thread.

3.2

Background

To situate this first inquiry, this section reviews two key areas of literature: the theoretical underpinnings of framing in design practice, and the current approaches to using AI as a partner in creative contexts.

● 3.2.1 Framing as a collaborative dialogue

Framing in design is seen as one of the key steps in a design process. The effort of framing a design situation is a mental act that offers possibilities for opening the problem space to find new approaches for solving it (Thurgood & Lulham, 2016). Consequently, creating this ‘problem frame’ can facilitate an alternative perspective on a problem and thus influence ideas generated in the ideation phase of a designer (Silk et al., 2021). In general, a frame can be described as a knowledge structure schema—characterized by “*expectations based on prior experience about*

³ A pre-trained artificial intelligence system that is publicly available and can be used directly without requiring extensive technical expertise, custom development, or training from scratch. These tools typically come with pre-configured models, standardized interfaces, and documentation that allow users to apply them to tasks immediately

objects, events, and settings” (Tannen, 1986, p. 106). In the context of designing, this expectation forms a view on a problematic situation and is characterized and followed by a series of design moves a designer can take, which allows the situation to ‘talk back’, causing novel perspective on the situation and allowing for the construction of new meanings and intentions (Schön, 1984). Adopting a new frame and performing design moves related to that frame can be described as ‘reframing’, which occurs as a result of reflection, throughout a design process (Paton & Dorst, 2011).

The design research literature describes framing as an individual mental act that occurs within a designer. This mental act of framing, however, is examined as a collaborative effort in the context of an interaction with an external agent, for example a fellow design student (Schön, 1984), a colleague, a client (Paton & Dorst, 2011) or even a design brief (Silk et al., 2021) that is framed in a certain way. This view of framing as an interaction raises a critical question for contemporary practice: could an Artificial Intelligence fulfill the role of this external agent?

● 3.2.2 Current approaches to designer-AI interaction

To understand how an AI might do this, it is useful to review the dominant ways it is currently positioned in design research. A prominent approach can be described as ‘designing for AI’, where the technology itself becomes the object of the design process. In this view, AI is treated as a new design material (e.g., Holmquist, 2017; Rozendaal et al., 2018). Designers are tasked with understanding its material qualities, opportunities, and limitations, similar to how I described ‘AI as a material’ in Chapter 2, but here the goal is to assess how it fits and embeds within products such as smart objects, or to deploy it in specific contexts where its benefit is not yet fully established. This approach also includes applying design skills to create the User Interface (UI) and User Experience (UX) for AI systems, aiming to make the technology’s complex workings less opaque to users (Dove et al., 2017; Rozendaal et al., 2018).

Another approach evaluates AI as an ideation partner or tool during a design process (Dove et al., 2017). This area attaches a process-oriented view to AI and aims to identify opportunities for AI tools to support a design process by having it interact with designers during a design-related task. In these setups, specific roles for the AI are identified beforehand. For example, Fu and Zhou (2020) explore the tutor capabilities of an AI which comments on a design task by giving a variety of suggestions (the AI in question is based on a Wizard of Oz scenario, i.e., the comments are actually

given by a human disguised as AI). Zhang et al. (2021) examine the potential of AI ‘performing’ as a collaborative tool in design teams and assess if this role improves team performance in solving a design problem. They argue that AI boosts the initial performance of low-performing teams as compared to high-performing teams. Such studies recommend interaction scripts that an AI should have to be beneficial to a designer’s process, based on fixed design tasks that are given to the participants.

While the results of such studies are interesting, they are also limited. There is no sense of the exploratory and fluid process of design thinking often exhibited in actual practice. One of the opportunities of AI that has not yet been widely explored in literature is how these types of AI-powered design processes can be set up (Malsattar et al., 2019; Chen et al., 2019), particularly processes that embrace rather than script the designer-AI interaction.

● 3.2.3 AI as a reflective material

These limitations show a critical gap. While the concept of ‘designing for AI’ treats AI as a material for a final product, and the ‘ideation partner’ model often confines the AI to scripted roles, neither fully leverages the AI’s potential to engage in the fluid, reflective dialogue that is central to framing.

This chapter synthesizes these two perspectives by proposing a novel approach: treating the AI as a *reflective material* within a partnership. It adopts the view of AI as a material, but reframes its purpose from product-oriented to process-oriented. In this view, the AI’s material properties, its unique logic and dataset-driven worldview, are not used to build a final product, but to serve as a direct input into the designer’s own reflective practice. This transforms the AI into a process material whose value is measured by its ability to provoke a designer’s thinking.

An AI’s ‘way of seeing’ can provide material for framing: its classifications, shaped by training data, offer alternative readings that a designer can take up as frames for reflection. As Malsattar et al. (2019) suggest, a designer’s understanding can change due to ‘seeing’ the world from the perspective of an AI. AI’s training data contains preset ideas about the world, and object detection algorithms often present results in an actual frame (i.e., a bounding box). The challenge, therefore, is to move beyond prescribed roles for AI and instead investigate how a designer’s interaction with this unique material can provoke the very framing and reframing central to design practice, an inquiry this chapter now addresses through a series of explorations in practice.

3.3

A practice-based inquiry

● *Starting Points*

The designer (hereafter referred to as ‘I’) initiates an open-ended, non-briefed interaction with the AI model. The results obtained from the AI are used as an additional information source for the designer to make decisions on which steps to take next. This explorative, non-guided approach is chosen because predefined goals or scripts could constrain both the designer’s responses and the range of possible interactions, making it difficult to observe how reflection emerges organically. The question driving these explorations was: How would the designer and the AI model respond to each other when there are little or no predefined expectations or rules for getting to a specific design outcome?

Two goals guided these explorations. First, to investigate how I might seek and use information from an AI model while developing a new design process collaboratively. Second, to identify when instances of framing or reframing occurred due to interactions between myself and that same AI model. The explorations described below are closely documented, with my reflections integrated throughout.

● *Object Detection Systems*

Given the visual nature of these explorations, the models I used fall within computer vision, a research area

investigating how to derive meaningful information from visual input in order to take predefined actions or make recommendations. Biometric software, for example, relies on computer vision to detect and recognize faces, classifying photographs according to gender or age, or matching a face to an existing dataset. One established application of computer vision is object detection, which deals with detecting instances of semantic objects an image might contain, such as humans, buildings, or everyday objects (Papageorgiou & Poggio, 2000).

Object detection draws on two core approaches, one based on neural networks and one on non-neural methods (Zou et al., 2023). Today, convolutional neural networks (CNNs) offer the preferred framework, and both models I worked with are CNNs. The first is TensorFlow 2.04, a software library (Abadi et al., 2016) that includes a CNN pre-trained for image classification, accompanied by a step-by-step tutorial called TensorFlow for Poets. This pre-training was performed on the ImageNet dataset (Krizhevsky et al., 2012), which contains 14,197,122 images and allows the model to differentiate 1,000 classes. The second is YOLOv4, pre-trained on COCO.

A model’s performance depends on how it is trained. Training a computer

vision model requires a well-annotated dataset that covers different image classes and represents each class well. Pre-trained models, however, can sometimes be retrained on much smaller datasets through transfer learning, a machine learning method that applies knowledge gained from one problem to a different but related

one. YOLOv4’s pre-training on COCO, for example, makes it possible to redirect the model toward objects outside its original classes, such as the objects in artworks. This flexibility became crucial for the explorations that follow, where I needed the AI to interpret creative content rather than everyday objects.

3.4

Explorations in designer-AI dialogue

This section details the three explorations that form the core of this practice-based inquiry. They are presented chronologically to demonstrate a deliberate progression in my interaction with the AI. The journey begins with passive interpretation of the AI’s outputs, moves to more focused interpretive analysis of designed objects, and culminates in an active, iterative dialogue that operationalizes Schön’s ‘seeing-moving-seeing’ cycle (Schön, 1992). This progression was designed to systematically build a model of interaction from the ground up.

● *Exploration 1:*

Establishing Surprise

The first exploration begins by examining the AI’s response to an open-ended creative context, in contrast to the predefined goals for which object detection systems are typically deployed. In creative processes, knowing a clear outcome in advance is often lacking due to their fluid and exploratory nature, but it is the

Schönian idea of ‘surprise’ that often drives what are termed ‘moving experiments’ (Schön 1983). This exploration sought to understand if and how an Object Detection system could generate this kind of productive surprise when confronted with creative content.

I applied YOLOv4⁶ (Bochkovskiy et al. 2020) — trained on the MS COCO dataset (Lin et al. 2014), a large-scale object detection dataset with 91 categories and 2.5 million labelled instances — to Salvador Dalí’s surrealist painting *Dream Caused by the Flight of a Bee Around a Pomegranate a Second Before Awakening* (1944). This painting, with its dreamlike imagery of a floating pomegranate, bees, tigers, and a resting figure, provided an ideal

⁴ Tensorflow is an end-to-end open-source platform for Machine Learning, containing a comprehensive ecosystem of tools to build and deploy machine learning models (<https://www.tensorflow.org/>)

⁵ <https://kiosk-dot-codelabs-site.appspot.com/codelabs/tensorflow-for-poets/#0>

⁶ <https://pjreddie.com/darknet/yolo/>

test case for how an object detection system trained on everyday objects would interpret highly symbolic and surreal visual content.

When the system ‘looked’ at this artwork, the results were ‘faulty’ interpretations compared to what I could clearly see (see *figure 3.1*). In other words, the AI ‘misclassified’ what it saw in human terms, though these unexpected labels could be taken up by the designer as frames for interaction. For example, the painted bunch of grapes emerging from the pomegranate was misclassified as ‘broccoli’, perhaps due to the dense, clustered depiction resembling broccoli florets, or due to its position in the composition. The tiger’s paw was misclassified as ‘skateboard’, a label even less related to the original object’s visual features. In both instances, although mislabeling occurred in human terms, the surprise of this mislabeling set a concrete frame for engagement. A question arose: what features is the AI seeing to make its classification?

I also applied YOLOv4 to Hieronymus Bosch’s *Garden of Earthly Delights* (c. 1490–1510), where the system’s interpretations became even more surreal (see *figure 3.2*). In this densely populated triptych of fantastical figures, the AI detected ‘dining table’, ‘cake’, ‘bed’, ‘dog’, and several configurations of ‘person’ labels, attempting to impose domestic, mundane categories onto Bosch’s otherworldly scenes of paradise and hell. The AI’s insistence on finding familiar objects in the unfamiliar again demonstrated how its training on everyday imagery created

productive friction when confronted with the surreal.

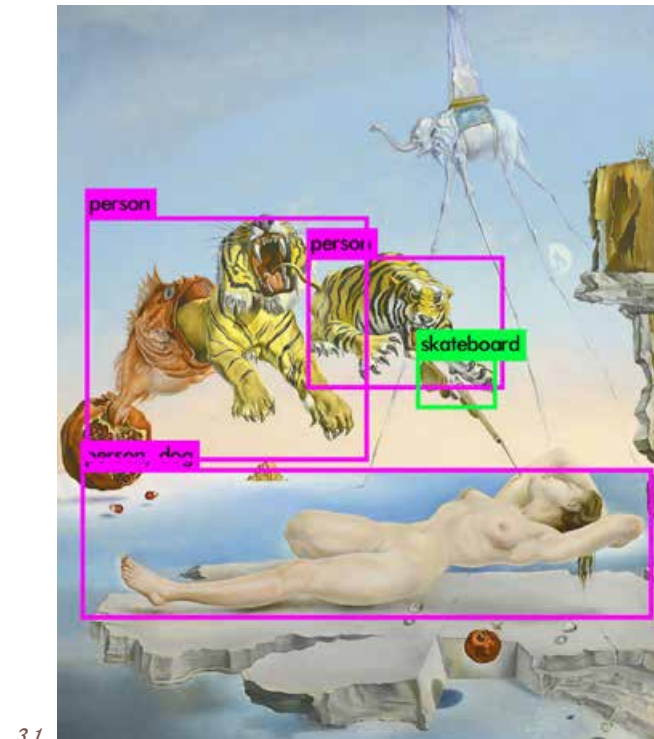
Even though the cause of this mislabeling can be explained through technological characteristics (what is detected minimally corresponds with features found in the dataset categories), it opens up a possible exploration space for design research. The mislabeling can be appreciated as providing an alternative view that the designer can adopt as a frame on the original content – a ‘surprising’ perspective due to its unexpectedness.

● *Exploration 2: Deepening interpretation with designed objects*
Building on the initial phenomenon of AI driven surprise, the second exploration narrows the focus to test the reframing potential of mislabeling intentionally designed objects. The central question here was: *What will the AI make of the functional and aesthetic intentions embedded in designed objects?* Where the first exploration tested the AI on surreal art, this exploration examined how the system would interpret objects crafted with specific design intentions, functional furniture pieces that, while unconventional in their aesthetic approach, still operate within recognizable typologies.

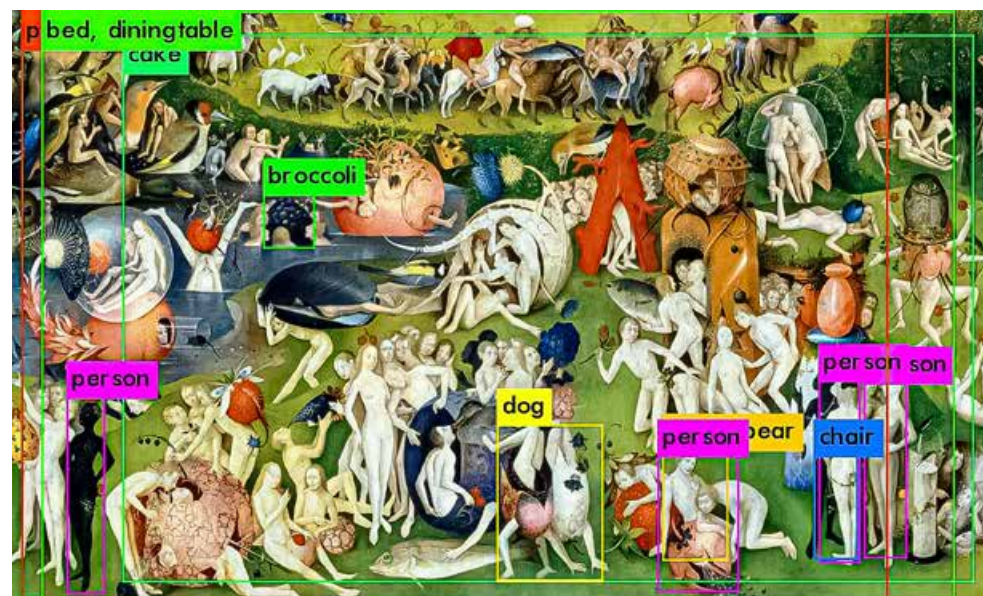
I applied YOLOv4 to images of two designed objects: Nacho

→ *Figure 3.1*
Objects detected via YOLOv4 in Salvador Dali, “Dream Caused by the Flight of a Bee around a Pomegranate the Second before Waking” (1944). The paw of the tiger, perhaps along with the stock of the rifle, is framed as a skateboard.

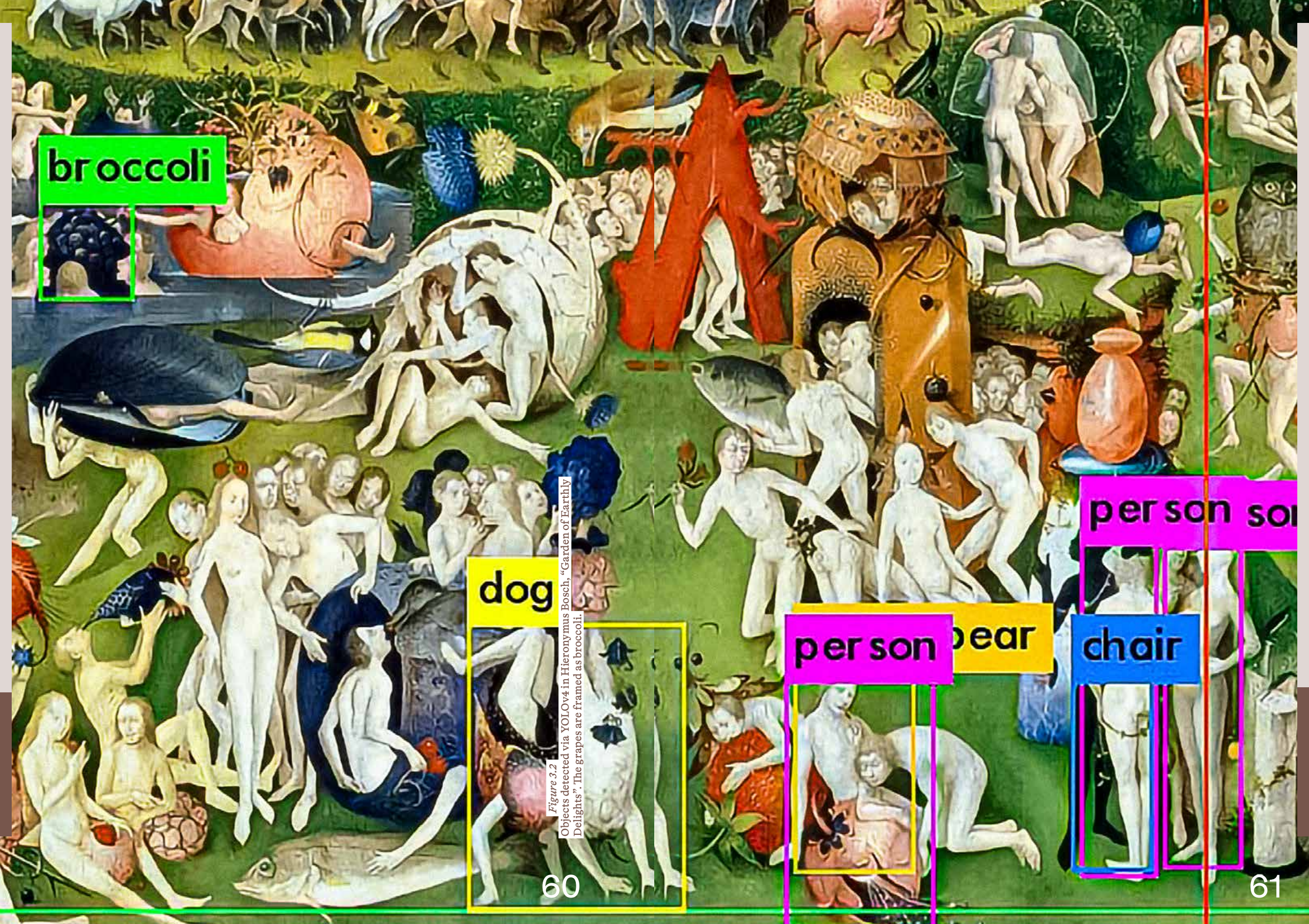
→ *Figure 3.2*
Objects detected via YOLOv4 in Hieronymus Bosch, “Garden of Earthly Delights”. The grapes are framed as broccoli.



3.1



3.2



broccoli



dog



60

Figure 3.2
Objects detected via YOLOv4 in Hieronymus Bosch, "Garden of Earthly Delights". The grapes are framed as broccoli.

person bear



person son



chair



61



3.3



3.4

Carbonell's *Light Mesh* series, a sculptural lighting object with an organic, mesh like structure emerging from a rounded base and supported by thin metal legs, and Maarten Baas' *Clay Furniture*, a cabinet with an intentionally rough, hand sculpted clay aesthetic. Both objects occupy an interesting space between functional furniture and sculptural art, making them ideal candidates for testing how an object detection system trained on conventional objects would parse their unconventional forms.

When I analyzed the AI's interpretations, two distinct types of behavior emerged. In the case of Carbonell's *Light Mesh* (see figure 3.3), the AI engaged in a surprisingly surreal reading of the object. Rather than recognizing it as a functional lighting fixture or piece of furniture, the system detected birdlike and animal-like forms within the organic mesh structure. The rounded elements from which the mesh tree-like form emerges were interpreted as birds, sheep, or other animal-like entities. The AI seemed to respond to the organic, almost nest-like quality of the mesh, reading biological life where I saw design craft. It was as if the AI looked at this carefully designed object and saw a flock of birds or a gathering of creatures, fragmenting the unified designed object into a collection of living things.

← Figure 3.3

AI interpretation of designed objects: *Light Mesh* Series by Nacho Carbonell. The system detects 'bird' and 'sheep' in the sculptural rock forms.

← Figure 3.4

AI interpretation of designed objects: *Clay Furniture* by Maarten Baas. The system detects 'chair,' 'elephant,' 'zebra,' and 'sofa'.

In Baas's *Clay Furniture* (see figure 3.4), a different pattern appeared: multiple detection boxes overlapped across almost the entire object, each associated with different predictions. The AI seemed uncertain, proposing various competing interpretations simultaneously, perhaps seeing elements of 'vase,' 'sofa,' or other not that everyday objects like 'zebra' or 'elephant' in the rough, hand formed clay surfaces.

These interpretations, though technically incorrect, provided a surprisingly rich perspective for reflection. The AI's reading of Carbonell's piece as populated by birds or animals resonated with something genuine about the object's aesthetic language. Carbonell's work does draw on organic, natural forms. The mesh evokes nests, cocoons, or organic growth patterns. The AI, in its 'misreading,' perhaps detected qualities that the designer intentionally embedded: a sense of the living, the grown rather than the made, the natural rather than the industrial. This made me question: Was the AI really wrong? Or had it detected an undercurrent of biological metaphor in the design that functions beneath the surface of its furniture typology? The fragmentation into multiple creatures, rather than seeing a single object, suggested that the piece might be experienced not as a static, unified form but as something more dynamic, a gathering, an emergence, a collection of organic elements in dialogue.

Similarly, the AI's multiple overlapping interpretations of Baas's cabinet highlighted the tension between its functional form and its sculptural,

almost archaeological aesthetic. These surprising perspectives could activate reflective thought about the original designs, presenting an opportunity to reevaluate them and ideate further on specific elements.

I extended this analysis to two additional iconic design objects: Martine Bedin's *Super Lamp*, a postmodern car shaped lighting object from the Memphis Group, and Sabine Marcelis's *Totem Lights*, a series of cylindrical resin columns with internal lighting (see figure 3.5). The AI interpreted Bedin's playful, toy like *Super Lamp* primarily as 'bottle', with multiple overlapping bottle detections across the object's rounded form, ignoring its wheels and car like silhouette. Marcelis's elegant, minimalist light columns were read as 'refrigerator', 'cell phone', and repeated 'refrigerator' labels, the AI seemingly responding to the vertical, appliance like proportions rather than the translucent, luminous qualities central to the design. These additional examples reinforced the pattern: the AI consistently struggled with objects that blur categorical boundaries, offering interpretations that, while technically incorrect, highlighted interesting tensions in how we categorize designed objects.

A tentative understanding of the AI began to emerge through these explorations, but the interaction so far remained limited, fundamentally one directional. I was observing the AI's interpretations, but not yet engaging in dialogue with them. How could I generate a richer, more reciprocal conversation?

● *Exploration 3: Activating an iterative dialogue of 'seeing-moving-seeing'*

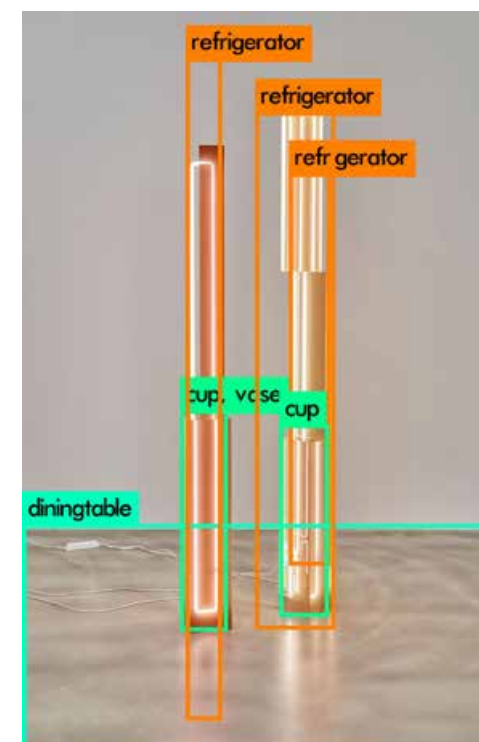
The previous explorations demonstrated that the labels given by the AI can be taken up by the designer as frames for dialogue, but the interaction remained one-way. To generate a richer conversation, a true dialogue needed to be established. This raised a new set of questions: *How to talk back to the AI? How to reframe its perspective? How to show the AI a different understanding?* To address this, I designed the final exploration to create an iterative, dialogic loop where the AI and I would respond to each other through cycles of interpretation and physical modification.

For this conversation, I retrained Tensorflow 2.0 with the CalTech101 dataset, which contains 101 object categories with 40 to 400 images each of mundane objects and natural artifacts (flowers, trees, animals). This retraining was crucial: it gave the AI a different vocabulary to 'see' with, one that was still grounded in everyday objects but offered a wider range of possible interpretations. The retrained AI would be given an image for object detection, predicting what objects it saw and assigning probabilities to each prediction.

To start the cycle of interaction, I selected a designed artifact to work with: a simple wooden chair. I

→ *Figure 3.5*

AI interpretation of designed objects: *Totem Lights* by Sabine Marcelis, where the system detects 'refrigerator', 'cup,' 'vase,' and 'dining table' in the monolithic resin light columns (top); *Super Lamp* by Martine Bedin, where the system detects 'car' for the wheeled base and 'bottle' for each colorful light bulb (bottom).



3.5



3.5

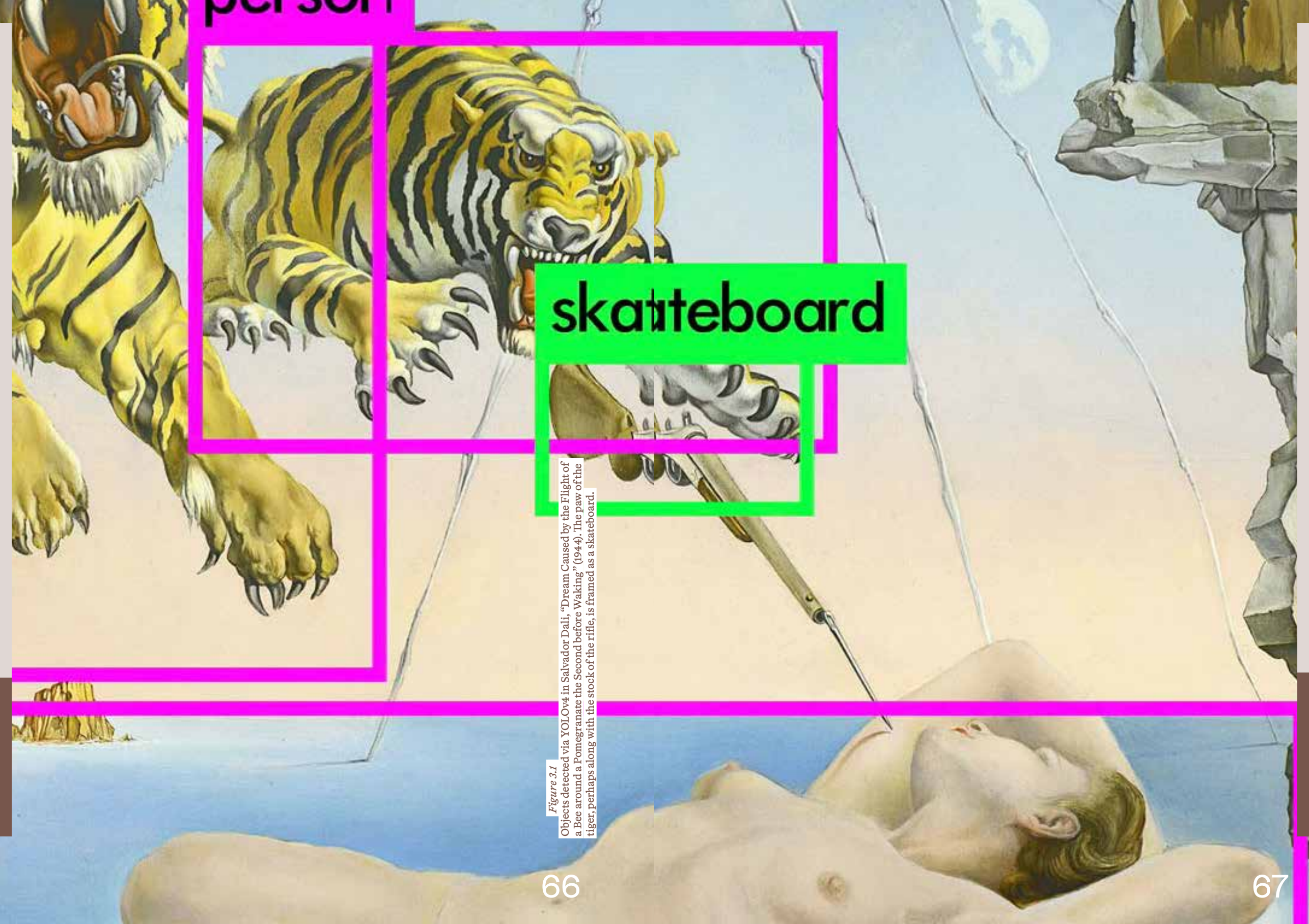


Figure 3.1
Objects detected via YOLOv4 in Salvador Dalí, "Dream Caused by the Flight of a Bee around a Pomegranate the Second before Waking" (1944). The paw of the tiger, perhaps along with the stock of the rifle, is framed as a skateboard.

skatteboard

*Designer interprets
and adjusts*

AI interprets

AI interprets

*Designer interprets
and adjusts*



Image 1



Image 2



Image 3

Table 3.1
Predictions for successive iterations of the designed object. Each modification responds to the AI's previous interpretation, creating a dialogue of seeing and making.

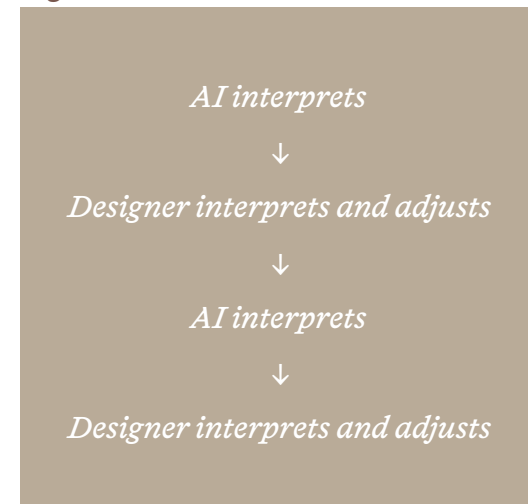
Prediction	Probability	Prediction	Probability	Prediction	Probability
Knife	0.766	Chair	0.609	Chandelier	0.973
Chair	0,087	Chandelier	0.119	Chair	0.001
Kitchen	0.080	Sunflower	0.072	Kitchen	0.001
Tulips	0.055	Accordion	0.030	Tulips	0.001
Roses	0.017	Menorah	0.023	Roses	0.0003

photographed the chair against a neutral background to prevent visual noise from disrupting the object recognition process, then input the image to the AI and noted the predictions and corresponding probabilities. The AI's response was surprising: while it did recognize elements of 'chair', it also saw 'knife' with a notable probability (see *Image 1 in table 3.1*). This unexpected label became the prompt for my first material intervention.

The ranking order of objects backed up with statistical predictions functioned as the AI's alternative 'way of seeing', its 'visual frame' on the object. But it was unclear how the AI had 'seen' these artifacts in the original image. To get a better understanding of this 'gaze', I treated the predictions as criteria for a design brief: Adjust the object (chair) by following loosely the given percentages. If the AI saw 30% knife and 70% chair, I would modify the physical object to incorporate knife like qualities while maintaining its chair like form.

This design brief initiated a cycle where I used my own interpretation of what the object categories could mean physically. When the AI saw 'knife' in the chair, I asked myself: How could I add knife like elements to this chair? What materials might embody this quality? How could the chair still function as a chair if it has 'knives' added to it? I made physical modifications, adding sharp, blade like metal elements to the chair's structure. Once I finished adjusting, I photographed the new artifact and gave the image file to the algorithm for classification (see *Image 2 in table 3.1*).

A dialogic practice emerged. The AI interpreted, I responded through making, the AI interpreted again. The nature of this interaction became circular, a feedback loop in which the AI and I iteratively interpreted and reinterpreted the evolving object. The three iterations that emerged from this dialogue are shown in *table 3.1*. With each cycle, the object became stranger, less recognizable as a conventional chair, more hybrid. This cycle continued for three iterations. The AI continued to offer surprising labels: 'chandelier', 'menorah', 'tulips', each detection prompting new questions about how to materially manifest these qualities. But some suggestions proved untranslatable: how does one physically realize a percentage of 'menorah' in a chair? The dialogue reached a natural boundary when the gap between the AI's interpretations and feasible physical responses became too large. The object had also transformed to a point where further modification would lose all connection to its original form. The creative dialogue reached its conclusion.



Having the AI give surprising interpretations of the designed object resulted in very specific, illogical actions, yet actions based on a level of pattern recognition. There was a sense of control and being controlled, of listening and acting. Surprise kept me from routine behavior, with the AI opening up problem and solution space, but also a space of reflection.

'Seeing' knives in a chair opened up a range of possibilities that I had not considered before. The AI moved me to a new space of possibilities through this process of framing and reframing, making the familiar strange and forcing me to negotiate between the object's original identity and its emerging, AI driven transformations.

3.5 *Seeing-moving-seeing: A model for designer-AI dialogue*

These explorations have shown how a practice-based dialogue between a human designer and an AI can be conceptualized through the idea of framing and associated concepts, in particular 'surprise'. Dorst and Cross (2001) emphasize that surprise keeps a designer from routine behaviour, independent of whether this occurs in problem or solution exploration. The practice developed by the designer in the present study aligns closely to Schön's description of designing as reflection-in-action whose basic structure—seeing-moving-seeing—is an interaction of both designing and discovering: "Working in some visual medium [...] the designer sees what is 'there' in some representation of site, draws in relation to it, and sees what has been drawn, thereby informing further designing" (1992, p. 135).

When I translate this theory to the interaction script of the designer and the AI, I find that the AI 'saw' the object which then provided a frame for the designer to explore in a 'move-experiment'. An iterative practice thus emerged whereby the AI then 'sees again' the reinterpreted design. It is intriguing that, even 30 years ago, Schön was considering how AI could work with designers, considering the forms of knowing-in-action that a 'knowledge-based system' would need to have in order to meaningfully support design activity: "*I conclude that the practitioners of Artificial Intelligence in design would do better to aim at producing design assistants rather than knowledge systems phenomenologically equivalent to those of designers*" (Schön, 1988, p. 131). Ironically, this project suggests that the intuitive reflection-in-action process that Schön describes when

humans design together provides a better model for how present day ‘off-the-shelf’ AI can interact with human designers.

● *The productive role of surprise*

A key insight from this inquiry is the (re)framing of the AI’s technical ‘failures’ as productive creative acts. Although the AI has agency in the design interaction, the roles of each party are not equal. Though the AI takes on a controlling role, it is largely passive and restricted to ‘seeing’. By ceding power to the AI the designer takes on a more active role, conducting the ‘moving’, before reframing takes place. Nevertheless, the AI does generate surprise from its necessarily limited training data. As noted, there is a balance to be struck between predictability and randomness in generating classifications, but in the present sequence of explorations surprise was triggered through the ‘mislabeling’ of image contents. Mislabeling, of course, depends on a human interpretation of the image, but as showed, whether the label is correct on human terms does not alter the fact that it can be used as a productive creative frame.

● *Limitations and future inquiries*

This inquiry has several limitations that, in turn, open up avenues for the subsequent inquiries in this thesis. First of all, the AI-models used are not unbiased. The retraining dataset has more images of knives than it does of chairs (mostly with the appearance of barstools). Should the mislabeling be trusted? Is there any reality to the

AI’s classifications? Furthermore, due to the uneven distribution of the images in each of the categories (some categories have more image examples) the AI tends to detect well-represented categories with higher-level of accuracy than those categories with a poorer representation. Additionally, object detection models like YOLO have known technical constraints, including lower precision when identifying small objects or features within complex compositions (Ji et al., 2023). In artworks like Bosch’s densely populated triptych, such limitations likely influenced which elements the system detected and which it missed entirely.

The second limitation is a paradox: a perfectly accurate model may be creatively inert. If a computer vision model that is trained on a wide range of chairs is shown an image of a chair, it detects a ‘chair’ with high accuracy. In other words, a well-working AI system, which is set up for detecting what objects are present in an image, will render predictable results by successfully listing all the objects correctly. On the other hand, a completely random prediction that shows no connection to the designer’s interpretation is unlikely to inspire the designer. What, then, is the right level of unpredictability that will provide ‘appropriate incongruity’ (Ludden et al., 2012) to prompt or inspire the designer?

These limitations lead to speculation on what kind of datasets could be relevant for the kinds of creative practices described in this chapter. The training dataset used for this

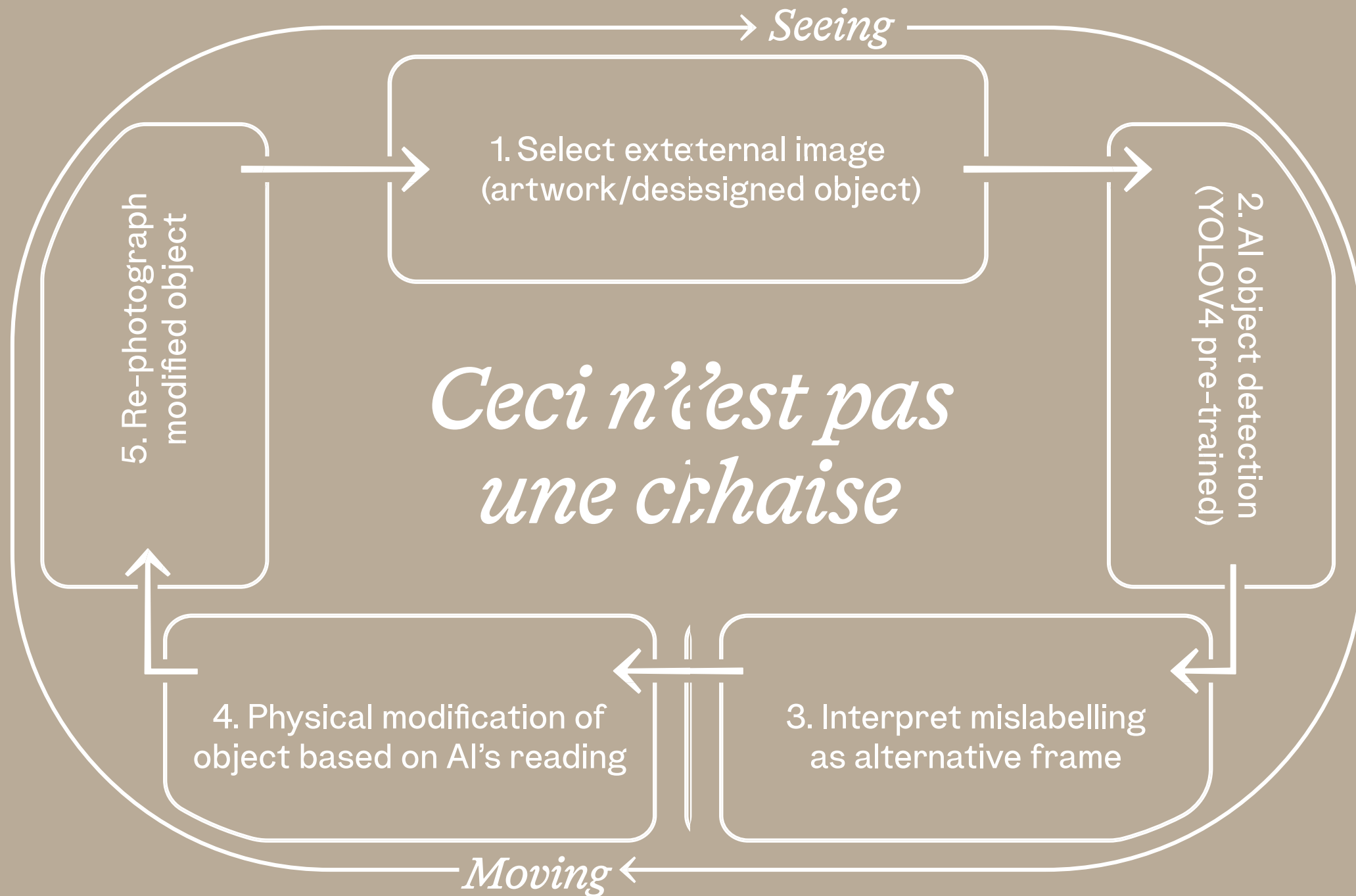
exploration contained images and labels of mundane objects. As this dataset is generally used for the goal of detecting objects in images, its usefulness for augmenting a designer’s creative practice can be questioned. If the goal of the AI is to offer unexpected inspiration, the datasets could be constructed and labeled differently with categories that could be more design-related. For example, instead of labeling a category ‘chair’, a category could be labeled according to a functionality, such as ‘sitting’.

Another approach could be to create datasets that contain different images, for example visual design styles or other design-related objects. Transfer learning offers considerable opportunities and comes closer to thinking of AI as a material that can be adapted and changed to enhance different kinds of creative practice. This concept of adapting the AI’s materiality for a specific creative context will be the central focus of the next chapter, which investigates a reflective dialogue with an AI specifically exploring the annotation and training process.

3.6 *Conclusion*

This chapter has established and demonstrated a foundational model for a designer-AI dialogue, operationalizing Schön’s ‘seeing-moving-seeing’ through the concepts of framing and surprise. It has shown how the perceived technical failures of an AI can be productively adopted by the designer as creative frames,

enabling to reflect on and reinterpret an external artifact. Having laid this groundwork with a focus on external objects, the thesis now turns this reflective lens inward. What happens when the object of reflection is not a chair in the world, but the complex, internal ‘design world’ of the practitioner?



4.

Developing the Objective Portrait Method

The previous chapter established a model for a designer-AI dialogue, demonstrating how the AI's surprising 'gaze' could provoke the reframing of an external object. This chapter turns that reflective inquiry inward. If an AI can help a designer see a chair differently, can it also help a designer see themselves differently? This chapter explores this question

and documents a practice-based inquiry into whether AI could help designers reflect on their own work. Working with two other designers, I explored training subjective AI models on our personal fascinations, the implicit aesthetic preferences and judgments that shape our practices. This exploration led to the development of a five-step process called Objective Portrait, which positions AI training as a site for reflective design practice.

4.1

Reading an image

If you look at the image at *Figure 4.01*, what do you see? Hands? An egg? Creases in a sheet? If you look longer, deeper, do you notice anything else? Do you wonder what the image means? Do you notice irregularities, do you find things odd, or even uncomfortable? Or does it give you a specific feeling? A feeling of joy, perhaps? Or are you rather enticed? What actually pulls you in, and what pushes

you away? What elements attract, feel warm, and what feels rather distanced and cold?

And then, if we zoom in and look at *4.02*, would you be able to point out exactly which element makes you feel this way? How do you identify instances in this image? Does that reveal something about how you read the image? This hand, is it uncanny? Or is it sensual?



4.1



4.2



And this little twirl on the right, does it make you laugh? Do you want to pull it, or make it go away? Does it add to the image, or do you think it could be erased? What could it represent? Is it stubborn, teasing, or merely obsolete?

Or this crease in the sheet, do you want to straighten it, therefore, does the discomfort make you want to touch it? What word would you give

this part of the image, if you would have to describe it? Would it help you if you had more insight in how you look at an image?

↑↑ *Figure 4.01*
An image of the painting "Geopoliticus Child Watching the Birth of the New Man" by Salvador Dalí (1943).

↑ *Figure 4.02*
Details of that same painting.

4.2

Internal design worlds

Our perspectives on images are deeply personal and nuanced. What John Berger termed ‘ways of seeing’ (Berger, 1972) conditions the feelings we get from gazing at things, and what aspects we choose to concentrate on. This is especially true for designers, who construct meanings from visual information by recognising patterns and exploring possibilities beyond what is shown.

These fascinations are a form of tacit knowledge (Wong & Radcliffe, 2000), though not in the sense of embodied procedural skill like throwing clay, or riding a bike. They are tacit in the sense that they guide decisions without being consciously

examined: preferences, attractions, and judgments that shape what a designer notices, selects, and values. These fascinations, in turn, shape the ‘design worlds’ designers construct (Schön & Wiggins, 1992): internal environments built upon personal beliefs, interpretations, and relationships between elements in their projects.

Can we gain deeper insight into these personal design worlds by making them explicit and teaching them to an AI? Can we teach a machine our ‘ways of seeing’? If we can, might it shed light on our own perspectives? By teaching a machine our design fascinations, could we get closer to our own design worlds?

4.3

The tension of the mirror metaphor

As mentioned earlier in this thesis, the mirror is often used as a metaphor for AI models. The mirror metaphor intrigued me because it captures a tension. AI development rightly focuses on eliminating harmful bias: patterns that replicate stereotypes or lead to discriminatory outcomes. But in this focus, personal perspective can get conflated with harmful bias, as if all subjectivity were a problem

to solve. The aesthetic judgments, fascinations, and ways of seeing that shape a designer’s work are not harmful bias but part of who we are. What if, instead of treating this subjectivity as noise, I deliberately amplified it? What if I built an AI mirror specifically customized to my own values, opinions, aesthetic preferences, and desires? Could such a personalized reflection offer

valuable insights precisely because it captured my perspective?

These questions connect to a broader conversation about subjectivity in computational approaches. Developments in powerful methods to classify, categorize, and label data have created a false sense of ‘objectivity’ often associated with computational thinking. Recent research has pushed back against this, emphasizing the inherently interpretive nature of categorizing and labeling. As Murray-Rust et al. (2022) note, metaphors in AI contexts both illuminate and obscure, centering certain ideas while marginalizing others. The language of ‘ground truth’ and ‘bias’ already carries assumptions worth questioning. Tanweer et al. (2021) provide several examples where disregarding context in interpretation leads to incorrect understanding of data, which in turn leads to incorrect generalizations. They make an argument for interpretivist

approaches, to “*probe the multiple and contingent ways that meaning is ascribed to objects, actions, and situations*” (Tanweer et al. 2021, p. 5). Scholars have proposed various ways to bridge this gap: Nelson (2020) offers a framework for computational grounded theory that combines automated text analysis with qualitative interpretation, while Baumer et al. (2020) developed visualization tools that support interpretivist scholarship in machine learning contexts.

The work presented here takes this interpretivist stance seriously. I was interested in deliberately amplifying my own perspective to make it visible to me. The reflection I would see would be shaped by the choices I made in building it: which data to include, how to categorize it, what to pay attention to. If I could train an AI to see the way *I* see, might its interpretation help me understand my own vision more clearly?

4.4

A comparative and collaborative research approach

The inquiry described in this chapter traces how three designers, including myself, each trained a subjective AI model to capture our individual ‘ways of seeing.’ Each of us followed the same cycle: collecting a dataset of images representing a personal fascination, annotating those images with subjective labels, training an object

detection model called YOLOv5, (Jocher et al., 2022) on these annotations, and using the trained model to analyze a composition of our own design work, which we called an ‘Object Portrait’. In other words, we tried to capture our subjective way of looking at the world and then mirror that gaze back at our own

work. The project was named ‘Objective Portrait’, both to capture how we portrayed our practice through objects, and to provoke reflection on the apparent objectivity of an algorithmic interpretation of a portrait.

Gijs de Boer and I developed this process together, therefore the ‘we’ in this writing refers to us. We involved Maurik Stomps at a later stage, to see how it would work across three designers with distinctly different aesthetics and design interests. As described in *Chapter 2*, first-person methods require keeping one’s own practice as the site of inquiry, yet articulating choices to others can create different conditions for reflection on that practice. During this project, there were moments of collaboration and moments where we worked individually. The annotation process (which I will describe in detail in the next sections) where subjective judgment was most at stake as this was the point where we would apply our subjective labels to training data, we each did alone. But we helped each other develop those personal labels that would describe our fascinations, because putting vague aesthetic preferences into words proved easier in conversation. When creating the

object portraits, we assisted each other practically, becoming temporary helpers in each other’s setups. Afterward, we compared how we had each experienced it through reflective conversations of which we made notes, but also through analyzing each other’s documentation of the entire process. All these steps will be described in detail in *section 4.6*.

The goal was to *see* these three practices alongside each other: to compare the resulting object portraits, annotation vocabularies, and processes visually. This visual comparison is central to how the research communicates, which is why this chapter stems from a pictorial publication. We found that pictorial representation was what designers, the envisioned audience for this research, could most readily engage with and understand. The chapter that follows relies also on this visual documentation of our processes. When we translated the research into an installation for Dutch Design Week 2022, this principle extended further: visitors encountered our three object portraits alongside the training datasets that produced them. Images of this installation are included at the end of this chapter.

4.5

Object Detection as a reflective tool

The choice of an object detection model was intentional and informed by the previous chapter, which

established how object detection results, and specifically the ‘mislabelling’ can serve as foundations

for interpretation, speculation, and reflection. Therefore, for the technical implementation, I selected YOLOv5 (Ultralytics, 2023). YOLOv5 offers several advantages for this inquiry. First, it can be relatively easily retrained with modest computational resources compared to larger models. Second, unlike generative AI technologies that create entirely new content from text prompts, YOLOv5 interprets existing images through a distinct visual language – bounding boxes with confidence percentages – that naturally invites reflection by offering interpretations of elements already present in the image.

To ensure that the model responded appropriately to our personally curated and annotated training data, we employed transfer learning (Pan & Yang, 2010). Transfer learning starts with a model already trained on a large general dataset, which has learned to recognize basic visual features like edges, shapes, and textures. This pre-trained model is then fine-tuned on a smaller, specialized dataset. For our purposes, this meant we could train personalized models with only a few hundred images each, adapting the model’s existing visual knowledge to recognize our subjective labels.

4.6

The reflective practice with AI: A five-step process

The following sections detail the five core steps of the Objective Portrait process. Each step is framed by a question that drove our inquiry at that stage. This structure is designed to trace the technical workflow of AI training and the reflective practice that unfolded alongside it.

- Step 1: Collecting a dataset representing our fascinations
Can fascinations turn into AI training data?

To teach an AI our ‘ways of seeing’, we had to translate our subjective fascinations into the language it understands: data. AI systems learn by establishing patterns in large datasets. In the case of object detection models, these are collections of digital images with corresponding labels. To teach the algorithm our own subjective view, we therefore started by collecting a training dataset that visually represented an important fascination in our making practice: a collection of images about which we would have opinions, experiences, or feelings relevant to our own work. Each of us chose a different fascination that was core to our design practice as the basis for our dataset:

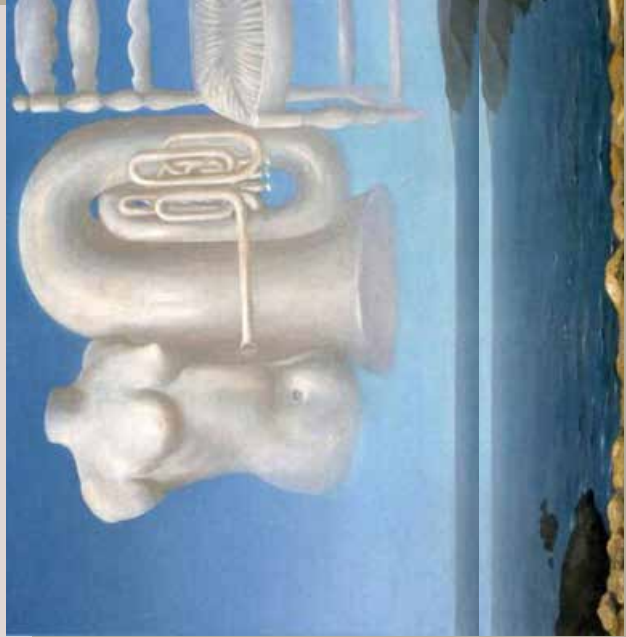
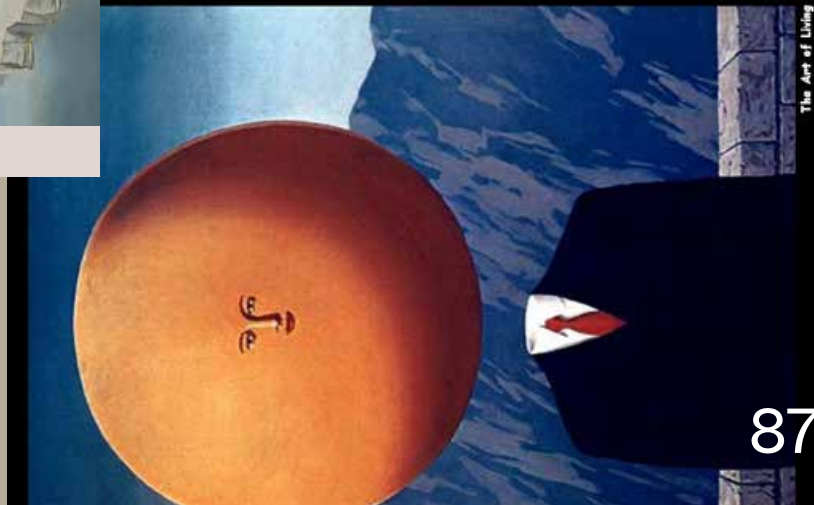
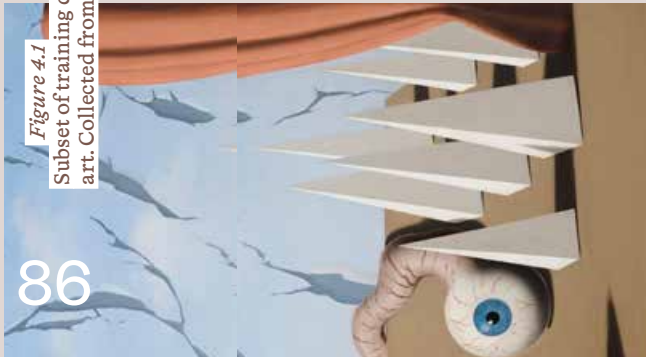


Figure 4.1
Subset of training dataset of Vera, representing her fascination for surrealist art. Collected from various sources online through Google Image Search.



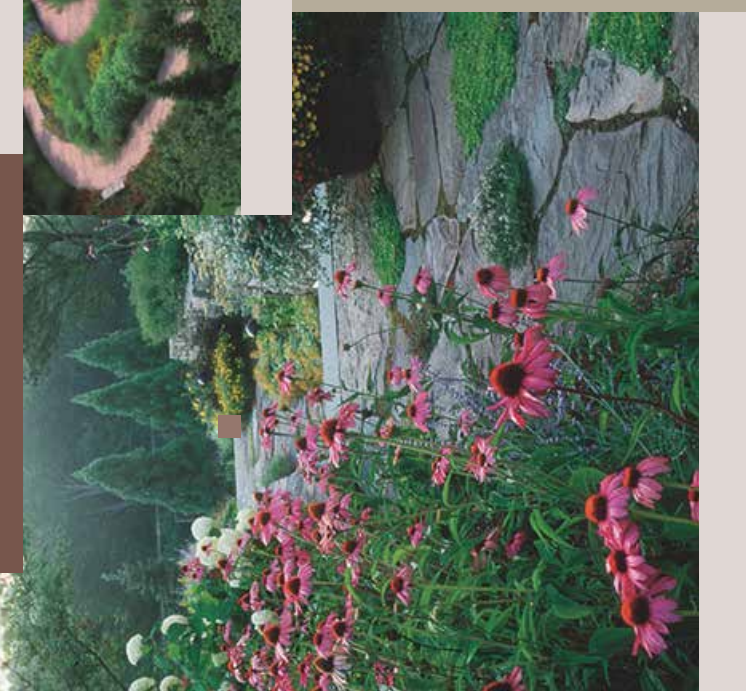
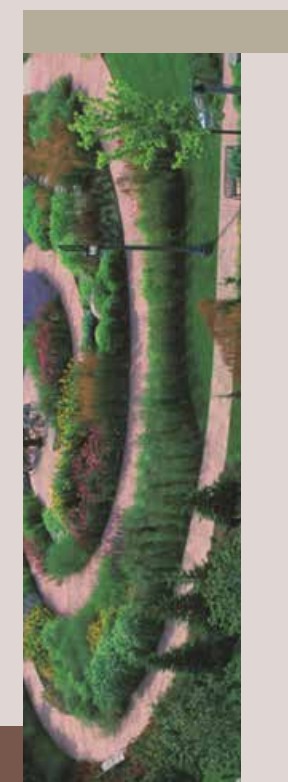
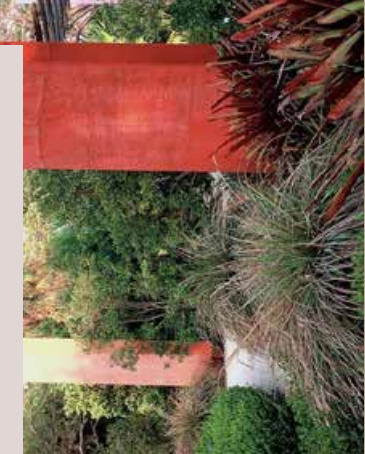


Figure 4.2
Subset of training dataset of Gijis, representing his fascination for gardens.
Collected from Gardens of the world.

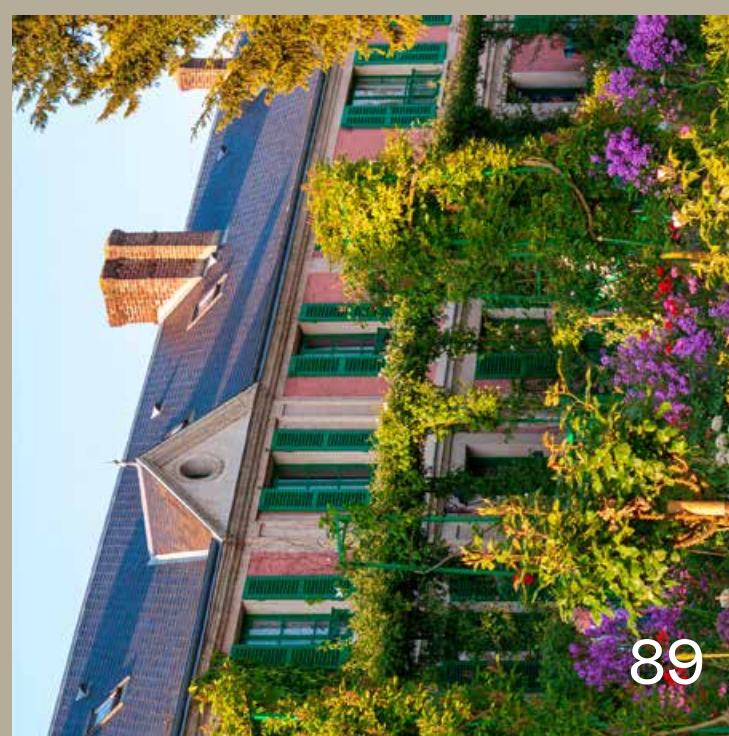
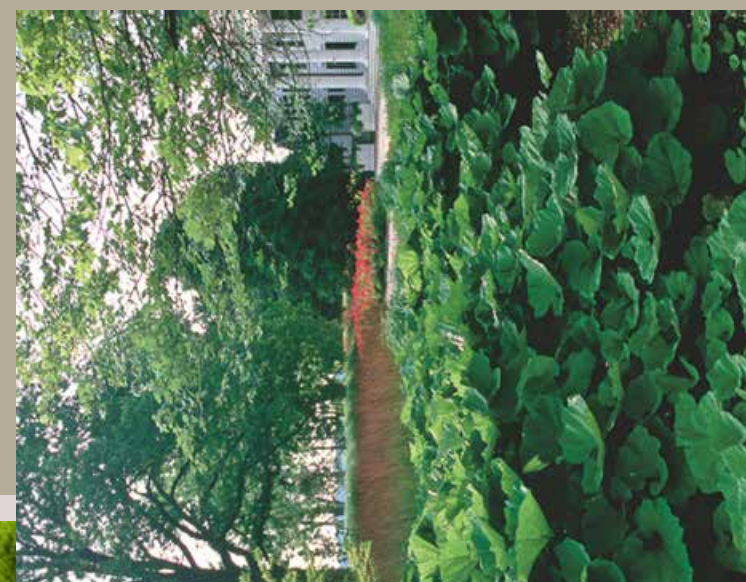
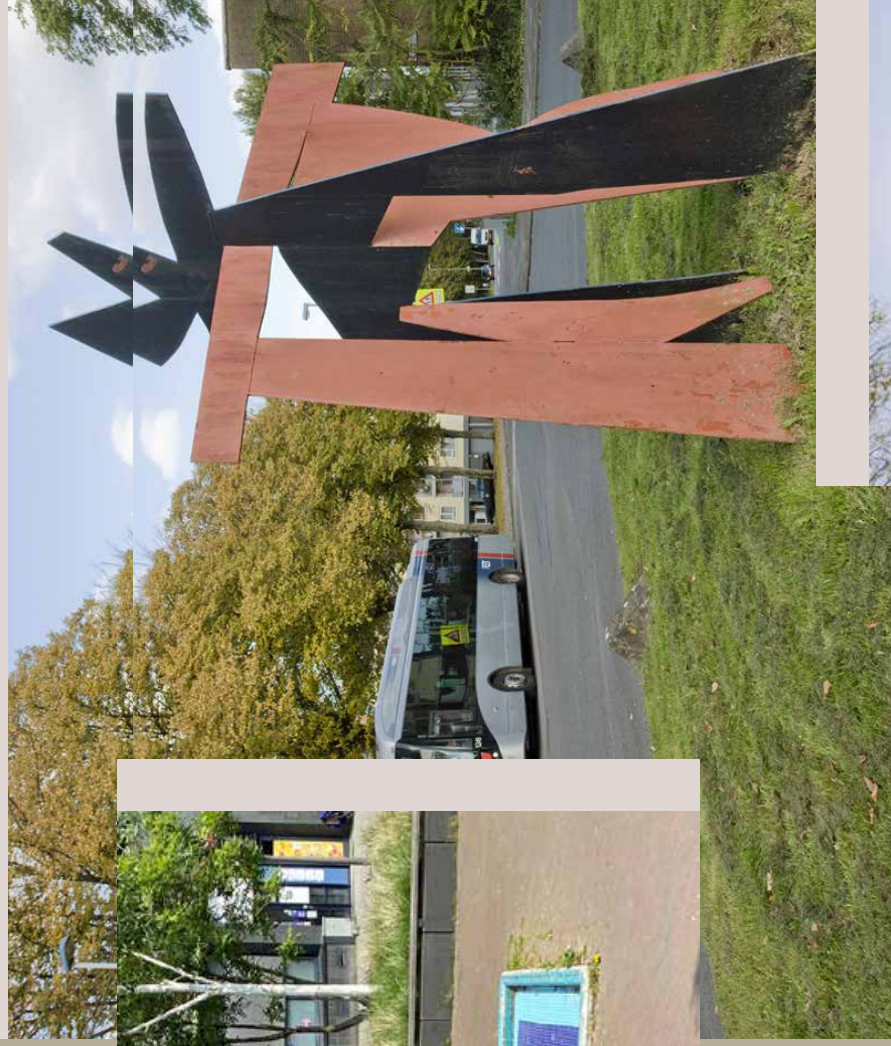




Figure 4.3
Subset of training dataset of Maurik, representing his
fascination for art in public space of Rotterdam.

90



91

- *Gijs de Boer*, a design researcher and writer who studies human-plant relations, chose images of gardens. As he described it: “A garden embodies the tension between care (letting things grow) and the control of giving shape.”
- *Maurik Stomps*, a designer whose work aims to give citizens agency over public space, chose images of art in public space in Rotterdam. He was curious to learn how he actually perceives public space, given his conflicting feelings about top-down design approaches versus appreciating public space as free commons.
- *Vera van der Burg*: In my making practice, I intuitively create and collect objects that tend to have surrealistic, strange shapes. To understand what could be hidden behind this urge or implicit attraction, I used an image dataset of surrealistic paintings by René Magritte and Salvador Dalí.

We each collected approximately 200–400 images. A subset of these images are shown in *figure 4.1*, *figure 4.2* and *figure 4.3*. This quantity was small enough to make annotation manageable for a single person, yet large enough to provide sufficient training data for the model. We sourced these images from online archives, museum databases, and personal collections, selecting them based on their relevance to our fascinations. Several practical considerations shaped our datasets: all images needed to depict visual scenes containing multiple elements for the algorithm to detect, rather than isolated close-ups of single objects. Consistency in visual style within each dataset (e.g., all paintings, or all photographs) helped the model learn patterns more effectively. The hands-on process of searching, selecting, and curating these images was itself a first moment of reflection, as it required us to make explicit what we considered representative of our fascinations.

● Step 2: Choosing subjective labels with the quadrant tool

How to categorize our subjective perspective?

State-of-the-art object detection models are trained on large datasets containing thousands of images, frequently annotated by multiple human workers. Miceli et al. (2020) define data annotation as a ‘sense-making practice’ where annotators assign meaning to data through labels. In industrial contexts, however, this sense-making is typically suppressed. Often outsourced or distributed

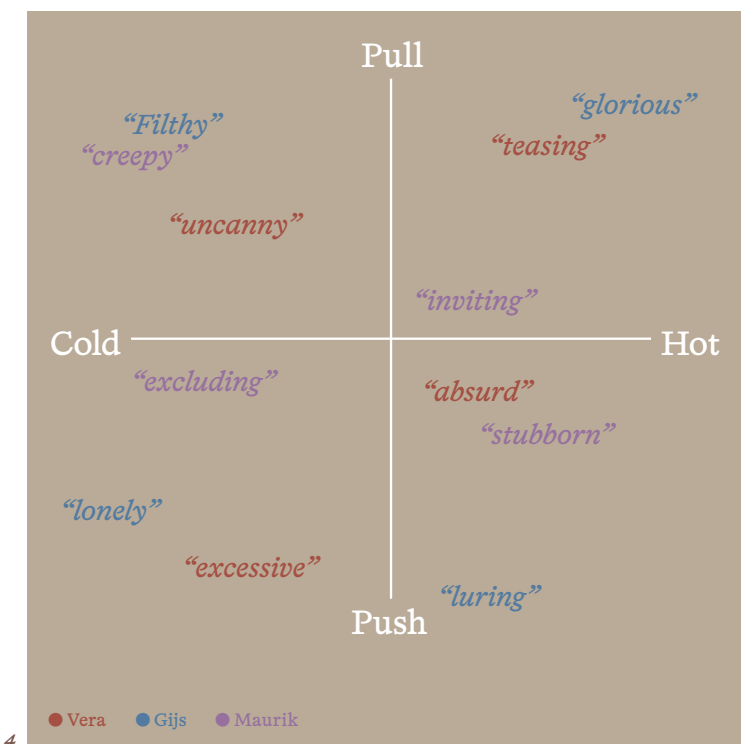
→ *Figure 4.4*

Comparative overview of the three designers’ quadrant tool outcomes, translated in digital. Labels are color-coded by designer: Vera, Gijs, and Maurik. The spatial distribution reveals how each of us navigated the warm/cold and pull/push axes differently based on their specific datasets and fascinations.

→ *Figure 4.5*

An overview of how we approached the selection of subjective labels with the help of the quadrant tool

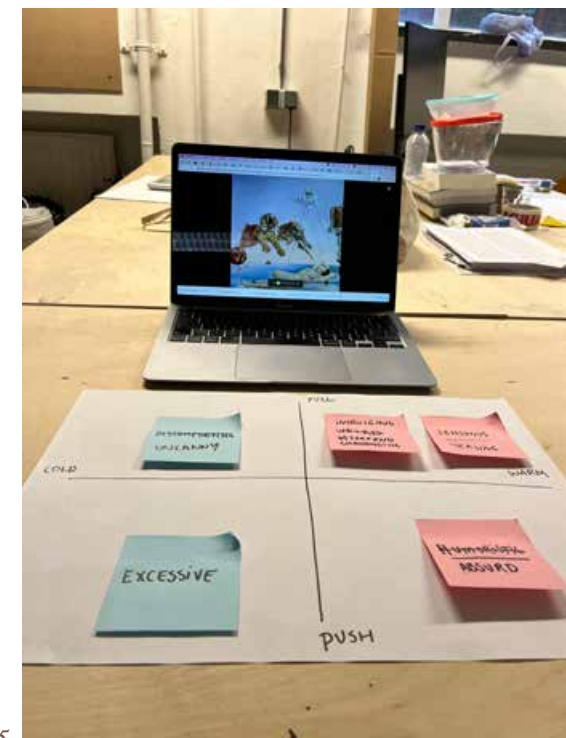
4.4



4.5



4.5



across many workers, annotation is structured to minimize individual interpretation and prioritize speed. Classifications are imposed from above: tasks are predetermined (e.g., ‘Tag all chairs’) and oriented toward providing objectively ‘correct’ labels. As Miceli et al. show, assigning meaning to data is often presented as a technical matter, when it is in fact shaped by power relations between those who define categories and those who apply them.

This step in the process inverted that logic. We were not labeling for someone else’s categories but developing our own. We were looking to capture the subjective gaze of only one person, which meant that only one of us would annotate our own training dataset. To teach our algorithms something about our subjective gaze, we needed labels that captured subjective interpretations that would describe elements in this dataset: feelings, opinions, or judgments we experienced when faced with an image. We limited ourselves to four labels each, a constraint that forced us to distill many possible feelings into a workable set rather than generating endless categories with too few examples to train on.

○ *The Quadrant Tool.*

To choose these four labels, we decided to work with a quadrant tool (see *figure 4.4* for an example). Where conventional annotation instructions tell workers what to see or find in the dataset, the quadrant tool asked us to articulate what we already felt or saw but had not yet named. Inspired by two-dimensional representations of emotional affect like Russell’s circumplex model (Russell, 1980), which typically organize emotions along orthogonal axes of valence (positive/negative) and arousal (high/low), we created a structure with deliberately undefined ‘push/pull’ and ‘warm/cold’ axes. We arrived at these dimensions by examining a subset of around 10 images from our training datasets and asking ourselves: “*Does this element in this image pull me in or push me away?*” and “*Does this element feel ‘warm’ or ‘cold’?*”

These axes were intentionally ambiguous. What ‘warm’ meant was up to each of us: inviting, comforting, familiar, nostalgic. The same for ‘push’ and ‘pull’, which captured something about attraction and distance but left room for interpretation. A label like ‘happy’ might sit in the warm-pull quadrant, while ‘content’ might feel warm but more settled, less actively drawing you in. By positioning candidate labels along both axes, we started to see where our words clustered and where gaps remained, which helped us settle on four labels spread across the quadrants. We kept asking ourselves: “*How would I describe, or label, this feeling?*”

We each drew the quadrant tool by hand. This made it something we could keep returning to, scribbling new words, crossing out old ones, shifting labels closer to one axis or another as our thinking evolved. Looking back at these hand-drawn sheets, we could trace how our interpretations had shifted over time: which words we had abandoned, which ones moved, which ones stuck (see *figure 4.5*). What started as a practical step in preparing labels for annotation became a reflective exercise in itself.

● Step 3: Annotating the dataset with subjective labels

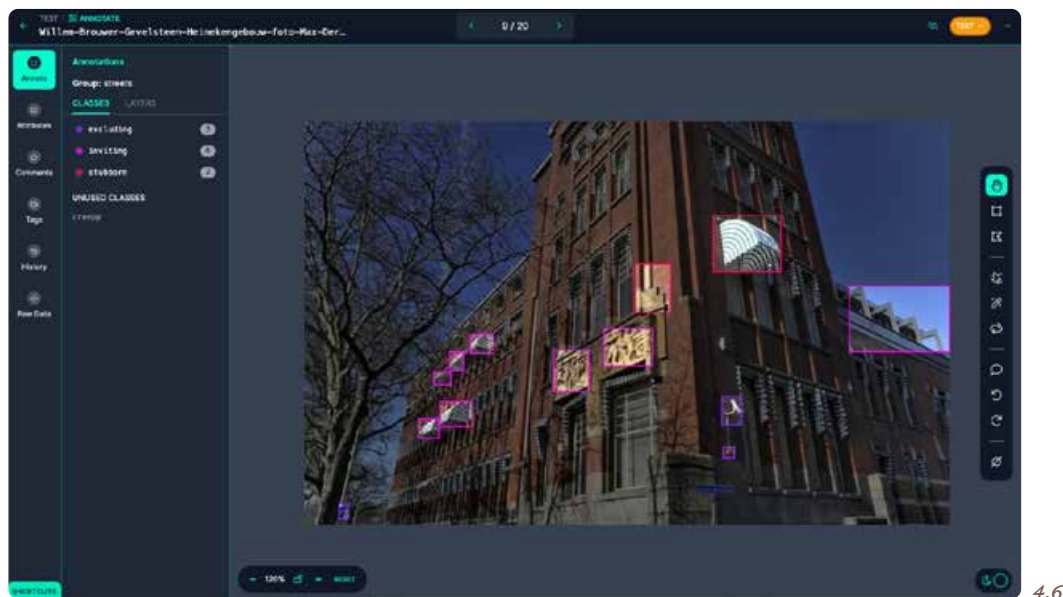
How to annotate a fascination?

YOLOv5 sees an image as a collection of parts. It scans a scene and identifies separate elements within it. To work with this model, we had to see our data the same way. This is why we chose images of scenes rather than isolated objects: a garden photograph contains plants, paths, shadows, edges; a surrealist painting contains figures, backgrounds, objects in relation. Even an isolated object, in theory, could be broken down into parts, but scenes made this way of looking more intuitive.

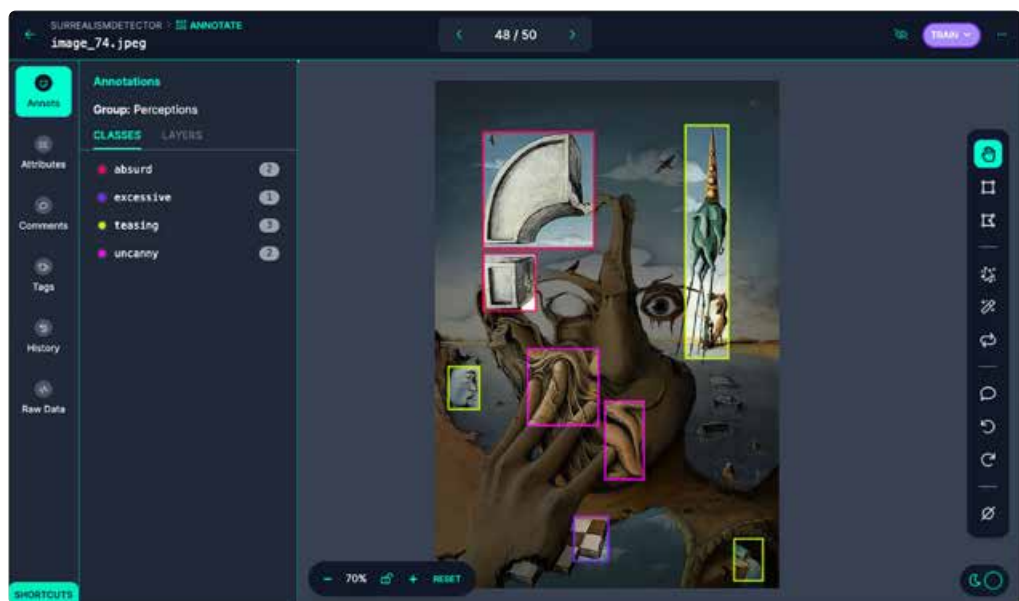
To annotate the data, we identified parts of these scenes that conveyed one of our predefined labels. We used Roboflow (Dwyer et al., 2022), an online annotation tool with bounding box functionality to segment parts of an image and assign our labels to them. The coordinates of these labeled bounding boxes could then be used to train our personal model. An example of the annotation tool interface is shown in *figure 4.6*.

This process required some getting used to. A feeling may arise from viewing a whole image, but annotation requires identifying a specific part. We had to ask: where exactly does this feeling come from? Is it the shadow, the figure, the empty space beside it? Often we could not fully separate content from context. A figure might feel ‘uncanny’ because of its relation to the surrounding space, yet the bounding box had to go somewhere. The tool forced a way of looking at our image data, that was not naturally how we would see images in general. This was uncomfortable at first, but it became generative: choosing where to place the box required us to make our reasoning explicit, and it would allow to really look at an image we were fascinated by in much detail.

We moved slowly through our datasets. Gijs and I each spent three to four hours per day over four days; Maurik spent two afternoons, resulting in a smaller annotated dataset of 124 images compared to our 300 each. Maurik found the annotation process tedious in ways that Gijs and I did not, which affected both the size



4.6



4.7

of his dataset and his engagement with that stage of the method. Yet we all spent considerable time with individual images, lingering longer on them when it felt necessary.

● Step 4: Creating an Object Portrait

How to start a conversation with a subjective algorithm?

After annotating our datasets with subjective labels, we trained three individual object detection models based on our personal fascinations: surrealistic art, gardens, and public spaces. The training process established patterns between the annotated images and their corresponding bounding box coordinates, taking approximately two hours per model.

Once trained, we faced the question of how to use these models to reflect on our own work. Our models had learned to recognize elements in images that resembled a surrealist painting, a garden, or a public space. We could have simply fed them existing images of our work, but this felt too passive. Instead, we chose to create new images specifically for the models to analyze: compositions that would represent our practices and the objects and materials central to them.

We called these compositions Object Portraits. They were carefully arranged still life artworks, composed in a manner consistent with our models' training datasets. I arranged surrealist objects from my collection and played with surreal perspectives and camera angles, echoing the compositional style of Magritte and Dalí. Gijs recreated a garden scene indoors, placing his own objects on a grass mat to play with garden-like features, evoking the aesthetic his model had learned. Maurik placed many of his artworks in an outdoor space. By incorporating stylistic elements from the training data into our portraits, we gave the models something they could 'read' through the visual language they had learned. We decided to create video portraits that could integrate small movements or changes, allowing us to place or remove objects and observe whether this would influence the model's predictions or alter the compositions and detections. The object portraits therefore occupied a space between still image and video (see figure 4.8, 4.9 and 4.10).

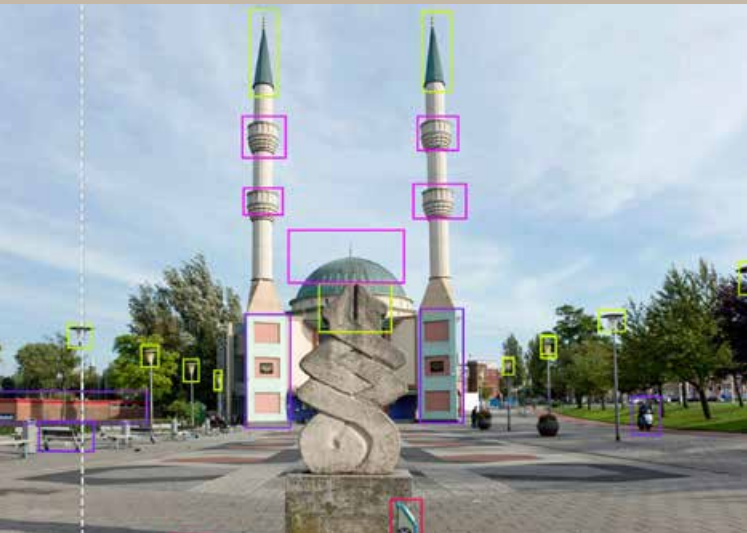
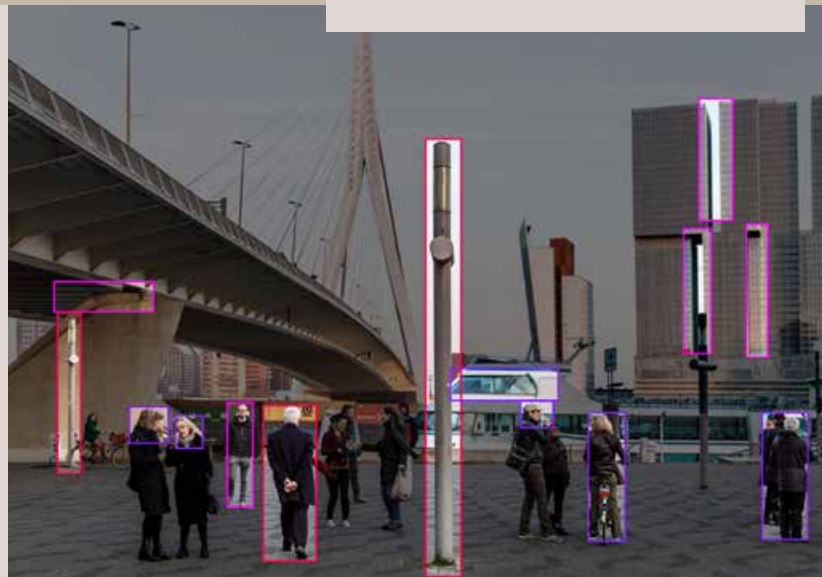
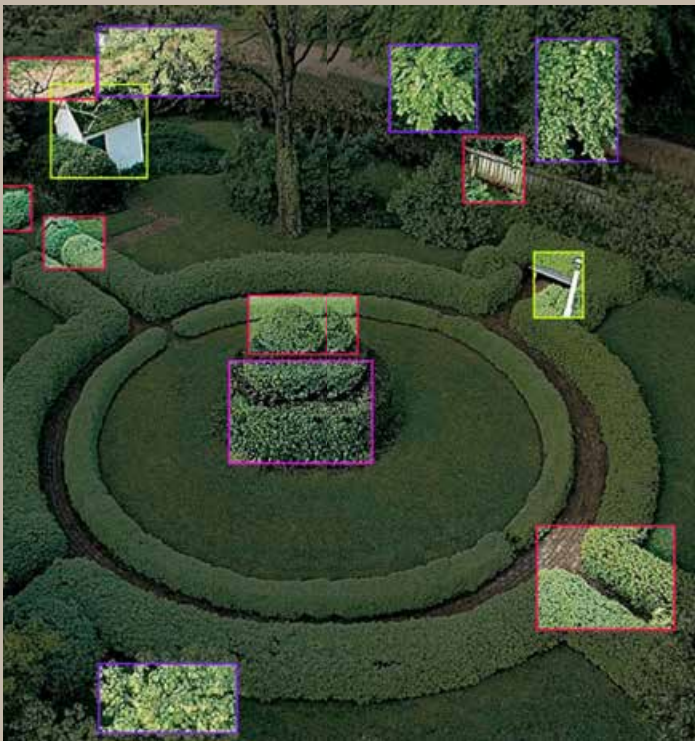
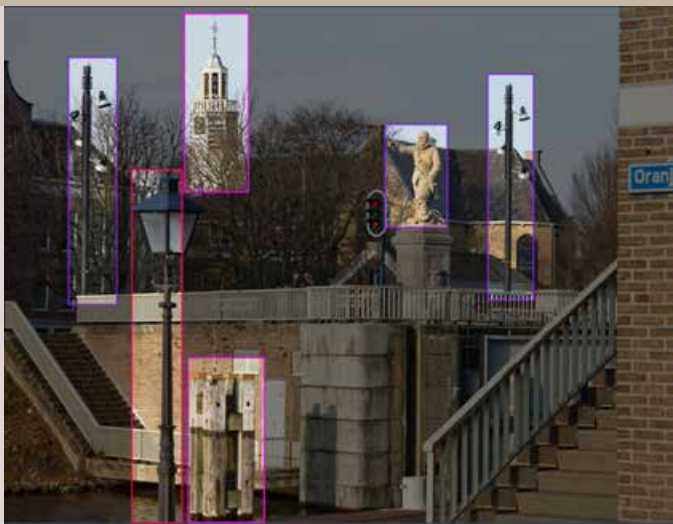
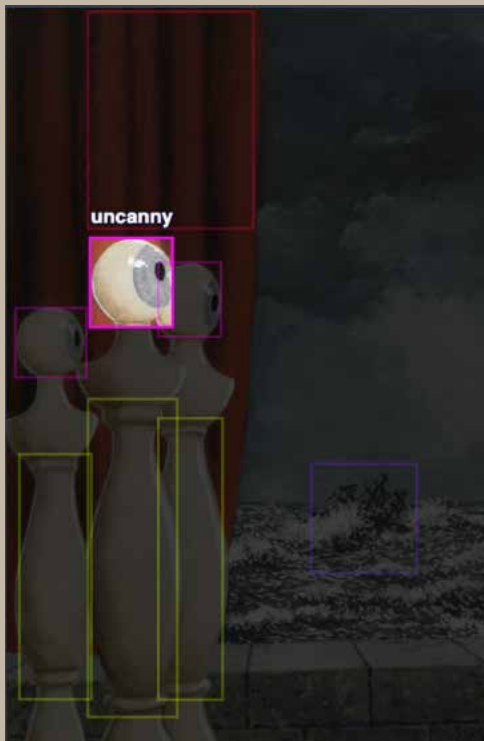
Creating the Object Portrait gave us a clear goal that reframed the whole process. We were making something that would be 'read' through the lens of our own training. The Object Portrait became the place

← Figure 4.6

The Roboflow annotation interface showing how subjective labels were applied to specific regions within images. Each bounding box assigns one of the four personal labels to a particular visual element, here for an image of the dataset of Maurik.

← Figure 4.7

The Roboflow annotation interface, here showing an image of the dataset of Vera.



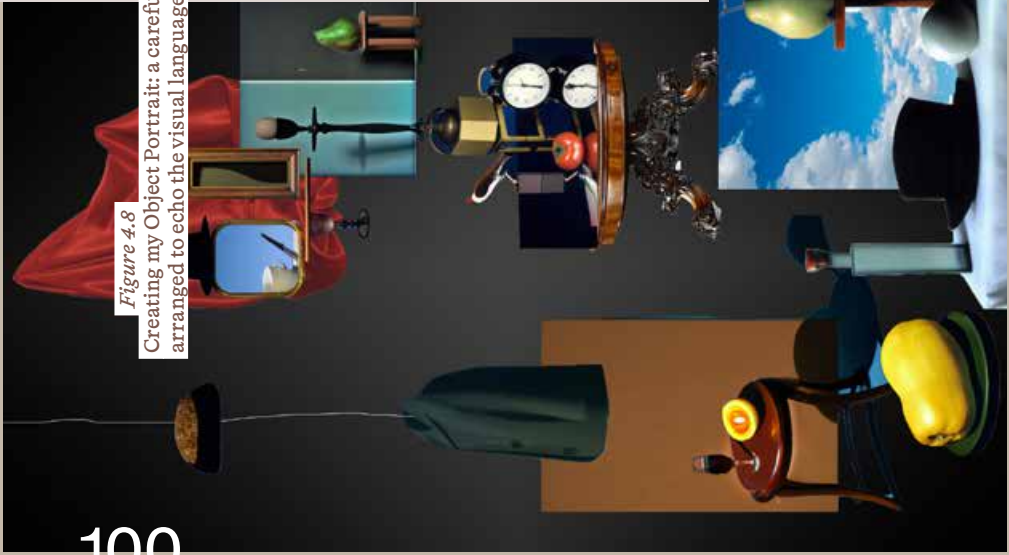
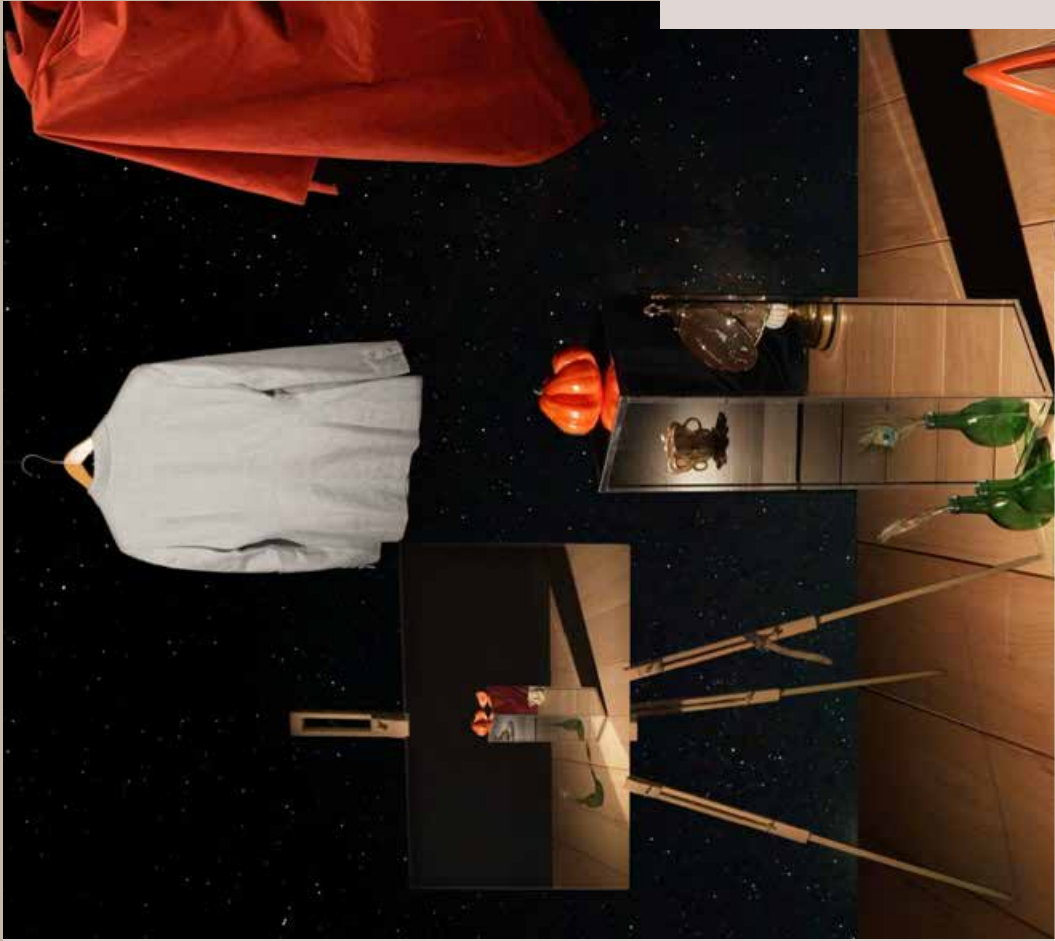


Figure 4.8

Creating my Object Portrait: a carefully composed still life of surrealistic objects from my personal collection, arranged to echo the visual language of my training dataset while representing my own making practice.





Figure 4.9

Creating *Gijs’ Object Portrait*: a carefully composed still life of objects from his personal collection, arranged to echo the visual language of his training dataset while representing his own making practice.

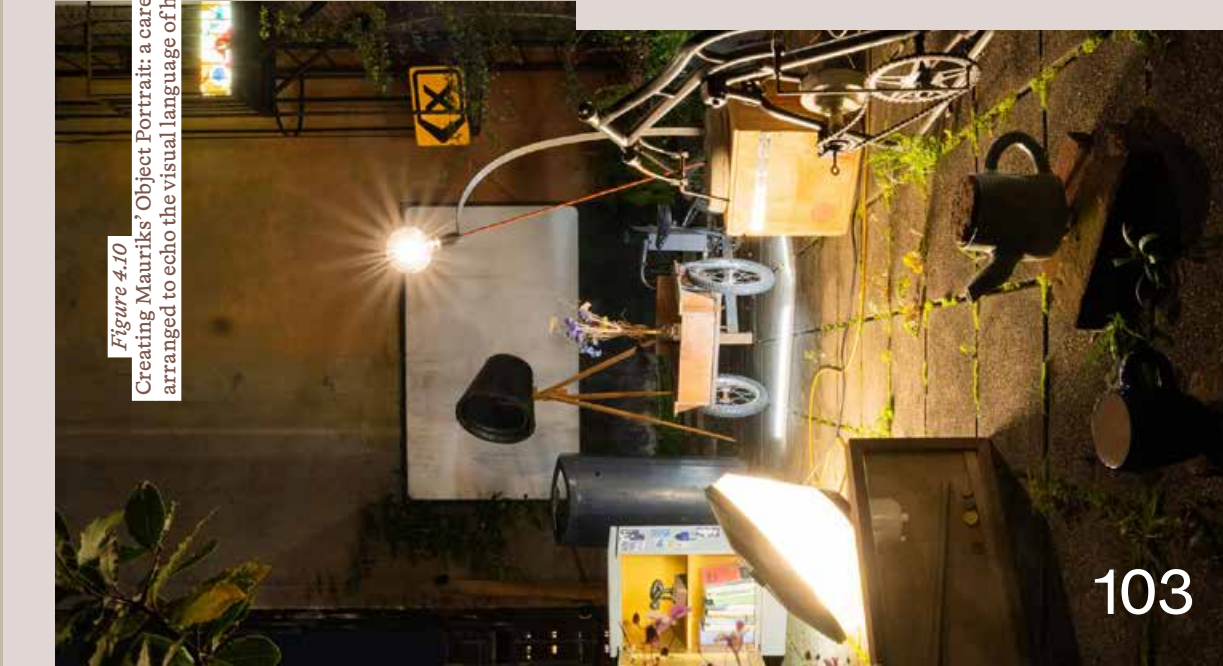


Figure 4.10

Creating Mauriks' Object Portrait: a carefully composed still life of objects from his personal collection, arranged to echo the visual language of his training dataset while representing his own making practice.

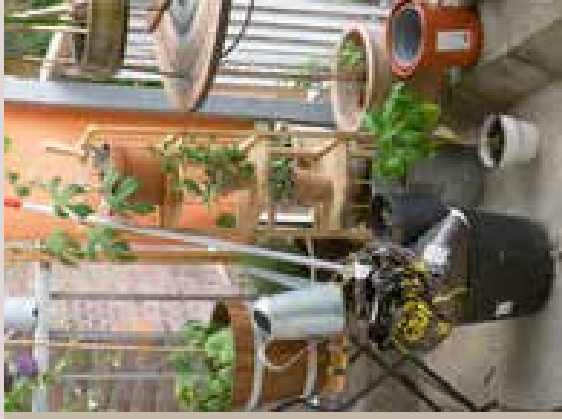




Figure 4.11

The three completed Object Portraits created for AI interpretation. Each composition reflects both the aesthetic of its creator's training dataset and the objects meaningful to their practice. Left to right: Vera, Maurik and Gjijs.



where our earlier decisions would show up as interpretation: how we had collected our data, which labels we had chosen, where we had drawn our bounding boxes. All of this had shaped what the model could now ‘see.’

● Step 5: Testing the model
What does a subjective model see?

○ *Training the models*

We trained our models on the Industrial Design Engineering department’s computing cluster at TU Delft, equipped with three NVIDIA RTX 3090 GPUs (see *Appendix A.2* for full hardware specifications). We configured each model to run for the necessary training epochs; the system stopped automatically when training completed. Training times ranged from 36 minutes to just over an hour and a half, depending on dataset size (see *table 4.1*).

The three models learned to associate visual patterns with our subjective labels, though with notably different results. Mean average precision (mAP) measures how accurately a model locates and classifies objects: a score of 100% would mean perfect detection, while random guessing on four labels would yield around 25%. Our models achieved modest scores ranging from 7% to 20% (see *table 4.1*). These differences potentially reflected the nature of our source material. My surrealist paintings share recurring visual motifs: bowler hats, floating objects, melting forms, consistent color palettes. The model could latch onto these patterns. Gijs’ garden photographs and Maurik’s public space images were more visually diverse, potentially giving their models less consistent material to learn from. Differences in annotation consistency may also have played a role.

In the context of this work, what does ‘accuracy’ mean when there is no ground truth? In conventional object detection, a model is correct when it matches an external standard: a cat is a cat, a stop sign is a stop sign. Our labels had no such external reference. When my model predicted ‘teasing’ at 73% confidence, correct according to whom? The only ground truth was my own prior annotation, a subjective judgment made on a particular day, in a particular mood. The model’s ‘errors’ were measured against my own inconsistency. A low accuracy might mean the model failed to learn, or that I failed to label consistently, or that categories like ‘absurd’ and ‘uncanny’ blur into one another. We suspected some combination of all three, and mainly, this uncertainty did not trouble us. We were

➤ *Table 4.1:*
Overview of the three training datasets, models and training performance

	Vera	Gijs	Maurik
Fascination	Surrealist art	Gardens	Public space
Images annotated	378	300	124
Labels	absurd, excessive, teasing, uncanny	filthy, glorious, lonely, luring	creepy, excluding, inviting, stubborn
Training epochs	200	300	338
Training time	43 min	1h 37 min	36 min
Model accuracy (mAP)	20%	7%	7%

interested in whether training a model on subjective labels, and then seeing its responses to our own work, could prompt reflection. The metrics that matter for conventional AI development did not apply here. Though it is interesting to present this data nonetheless: the differences between our three models hint at questions about visual consistency, annotation patterns, and what it takes for a machine to learn something as slippery as a personal judgment (see *Appendix A.2* for more detailed training results).

○ *Analyzing the Object Portraits*
After training, we used our models to analyze our Object Portraits (see *figure 4.11*) models identified elements in our compositions and assigned our subjective labels with confidence scores between 0 and 1. A score of 0.86 ‘luring’ meant the visual features strongly resembled what had been labeled ‘luring’ in the training data. Our question was not whether these predictions were accurate, but whether they could assist reflection. Did we agree with the model’s readings of our own work? Did its interpretations surprise us, confirm what we already knew, or reveal something we had not noticed? Was this the reflection we were looking for?

4.7
Insights: The process as the portrait

This project was where the central insight of this thesis first emerged. When we began, we expected that the trained model’s analysis of our Object

Portraits would be the primary reflective moment: the AI mirror showing us something about ourselves through its interpretation of our design work.

It did not seem to be so. While the final output provided a starting point, the insight that emerged was that the most profound reflection came from the training process itself. By comparing our experiences across our distinct fascinations (gardens, public space, surrealism), patterns became visible. The following insights describe where in the training process this reflection occurred.

● *Insight 1: Reflection is in the annotation*

Although it was intriguing to see the AI label objects in our final compositions with probability scores, the implications of these predictions alone were unclear. What did it mean that the model detected ‘absurd’ with 0.05 confidence? To understand where the predictions came from, we had to go back into our training datasets and reflect on how we had labeled them.

When we reflected on the annotation task afterwards, we all recognized a similar pattern: labels that seemed clear at the start became ambiguous during annotation, requiring us to articulate what we meant visually. Because we annotated with bounding boxes rather than whole-image descriptions, we had to locate our feelings in specific parts of each image. This granular attention made the process slower, but also more considered. What Miceli et al. (2020) call annotation as ‘sense-making’ became visible in our own experience: by annotating our own data with our own labels, we reclaimed the interpretive work that industrial annotation typically minimizes.

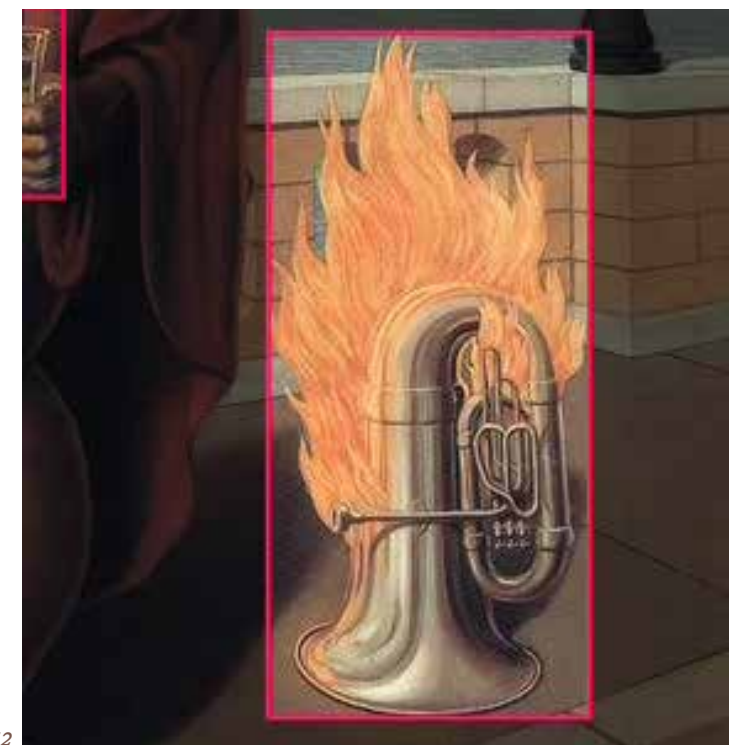
I could not fully explain what I meant by ‘absurd’ at the beginning. At times it felt like a joke, overlapping with my label ‘teasing.’ Through repeated use, following intuition, I gradually developed understanding: I realized I saw all body parts in surrealistic art as ‘uncanny,’ and ‘absurd’ as representing a visual joke (see figure 4.12). Once this became clear, my annotation shifted: I started actively looking for body parts to label as ‘uncanny,’ rather than waiting to stumble upon the feeling. Understanding changed how I saw the images. Gijs considered well-raked and organized gardens ‘lonely’ and blooming flowers ‘glorious,’ but came to understand these distinctions through the act of repeatedly applying labels (see figure 4.13). Maurik realized that ‘stubborn’ described something out of the ordinary in a scene, particularly elements that seemed to resist or do something different from their surroundings. Teaching the machine our subjective associations helped us understand our own visual concepts. Annotation forced tacit visual concepts into explicit articulation, and that articulation, in turn, shaped what we noticed.

● *Insight 2: Subjective judgment is dynamic while annotation is static*

A tension ran through the annotation process: our feelings about images sometimes changed while we were labeling them, but the model needed

→ *Figure 4.12:* Comparison between Object Portrait detection (bottom) and training data annotation (top) for my label ‘absurd.’ The model detected a ‘visual joke’ form in my Object Portrait with this label.

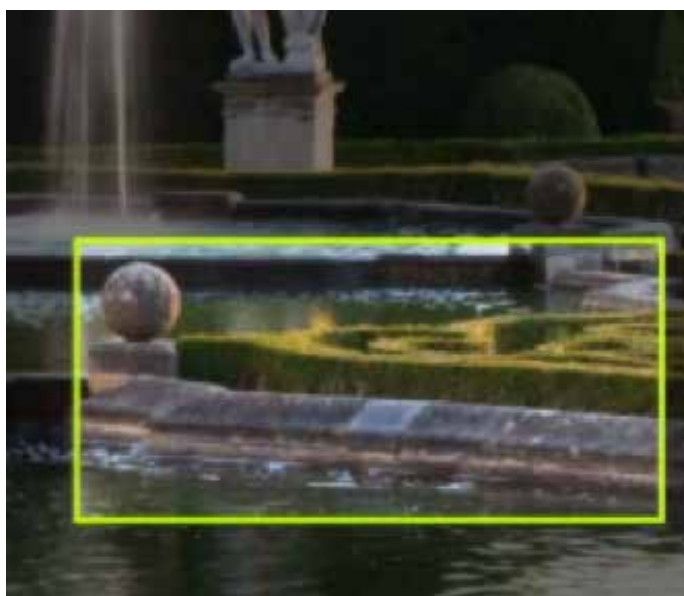
4.12



4.12



4.13



4.13

consistent patterns to learn from. Gijs experienced this tension most explicitly. He initially labeled well-raked gardens ‘lonely,’ but as he continued annotating, he started finding them appealing. He faced a choice: stay consistent for the model, or be honest to his shifting judgment. He chose the second, and started labeling those instances ‘luring’ instead.

This created confusion for the model but clarity for Gijs. The model learns by finding visual patterns that correlate with labels. When similar-looking images receive different labels, the pattern becomes inconsistent, and the model struggles to learn what the label ‘means’ visually. The inconsistency revealed that Gijs’s relationship to well-raked gardens had shifted during the annotation task itself.

In annotation research, this is called ‘annotation drift’. Chang et al. (2017) describe how annotators’ interpretations of categories can shift over time, especially during long labeling sessions, leading to inconsistencies in the training data. In industrial contexts, this is treated as a problem: protocols are designed to minimize drift and maximize agreement between annotators. For us, it was productive. The drift captured something real: our judgments were not fixed, and the friction of having to choose revealed when and how they changed.

← Figure 4.13

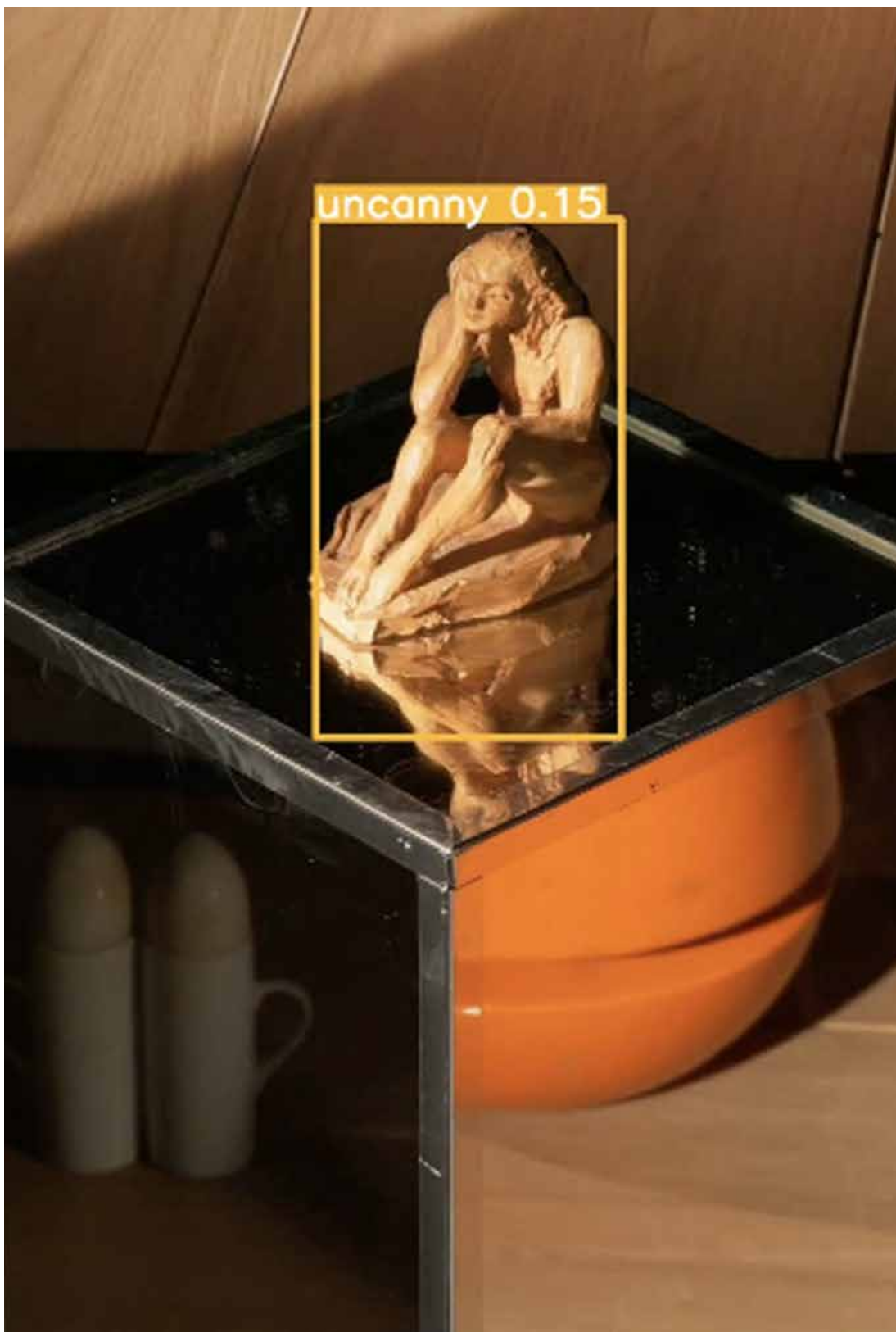
Comparison between Object Portrait detection (Top) and training data annotation (Bottom) for Gijs’s label ‘lonely.’ The model learned to associate well-raked, organized garden elements with loneliness based on his annotation choices, then applied this interpretation to objects in his Object Portrait.

We all experienced this tension differently. In some cases, like Gijs, a switch was made. In others, we aimed to maintain consistency. I experienced similar friction when distinguishing ‘absurd’ from ‘teasing,’ but chose differently than Gijs: I kept my original patterns, prioritizing what the model needed to learn even when my own judgment had started to shift. Either way, balancing subjectivity with the model’s technical requirements surfaced questions about how fixed or fluid our aesthetic judgments really were.

● *Insight 3: Creating the Object Portrait as a reflective journey*

Creating the Object Portrait evolved into a reflective journey of its own. To create the portrait, we carefully selected objects using several criteria. One basic criterion was the technical capabilities of the model: we had to ensure that the objects we selected could be detected. But the more important criteria involved our curiosity, our attachment to the objects, and the personal connection we had with them as designers. We selected objects that represented us and to which we felt a strong attachment, whether they were objects we had created ourselves or cherished items in our collections.

We were driven by curiosity, wondering if the model would recognize these objects as manifestations of the fascinations we had trained it on. We each chose objects that sparked a question: ‘What would the AI perceive in this object that is intimately connected to me?’ Gijs pondered how



4.14

the model would interpret an object he felt he had created with excessive control. If the model's interpretation differed, what would that reveal about his style, fascinations, or judgment? I selected surrealistic objects from my own collection, curious whether the model would detect the same 'absurd' qualities I had been training it to recognize. Maurik placed his artworks in an outdoor public space, wondering how his trained model would read his own response to that environment.

When a bounding box appeared around one of Gijs's objects labeled 'lonely,' it prompted a question: "*Do I agree? And if not, why did the model see it that way?*" These moments sent us back into our training data, looking

for the annotations that had led to this interpretation. Sometimes we recognized our own logic; sometimes the model's reading surprised us; sometimes it raised doubts about whether the objects we had selected truly reflected our fascinations at all.

When we finally saw our Object Portraits analyzed with bounding boxes and confidence scores, the AI's reading of the image became visible. The model offered a perspective on what was already there, and we could trace that perspective back to our own choices: our data, our labels, our annotations. Seeing our training decisions manifest as interpretation brought together what we had built through the process.

4.8 *Discussion*

The three insights point in the same direction: reflection emerged from the training process itself. This section explores what this finding means in relation to broader discussions about annotation, slowness in AI interaction, and the value of working comparatively.

● *Leveraging annotation bias*

At the start of this chapter, I described our approach as deliberately amplifying subjectivity rather than eliminating

it. The insights above show what this looked like in practice: labels that seemed clear became ambiguous, judgments shifted during annotation, and the friction of having to choose surfaced patterns we had not consciously recognized.

These experiences connect to broader discussions in AI research about the traces that human annotators leave in training data. Annotation bias describes how individual assumptions, interpretations, and inconsistencies shape the datasets that models learn from (Baumer et al., 2020). A related phenomenon is annotation drift, which describes

← *Figure 4.14*

Detection of "uncanny" in my Object Portrait. The model applied this subjective label to a sculptural figure, reflecting patterns learned from my annotations of surrealist paintings where I had associated human and creature forms with uncanniness.

how annotators' interpretations of categories can shift over time, especially during long labeling sessions (Chang et al., 2017). In industrial annotation contexts, both are treated as problems to minimize. Protocols use multiple annotators and majority voting to suppress individual variation and establish what counts as 'correctly annotated' (Mujtaba & Mahapatra, 2019). The goal is consistency and agreement; variation between or within annotators is noise to be filtered out.

Our process worked differently. We had no external ground truth to measure against, only our own prior judgments. And rather than treating variation as error, we paid attention to it. When Gijs' feeling about well-raked gardens shifted from 'lonely' to 'luring' halfway through annotation, the drift itself became useful. The inconsistency suggested that his relationship to control in gardens had changed during the task, something he might not have noticed without the pressure of having to label consistently. Similarly, when I began actively searching for body parts to label as 'uncanny' after recognizing my own pattern, my annotation shifted from open exploration to active confirmation. Understanding a label changed how I applied it, and the drift captured something about how meaning develops through the act of repeated labeling.

We treated the drift as material for reflection. The friction of choosing labels, locating feelings, maintaining or abandoning consistency surfaced questions about our own judgments. In this sense, we were leveraging the bias rather than suppressing it.

● *Slowness as a condition for reflection*

This process unfolded over a period of roughly three months: collecting datasets, developing the quadrant tool, annotating hundreds of images, training models, creating Object Portraits, and reflecting on the results together. This duration stands in contrast to the fast workflows associated with generative AI tools, where a single prompt produces output in seconds.

Chapter 2 introduced the principles of Slow Technology as articulated by Hallnäs and Redström (2001), who argue for technology that supports reflection and contemplation rather than efficiency. The Objective Portrait method operationalized these principles in a specific way. The repetitive nature of annotating hundreds of images, one bounding box at a time, created a rhythm that allowed patterns to emerge gradually. We tried not to rush the annotation. Each image required a decision, and those decisions accumulated into something like a conversation with myself about what I actually meant by labels like 'absurd' or 'uncanny.' Returning to earlier annotations, questioning whether a label still felt right, noticing when feelings had shifted: these reflective moments required duration. The slowness itself was what allowed reflection to occur.

● *Working comparatively*

Working alongside Gijs and Maurik shaped how the insights emerged. By comparing our parallel processes, we could observe that certain patterns (labels becoming ambiguous,

judgments shifting during annotation, anticipation shaping object selection) appeared across our different practices. This suggested that these were features of the method rather than specific to one person's experience.

The comparative approach also created contrast. Gijs chose to follow his shifting judgment and relabel images; I chose to maintain consistency for the model. Observing these different responses to the same tension helped articulate what was at stake in the choice. Similarly, comparing our labels side by side (uncanny, absurd, teasing, excessive; glorious, filthy, lonely, luring; creepy, inviting, excluding, stubborn) made visible how differently we each structured our subjective responses, even when following the same method. That said,

the depth of engagement varied across the three of us. Gijs and I maintained ongoing dialogue throughout the process, returning to our annotations and outputs in conversation. Maurik approached the project more independently and found the annotation process less generative. His smaller dataset (124 images compared to our 300 each) reflected this difference. His process provided valuable comparative data about how the method functions across different aesthetic domains, but the richer reflective insights in this chapter emerged from those of us who found the annotation process itself meaningful. This suggests that completing the steps does not guarantee reflection; the practitioner's willingness to dwell with the process appears to matter.

4.9 Conclusion

The contribution of the Objective Portrait project is a process: a hands-on method to reflect on fascinations by engaging with the construction of an AI. The insights showed that the most profound reflection emerged from the training process itself: the shifting judgments, the emergent patterns of annotation, and the confrontation with one's own inconsistencies.

One aspect of the method did not work as we had hoped. We expected the model's analysis of our Object Portraits to be a significant reflective

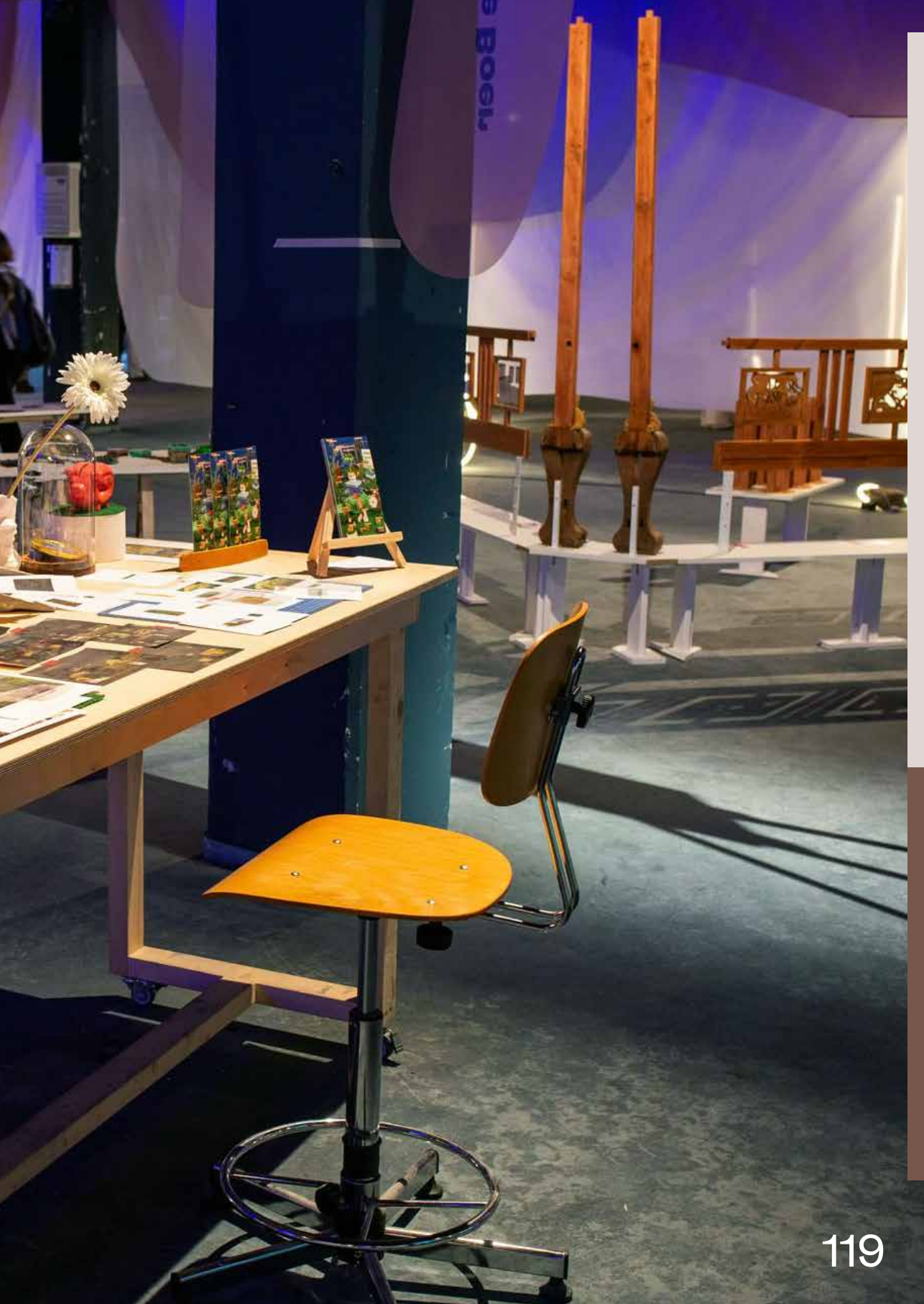
moment, the AI mirror showing us something about ourselves. In practice, the predictions alone were not challenging enough to prompt deep self-reflection. At times, we understood the AI's evaluation, but the implications remained unclear. To truly understand a prediction, we had to return to the dataset and examine our annotation behavior. The reflection happened in this return, not in the prediction itself. This suggests that the value of the method lies in the circular movement



Dutch Design Week 2022



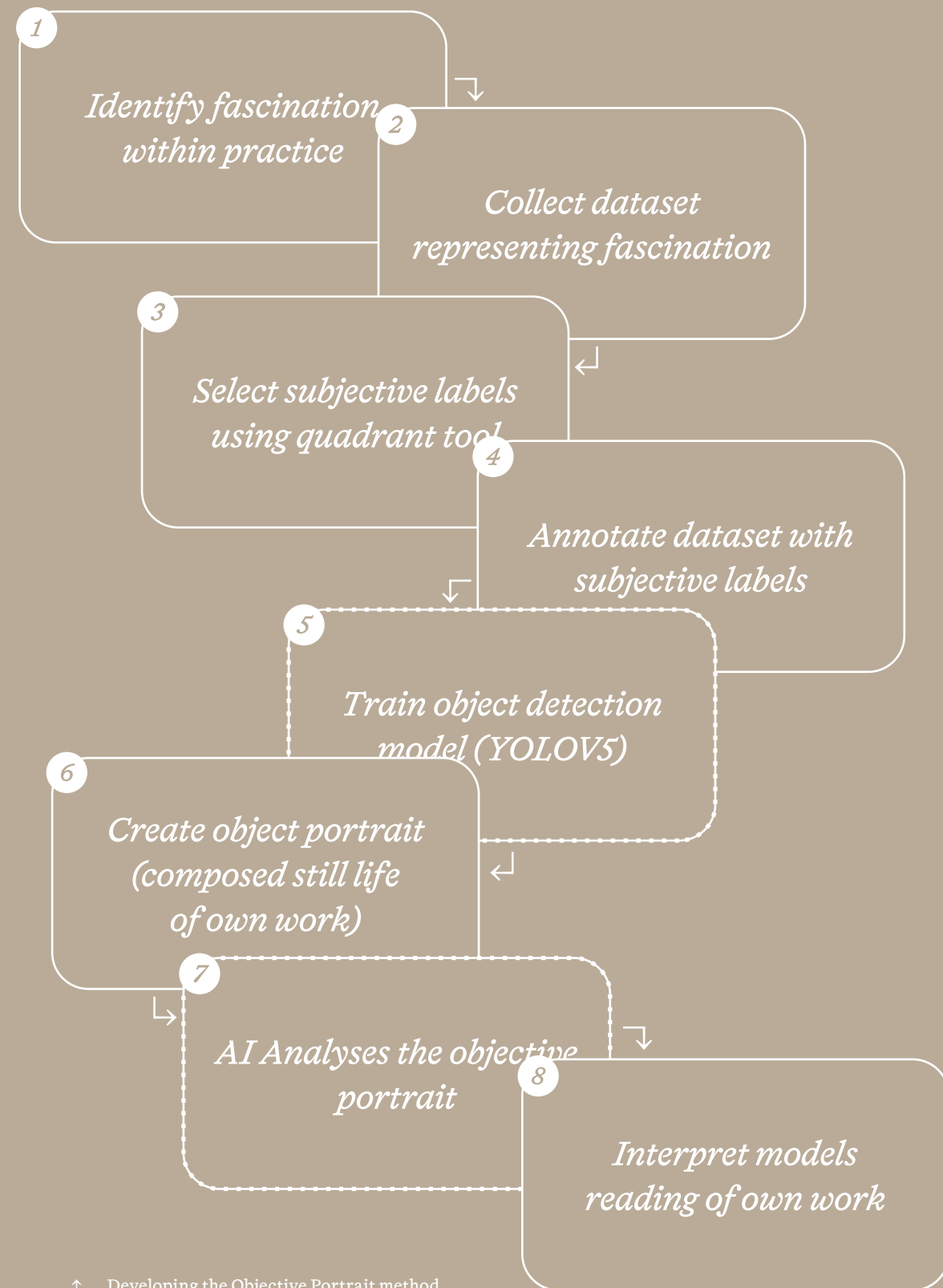
Dutch Design Week 2022



between training choices and output, rather than in the output alone.

This also raised questions about the labels themselves. Would different labels, more assertive or more opinion-based, produce more provocative predictions? Would a different relationship between annotator and subject

matter create different conditions for reflection? We had trained models on external fascinations: surrealism, gardens, public space. What would happen if participants annotated image material closer to their current creative work?



5. The Objective Portrait Workshop

The previous chapter established the Objective Portrait project as a hands-on process for developing a reflective practice with an AI model. Developed with two designers and myself, it demonstrated how the act of training an AI could become a process for self-reflection. However, the insights from that study were deeply situated within the experiences of us three

individuals who knew each other and explored the process on our own terms. We possessed established creative practices to reflect upon. Design students, in contrast, are actively forming their creative identities while learning to become reflective practitioners. To understand the broader viability of this process as a transferable method, it became necessary to test it with a larger group of participants in a new setting. Design education offered a suitable context: a setting where the development of reflective capabilities is a core pedagogical goal, and where novice practitioners with an interest in technology could benefit from this hands-on approach. This chapter therefore documents the adaptation of the Objective Portrait

chapter from an exploratory research method into a structured three-day educational workshop.

5.1

Reflective AI use in design education

Design education has traditionally cultivated reflective capacity through iterative processes that encourage students to examine their thinking and creative decisions over time. In design studio pedagogy, reflective practice is both the learning objective and the pedagogical strategy: students learn to design by doing, with reflection developed through repeated cycles of making, feedback, and revision (Milovanovic & Gero, 2020; Lousberg et al., 2020). A student might spend weeks iterating on a single concept through sketches and prototypes, with each studio critique prompting reconsideration of assumptions and alternatives. This pedagogical approach assumes that designers learn by developing critical awareness of how and why they make creative choices. Yet as I worked as a tutor in design education throughout my PhD during the rise of popular generative tools like ChatGPT and Midjourney, I witnessed how the integration of efficiency-driven AI tools introduced tensions with these pedagogical goals.

As established in Chapter 1, contemporary AI tools often lead to

design fixation, where the speed and perceived finish of their outputs cut short the exploration of design spaces (Hwang, 2022). Designers experience creativity exhaustion, feeling pressured to match the pace of generative AI outputs (Li et al., 2024). Additionally, it has been reported that the technical opacity of these systems limits designers' sense of control and agency over the creative process (Yang et al., 2020; Hu et al., 2025). While these challenges affect professional designers, they become particularly acute in educational contexts. Sandhaus et al. (2025) studied the unprompted use of generative AI by student groups in an HCI design course and found that despite benefits in enhanced creativity and faster design iterations, the way students engaged with these tools often resulted in shallow learning, particularly when not accompanied by reflection on how and why they used them. This is concerning for design education in the context of HCI practice, where students are expected to develop technical skills alongside critical sensitivities toward the social, ethical, and experiential

dimensions of technology. The general technical opacity, and even more so the ‘opacity at runtime’ of automatic generation (Burrell, 2016), hinders the in depth, reflective integration that design education seeks to cultivate.

While research on AI technologies within educational contexts has emerged, much of this work has focused on K-12 education, exploring how children and adolescents can develop AI literacy through tangible and constructionist approaches (Bilstrup et al., 2025; Lindrup et al., 2025; Benjamin et al., 2024). HCI-related design education itself, with its unique requirements for cultivating both technical understanding and critical reflection on technology’s role in practice, has received comparatively little attention.

Existing work has begun addressing these challenges in design education. Murray-Rust et al. (2023) developed ‘Grasping AI’ exercises that use enactment, metaphor work, and role playing to help design students develop critical perspectives on AI as a socio-technical system. This aligns

with broader efforts emphasizing experiential and constructionist educational approaches to AI and machine learning technologies. These valuable approaches successfully create reflective spaces around AI’s socio-technical dimensions, examining its social implications and broader systemic effects. However, existing approaches primarily address the technical opacity and social implications of AI systems. They do not specifically address how the efficiency focused paradigm of current AI tools undermines the slow, reflective practices foundational to design education.

This shows an opportunity to explore engagement with these technologies through slower, more reflective interactions. Building on the slow technology principles established in Chapter 2, we wanted to investigate whether treating AI’s technical pipeline as distinct sites for engagement could support the development of both reflective practice and critical technical awareness in design students, similarly to how it served us in the previous chapter.

5.2

Adapting the method for design students

The Objective Portrait Workshop was designed as an educational adaptation of the method documented in Chapter 4, responding to the concerns described in the previous section. Where consumer-facing AI tools present a seamless interface

from prompt to output, the workshop creates intentional slowness by separating data collection, annotation, training, and interpretation into

→ *Figure 5.1*
The students of the workshop at Design Academy Eindhoven



5.1

distinct activities across three days. This structure addresses the opacity that leads to shallow learning: by engaging with each phase separately, students can develop a sense of how their curatorial and interpretive choices propagate through the technical pipeline.

The workshop pursued two learning objectives. The first, reflective positioning, encouraged students to deepen their relationship with their emerging research interests through the external lens of data curation and annotation. By articulating what draws their attention and why, students could develop more conscious awareness of their aesthetic judgments and design values. The second, critical technical awareness, helped students recognize how their own choices shape what the model eventually does, cultivating designerly agency over AI systems. The extended time spent with each phase creates what Hallnäs and Redström call ‘time presence’ (Hallnäs and Redström, 2001): presence in the activity itself, where the slowness of annotation and interpretation opens space for reflection.

These objectives required adaptation from the professional context documented in Chapter 4. Where we as expert designers engaged in open exploration following established fascinations, students were still forming their creative identities and needed structured scaffolding to engage with each phase without losing sight of the whole. The workshop therefore provided clear phases and deliverables while maintaining space for personal interpretation. Students would train custom object detection models on personally meaningful datasets, then use these models to analyze their own creative work in an ‘Object Portrait,’ treating this entire process as an opportunity to reflect both on their emerging creative identities and on the technology itself.

This chapter documents the workshop through: the methodology (Section 5.3), a day-by-day description (Section 5.4), reporting on the reflective practices that emerged among eleven design students (Section 5.5), and the implications for integrating Reflective AI into design education (Section 5.6).

5.3 Methodology

The following sections document observations from a group of design students navigating the AI-based reflective process developed in Chapter 4. Building on the practice-based research approach established in

Chapter 2, this workshop created a structured scaffold enabling multiple participants to engage in first-person exploration simultaneously. This methodological shift offered observational distance: by watching students

Name	Topic/Question	Image Dataset	Labels
Eva	Humor in memes	Memes	ridiculous, amusing, appoling, facetious
Aisha	How would AI react to humanistic beliefs, cultural oppression, and social etiquettes?	Religious events	purity, salient, dominant, synchronous
Carlo	Queer identity expression	People in extravagant outfits or drag fashion	enticing, flamboyant, post-human, cloying
Sascha	Vibes at a festival	Festival stages & concert stages	commercial saturation, convivial, lifeless, otherwordly
Yann	What can the refusal of advancement actually look like?	Abandoned buildings, artworks of Gordon Matta-Clark and Alvin Baltrop	devotion, inconsistent, pretentious, rebelling
Borisz	Rubens' paintings and compositions	Rubens' paintings	majestic, repulsive, sublime, terrifying
Pepa	Age limitations on digital mediums	Children's cartoons from age 4–6	comforting, disruptive, frightening, overwhelming
Jassir	"Morphed by force"	Faces morphed by force	forged, mythical, reject, thrill
Rey	AI comedian	Wildlife photography	compromise, drastic, gloomy, mysterious
Yuky	Judgment in fashion design	Models on runways wearing biomimicry fashion	charismatic, generic, premature, tasteless
Agustina	Birthdays	Stock photography of birthday celebrations	isolated, neglect, safe, vain

navigate the process, I could examine how principles from my earlier work translate to the practices of novice designers. The workshop was designed as a step-by-step process that was prepared and guided, yet broad enough for students to fill with their own interests.

The primary insights emerge directly from participants’ reflections as they integrated AI into their own explorations. The eleven Object Portraits created during the workshop serve as visual records of each student’s reflective inquiry, yet the knowledge here is generated through careful interpretation

of the processes that led to those portraits, documented throughout. Rather than drawing from a single designer’s practice, patterns across multiple practices become visible. Ultimately, this workshop both validates and extends the conceptual frameworks from earlier chapters, examining how Reflective AI functions within design education.

5.3.1 Participants

We recruited eleven students from a bachelor’s program in digital systems and new media at the Design Academy Eindhoven. All participants had diverse cultural backgrounds and a shared interest in technology. The workshop’s timing was strategic, coinciding with the early phase of the

↑ Table 5.1
Overview of workshop students and their chosen topics, image datasets and labels for annotation



students' academic year when they were tasked with exploring topics for their individual research interests. This alignment meant students could integrate the workshop directly into their ongoing coursework, selecting training datasets that reflected their personal areas of exploration.

The diversity of topics students chose to explore is shown in *table 5.1*, which presents each participant's research question, their chosen image dataset, and the four subjective labels they developed. We received informed consent from all participants, and at their request, their first names have not been anonymized as they retain authorship of their Object Portraits⁷.

5.3.2 Methodological components

Our primary interest was in understanding how a comparatively slow and deliberate setup could become reflectively meaningful for these students. To achieve this, we designed the workshop around three methodological elements that each create opportunities for students to pause, articulate, and examine their own judgments.

First, the choice of object detection as the technical system was intentional. We used the same object detection model, YOLOv5 (Ultralytics, 2023), as in the previous chapter. The pedagogical advantage lies in the accessibility

of its annotation process: students create multiple annotations with bounding boxes within individual images, a process that is inherently slow and deliberate. This contrasts with generative models where annotation involves generating a single sentence per image.

Second, the quadrant tool, adapted from Chapter 4 (*see section 4.6, step 2*), was used to intentionally decelerate the label selection process, encouraging students to develop deeper awareness of their subjective judgments. The tool's intentionally undefined 'push/pull' and 'warm/cold' axes required students to articulate why specific feelings occupied particular positions, forcing commitment to a vocabulary that would later serve as classification labels.

Third, the Object Portrait itself—the carefully composed image that students create to be analyzed by their trained model—serves as both the workshop's culminating activity and its conceptual anchor. Students could create either a still image or a short film for analysis. This artifact brings together the previous steps of data collection and annotation, transforming the technical exercise into a reflective dialogue with their own work.⁸

The day by day implementation of this process is detailed in the following sections.

⁷ Research data supporting this chapter are available in 4TU.ResearchData at <https://doi.org/10.4121/6ea01644-acc6-4478-a581-1f6788dec1f7>.

⁸ Note on terminology: The broader methodology is referred to as the Objective Portrait, while the specific visual artifact created by the student is referred to as the Object Portrait.

← Figure 5.2

Quadrant tool for determining labels with which to annotate their data, shown here as filled by Pepa, a participant interested in the topic of age limitations on digital media. The tool helped students think of labels as occupying opposing ends of two orthogonal axes, which helped them come up with minimally-overlapping labels.

5.4

Day by day:

The Objective Portrait workshop

Given our focus on integrating this approach with participants’ emerging design practices, we made a decision to have students complete certain preparatory steps before the formal workshop began. Specifically, we asked participants to engage in initial topic selection, identifying a core concern within their creative practice, fitting the early stage of the academic year, and preliminary data collection prior to attending. To create a balanced learning experience that allowed sufficient time for reflection while maintaining momentum, we structured the workshop across three distinct days, each with specific focus areas. This format provided natural pauses for consideration while ensuring participants progressed through all necessary stages of the Object Portrait creation process. The detailed structure and activities for each workshop day are outlined in *Table 5.2*.

● *Day 1: Preparing the AI model*
The first day was dedicated to building the subjective foundation for each student’s AI model. As mentioned, we asked the students to arrive at the workshop with a pre-collected dataset of images based on criteria we provided beforehand. First, the

images had to be thematically aligned with each participant’s personal research topic or question. Second, to ensure the effective functioning of the object detection model, the dataset needed to have at least 100 images with consistent sizes and styles (e.g., no mixing of photos with paintings or screenshots). This approach reflects what Abuzuraiq and Pasquier (2024) call a ‘*small data mindset*’: recognizing the value of carefully curated small datasets in creative domains, where the reduced scale allows practitioners to browse, manipulate, and appreciate the influence their training data has on the model. We embraced this mindset for pedagogical reasons. Where these models are typically trained on thousands of images, we chose a lower number to make the annotation task manageable for one person, prioritizing deliberate annotation of 100 images over fast-pacing through 400. This also gave the students a comprehensible overview of what goes into their model.

The day began with a group reflection where participants discussed the datasets they had chosen and how these related to their research interests, familiarizing the entire group with everyone’s topics (*Task 1.1*). Following this, in the first half of the day, we

Day	Agenda	Image Dataset	Duration	Purpose
1	Training the Object Detection Model Reflectively	1.1 Collect image dataset	Prior	Obtain thematically-grounded training data for the AI model
		1.2 Label Selection	1h	Crystallise emotions and subjective judgements about the dataset
		1.3 Dataset annotation	3h	Assign above labels to training data to inform AI model "interpretation"
		1.4 Individual Reflection	30m	Reflect on (change in) perception of chosen labels
2	Creating the Object Portrait	2.1 Object Selection	2h	Reflect on own practice and material
		2.2 Object Portrait creating	3h	Reflect on own practice in relation to the dataset and labels
		2.3 Collective reflection	30m	Gather insights on the process of creating an image for AI interpretation
3	Integrating the AI precitions	3.1 Editing Object Portrait based on AI output	1h	Reflect on expected and unexpected aspects of AI model interpretation
		3.2 Exhibiting work	1h	Adapting to changes in perspective based on AI model output
		3.3 Collective Interpretation	1h	Shared reflection on AI "interpretation" of their work
		3.4 Presenting the work	1h	Articulating the essence of the work and what they learned
		3.5 Individual reflection	30m	Reflecting on change in perspective on AI

tasked the students with creating four labels to annotate their dataset with (*Task 1.2*). These labels explicitly needed to be subjective, emotional descriptions of their feelings and judgments when looking at and thinking about their datasets. They began this process by closely examining 10 images from their dataset (*see figure 5.3*) and then used a handwritten form of the quadrant tool, which allowed them to iterate on their labels and served as documentation of that iterative process (*see figure 5.2 for an example of Sasha*). In the second half of the day, students annotated their image datasets using the labels they had

developed (*Task 1.3*). They employed the annotation software Roboflow⁹ to select specific parts of their images and assign the corresponding labels. This task took approximately two hours, during which we encouraged a slow and deliberate process. Participants were allowed to annotate as many images as they could within this time and could also revisit and modify their labels if needed. To conclude the day, we asked participants to reflect individually by annotating their quadrant tools with notes, focusing on how the labeling process influenced their perception of the labels and their research topics (*Task 1.4*).

↑ *Table 5.2*
Overview of the tasks, estimated duration, and underlying purpose over the three days of the workshop.

9 Roboflow Annotate: <https://roboflow.com/annotate>



5.3



5.3



● *Day 2: Creating the Object Portrait*
Day 2 centered on the creation of the Object Portrait, an artwork explicitly made for AI interpretation, as introduced in the previous chapter. We guided students through three key requirements for their portrait: (1) it must feature objects related to their personal practice (*Task 2.1*); (2) it had to mimic aesthetic elements of their training dataset to ensure model recognition (*Task 2.2*); and (3) it could be either a single image or a 10-second video.

Crucially, we asked students to craft this portrait during the workshop, rather than using a pre-existing piece. In our own experience described in Chapter 4, creating the Object Portrait gave us a clear goal that reframed the whole process: we were making something that would be ‘read’ through the lens of our own training. We wanted students to have this same framing. The students were given one day to craft their portraits (*see figure 5.4*). The day concluded with a group reflection session where participants shared their insights on the experience of creating an image for an AI model.

● *Day 3: Integrating the AI Model into the Object Portrait*

○ *Model Training*

Between Days 2 and 3, we trained an individual YOLOv5 model on each student’s annotated dataset using

the same technical setup described in Chapter 4 (see Appendix B for full specifications). All models were configured identically: 300 maximum epochs with early stopping (patience of 100 epochs), meaning training would halt automatically if performance stopped improving.

Table 5.3 summarizes the training outcomes. Model performance varied considerably, with mean average precision (mAP@0.5) ranging from 0.009 to 0.384. Training duration also varied: some models stopped early at around 120 epochs while others continued to 266 epochs before the early stopping criterion was met.

As established in Chapter 4, conventional accuracy metrics have limited meaning when there is no external ground truth: a prediction is measured against the student’s own prior annotation, itself a subjective judgment. Mean average precision (mAP@0.5) measures how accurately a model locates and classifies objects; a score of 1.0 would indicate perfect detection, while our models ranged from 0.009 to 0.384. The variation in model performance likely reflects differences in visual consistency within datasets, annotation strategies, and the inherent difficulty of learning subjective concepts. Eva’s meme dataset, with its consistent visual format, achieved the highest performance; Aisha’s abstract religious concepts proved hardest for the model to learn, potentially due to her dataset being filled with extreme detailed imagery.

Crucially, this variation did not determine the reflective value of the process. Students whose models

← *Figure 5.3*

Top image: A student analyzing their printed out datasets to figure out what words to use for labeling. Bottom images: students engaging with the annotation task



Figure 5.4
Object Portrait preparation on Day 2. A student selected and arranged objects from their personal practice, anticipating how their trained models would interpret the compositions.



Student	Topic	Images	mAP@0.5	Epochs
Eva	Memes	92	0.384	122
Jassir	Morphed faces	87	0.209	209
Agustina	Birthday celebrations	88	0.180	141
Borisz	Rubens paintings	110	0.149	128
Sasha	Festival stages	102	0.065	159
Yann	Abandoned buildings	87	0.063	124
Pepa	Children's cartoons	71	0.057	203
Carlo	Queer identity/drag	88	0.057	266
Rey	Wildlife photography	78	0.029	150
Yuky	Biomimicry fashion	56	0.029	120
Aisha	Religious events	85	0.009	125

performed poorly still reported meaningful insights, sometimes precisely because the model’s struggles revealed something about the nature of their labels. The technical outcomes served as material for reflection, as starting points for questions about why certain concepts translated into visual patterns more readily than others. See Appendix B.2

↑ *Table 5.3*
Training outcomes for the eleven workshop models, sorted by mAP@0.5. Image counts reflect the original annotated images.

for detailed training visualizations including label distributions and confusion matrices for each student. The final workshop day began with the students testing their Object Portraits with their personalized models, interpreting the results as the AI’s ‘reading’ of their work. Based on these insights, students were given the option to modify their portraits (*Task 3.1*). The day then transitioned into a collective reflection structured as a final exhibition (*Task 3.2*). In a group discussion, participants

presented their work to their peers, focusing on what the AI’s interpretation had revealed—especially surprising or previously unconsidered aspects of their practice (*Task 3.3*). Recognizing the vulnerability in discussing overlooked elements, we encouraged them to present their work confidently, as if delivering a pitch. After presenting to their tutors (*Task 3.4*), the workshop concluded with a final, individual structured reflection. Students filled in ‘reflection cards’ (*See figure 5.5*) that prompted them to document their key learnings, shifts in perspective, and their evolving relationship with their AI model and AI in general (*Task 3.5*).

○ *Data Collection and Analysis*
In designing this workshop study, we anticipated large amounts of qualitative data and prepared for emergent findings, recognizing that insights often arise unpredictably in design research contexts (Gaver et al., 2022). Since we could not predict exactly when or how students would reflect, we identified key moments to explicitly prompt written or group reflections. These moments were captured using a voice recorder. We also maintained tutor notes throughout the three days to document our guidance and observations.

Data collected included the students’ Object Portraits, their image datasets, annotated quadrant tools, audio recorded reflections, written reflection cards, and tutor notes. We began by segmenting the collected data according to each stage of the workshop using a collaborative online whiteboard. We also

contributed our own reflections as tutors to contextualize the students’ processes. Transcribed reflections were mapped onto the corresponding workshop steps and linked with students’ images and labels. We employed thematic analysis approach (Peterson, 2017) to identify patterns across the collected data. Gijs and myself, who also served as workshop tutors, collaboratively mapped transcribed and written reflections onto a visual representation of the workshop stages, linking them to the corresponding steps, labels, and images. We adopted a collaborative approach inspired by reflexive thematic analysis (Braun et al., 2023). A third author, with a background in philosophy of technology, reviewed the initial thematic mapping to challenge our interpretations from a theoretical standpoint. Rather than applying a rigid coding scheme or calculating inter-rater reliability scores, we prioritized developing a rich interpretation of the data through iterative discussion until we reached a shared reading of the material. The resulting thematic structure is documented in Appendix B. Our analytic approach addressed two areas of inquiry: (1) how students reflected on their own creative practice and design thinking through this method, and (2) how engagement with AI systems facilitated insights, particularly across the processes of data collection, annotation, training, and interpretation. We paid close attention to moments where students’ understanding shifted, whether in their research focus, annotation strategies, or relationship to AI technology.

Question 1: When I started collecting the dataset I had a general idea of who Rubens is and what his paintings are like.

As I was seeing more and more I started to be more interested in the composition than the figures. I also was curious about the meaning behind each little symbolic detail.

I realized during the annotation that I really appreciate ~~the~~ the perfection of a technique, as Rubens is a master of painting. I also found the lighting important to me.

Back to Rubens, I was trying to notice, what my research question originally was related to, if I can tell apart Rubens and his pupils' works. I am not sure if I could.

Question 2: moments where you felt your position changed (labeling, portraying, interpreting)?

Some areas of my downloaded pictures I annotated as "commercial saturation", specifically within stage designs. At first this label was supposed to be repulsive, but in some cases I found that this commercial overload was intriguing and attractive.

It's also interesting to think about how the AI can make "mistakes" by associating a word with an area of the ^{completely unrelated} image, but those "misinterpretations" can also reflect a deeper meaning that goes beyond the associations we humans would usually make.

Figure 5.5

Two reflection cards documenting learnings and shifts in perspective. These questions aimed to capture individual insights at the workshop's conclusion.

Figure 5.6

Two portraits presented on screens in the exhibition space



5.5

Findings: Learning through and about AI

As Gijs and I analyzed the data from the workshop, it became clear that the slow, hands-on process of dataset curation, manual annotation, model training, and Object Portrait creation facilitated two distinct yet interconnected forms of learning. Some findings echoed Chapter 4, confirming these as features of the method itself; others revealed new dimensions specific to the educational context. First, students engaged in learning *through* the AI: reflection on their

own creative practices, with AI serving as a reflective medium—the workshop’s primary pedagogical goal. Second, students experienced learning *about* the AI: developing a deeper, practice-based understanding of AI systems and their subjectivities. This section presents these outcomes in turn. To ground these findings in the students’ lived experiences, we draw extensively on direct quotations from their reflections.

● 5.5.1 Learning through AI as a reflective medium

Learning *through* the AI captures the moments that were most exciting to me as a facilitator. These were the insights where the workshop structure became a vehicle for self-reflection on one’s own creative process, enabling students to discover new understandings, aesthetic judgments, and ways of seeing. Here, the AI served as a reflective medium rather than merely a technical tool, with learning outcomes extending beyond AI literacy itself.

○ Articulating subjective judgments through structured label selection

The quadrant tool and label selection process helped students articulate and understand their own complex judgments about their topics of interest by encouraging them to consider labels that represented potentially conflicting emotions and perspectives. This task revealed a significant shift in how labels are typically conceptualized, transitioning from static, object-based categories to temporal and fluid ones in a dynamic process.

Carlo’s experience exemplifies this process. Working with a dataset focused on queer identity expression and featuring photographs of people in queer and drag fashion, Carlo initially felt overwhelmed by their “*complex and overwhelming admiration*” for the individuals in their dataset. Through multiple iterations with the quadrant tool,

Carlo arrived at two contrasting labels: ‘*enticing*’ and ‘*cloying*.’ The breakthrough came when Carlo began ‘personifying’ the labels: “*I’ve sort of thought like if these [labels] were people, how would I describe them? For some reason that really helped me. So for example for [the label] ‘enticing,’ I say that it’s someone who’s intriguing and calm. And then, cloying was when it’s like a sweet, you know, like too much. It’s so much that it overpowers your senses and you’re just drowning in the amazingness.*” This process helped Carlo recognize that while their admiration was consistently positive, it contained an important distinction between measured appreciation and overwhelming intensity.

Sasha’s exploration of festival ‘*vibes*’ demonstrates how the structured approach revealed nuanced emotional landscapes. Working with festival stage imagery, Sasha used the label selection process to map different feelings he encountered at festivals, ranging from ‘*otherworldly*’ magical atmospheres to moments of “*commercial saturation*” that lacked depth. As he explained: “*So ‘lifeless’ and ‘commercial saturation’ are two kinds of themes in the festival stages that I thought didn’t really draw me forward and were not really appealing to me. And then ‘otherworldly’ and ‘convivial’ are two words that I thought were warmer and kind of intriguing, kind of as a reflection of the music and the artist.*” By organizing these feelings on contrasting axes, Sasha created a structured understanding of his aesthetic preferences that he had not previously articulated. Sasha’s iterations over the labels, followed by his reflections on the selected labels are shown in Figure 5.8 (a) and (b) respectively, capturing how he tried to search for the meaning of these labels in the quadrant tool.

The process of articulating complex judgments into single- or two-word terms did not simplify students’ positions; instead, it distilled them. From my perspective, this was a key pedagogical success: Positioning labels on the quadrant revealed connections, whether aligned or opposing, and brought latent feelings and judgments to consciousness, enabling students to develop more nuanced self-awareness about their aesthetic and emotional responses.

○ Anticipating AI analysis and rethinking design decisions

We noticed that the process of preparing data for the AI model influenced students’ approach to creating their object portraits. The anticipation of their work being ‘interpreted’ by the AI triggered a form of self-critique, as students began to view the creative process through a lens shaped by the annotation experience.

Sasha’s reflection encapsulates how participating in data labeling reshaped his creative orientation: “*It’s so funny how you’re not just training an algorithm. You’re also training your own gaze.*”



Figure 5.7
Images showing the exhibition of the portraits and students presenting their work to their tutors and peers.

(a) Iteration over labels



(b) Reflection on labels

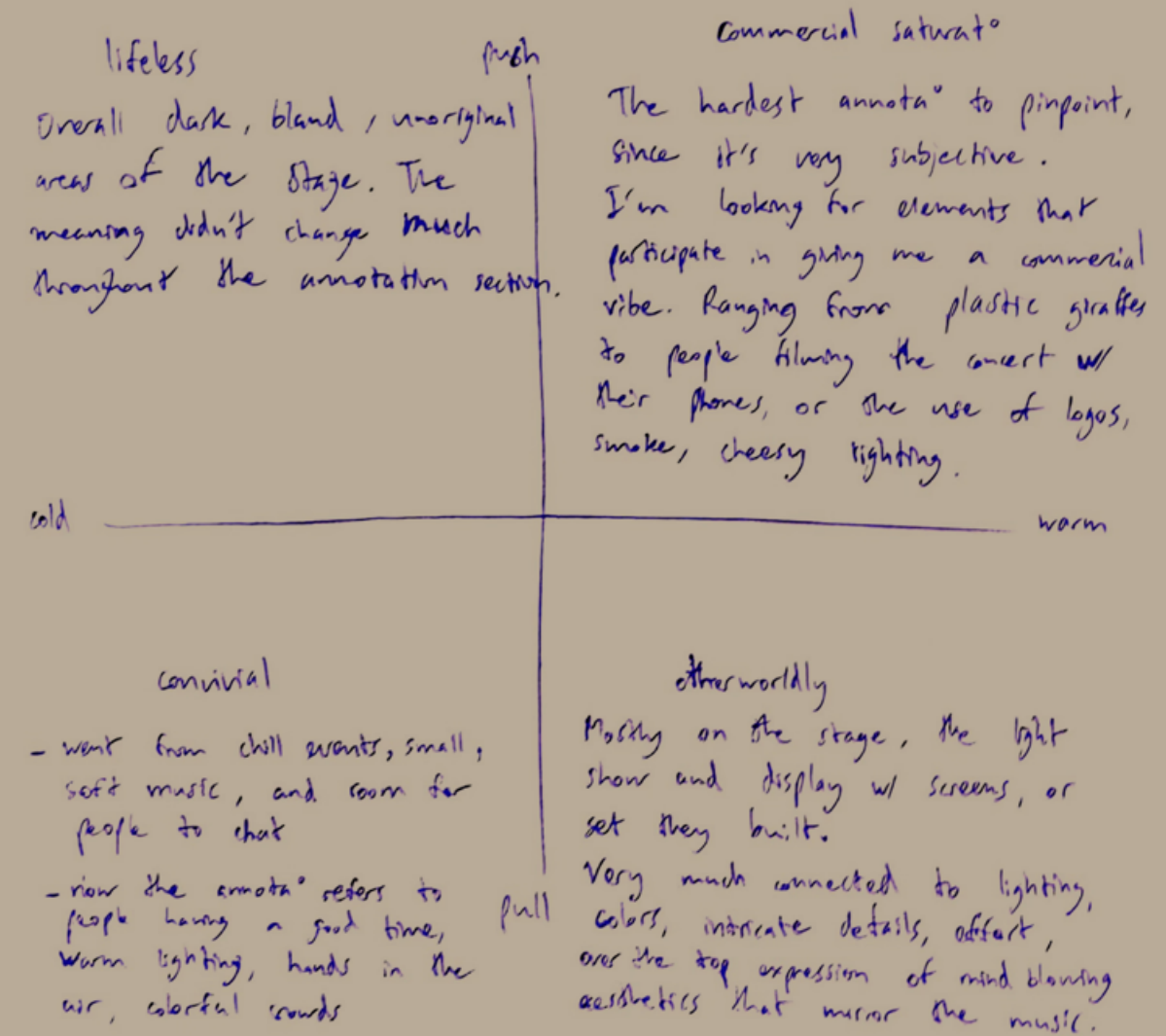


Figure 5.8
Labels chosen by Sasha relating to his topic 'Vibes at a festival', with photographs of festival and concert stages as his dataset. The figure shows (a) his iteration over the labels using his quadrant tool, and (b) his reflection on the final choice of labels.

You might start to see whatever you're looking for in the images also in the real world." This insight reveals how the labeling process extends beyond technical annotation to become a form of perceptual training for the students themselves. When creating his object portrait, Sasha deliberately incorporated elements he had labeled *'otherworldly,'* using altered blue lighting, scaled-up toys and a "Mark Zuckerberg in the style of a cyborg" figure, anticipating how his trained model might interpret these surreal elements. Similarly, Yann noticed his use of what he called *'DIY aesthetics'* when laying out objects from his previous design work, and started to wonder about how these might embody *'rebellion,'* one of his chosen labels. The knowledge that their works would be 'seen' by their own AI model seemed to instigate students to anticipate interpretations, which guided how they would see and create their object portraits.

○ Symbolic and metaphorical meaning discovery through AI interpretation

One of the most surprising findings for me was how students reported that the model predictions, regardless of whether they were anticipated or unexpected, served as starting points for speculation about the artistic and symbolic meaning of their work. Even though they had prepared the training data themselves, they still perceived the AI model as offering new, unexpected perspectives on their creative work. These ideas were not based on the technical accuracy of the model, but rather abstracted insights about the meaning of their work.

For example, Aisha's reflection on her object portrait analysis demonstrates this symbolic discovery process. When her carefully considered bluish white lantern was not detected but a candle received all four of her labels unexpectedly, this made her reflect on its metaphorical meaning: *"in a metaphorical way, I liked this part. When I light up the candle, it's turning into 'dominant', and when I blow it out it disappears. Cultural or religious activities see fear and hope at the same time. So here, I can really see what I want to say in this work."*

Similarly, Pepa's initial disappointment with her model's interpretation led to broader insights about cartoon aesthetics and age-appropriate media. Her model detected only a small area in her portrait using just one label (*'comforting'*), despite her having labeled most faces in her dataset

→ Figure 5.9
Samples from annotated training data and object portraits by Aisha and Pepa

→ Figure 5.10
Samples from annotated training data and object portraits by Yuki and Borisz. (a) Sample images from Borisz's annotated training data; (b) Sample images from Yuki's annotated training data; (c) Borisz's object portrait interpreted by his AI model; (d) Yuki's object portrait interpreted by her AI model.

5.9



5.10

as *'frightening.'* "I think this shows that the type of media we would call age-appropriate for children is very far from real life." The AI's analysis became a catalyst for reflecting on the broader cultural implications of her chosen visual domain.

The deliberate evaluation of the model's interpretations by the students—whether anticipated or unanticipated, and whether fully understood based on their annotator behavior or not—seemed to create a space for speculating on multiple levels of the creative process they had undergone. This ranged from attempting to grasp their annotator behavior through the model's detections, to engaging in higher-level speculation on the potential meanings of their work, where the model's detections served as inspiration or departure points for further reflection. Thus by training the model and anticipating on the model analyzing their own work, the students had also taught themselves new ways of looking at their work.

● 5.5.2

Learning about AI through hands-on practice

Learning about the AI encompasses findings where we saw students gain explicit knowledge about AI systems themselves: understanding how machine learning models function, discovering the challenges of data annotation, and attaining a form of AI literacy through hands-on engagement.

○ *Discovering machine learning principles through model performance analysis*

Some of the most substantial learning moments about AI in the workshop occurred when students gained insights into fundamental machine learning principles by analyzing instances where their models performed in unexpected ways, leading to direct understanding of how these systems process and interpret information. These moments of surprise created opportunities to grasp core concepts about machine learning that typically remain abstract in theoretical explanations.

Eva's observation of her meme dataset analysis exemplifies this learning. She noted that her model recognized *"a whole text paragraph, even though I didn't always annotate the full piece of text. It was sometimes just a singular word..."* This unexpected behavior taught her about how AI systems can generalize beyond exact training examples, learning to detect broader visual patterns rather than only the specific elements that were annotated. Her reflection that the AI *"did put the labels in almost the right place every single time"* despite not being able to *"read the text because it's only analyzing*

the pixels" revealed core principles about how computer vision systems process information—focusing on visual patterns rather than semantic content.

Similarly, students learned about pattern recognition principles when their models detected similarities they hadn't explicitly taught. Borisz's realization that his model labeled the same visual element as both *'majestic'* and *'repulsive'* led him to understand: *"I labeled animal legs majestic, and satyr legs repulsive, but they are actually the same if you just look at the legs. Now that I think about it, I understand why."* This moment illuminated how machine learning systems focus on visual features rather than contextual meaning—a fundamental principle about AI processing that emerged through hands-on analysis.

○ *Discovering the annotation consistency challenge through practice*

An interesting dynamic to observe was the central learning outcome concerning the fundamental tension between subjective human judgment and the requirements of machine learning training. Students discovered firsthand that while personal judgment is inherently fluid, consistent annotation remains crucial for reliable model performance. This challenge manifested as students found their understanding of their chosen labels shifting over the course of annotating a large amount of images. The repetitive and enduring nature of the labeling task led to a gradual evolution in their interpretations, creating a practical dilemma that illuminated core principles of AI training. As Yann reflected: *"for me all the words change their meaning a bit"* during the annotation process. Students developed three distinct strategies to navigate this tension:

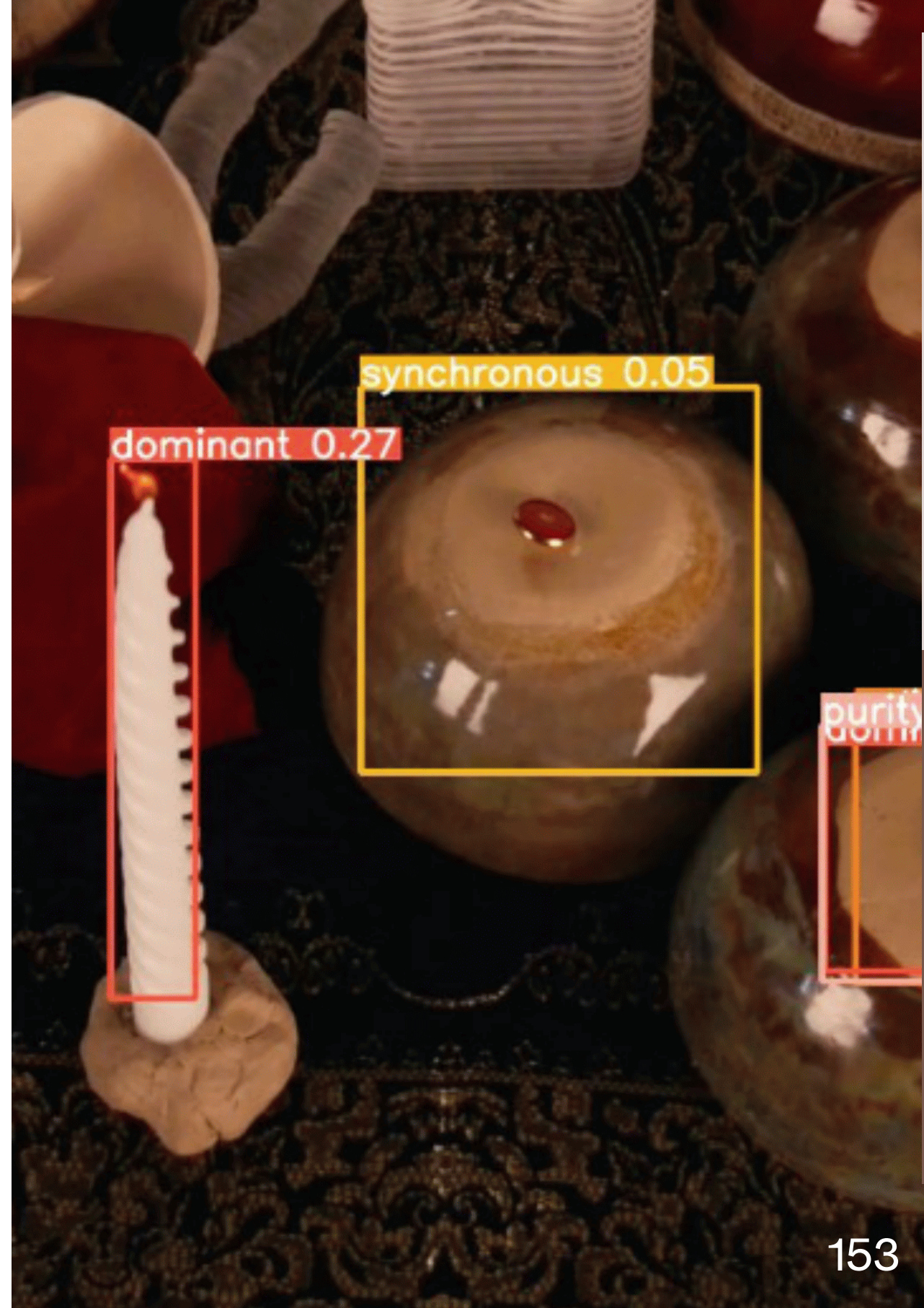
↳ *Revision strategy:* Some students chose to adapt by revising their labels to better capture their evolving judgments. Yann, working with a dataset of abandoned buildings, experienced a significant shift in his perception of labels during the annotation process. His understanding of *'pretentious'* evolved: *"I would now understand [it] as 'hiding' and not really the 'pretentious' part. Hiding from something that it is not,"* he explained, suggesting that modern buildings began to feel like they were hiding from a certain "reality" that abandoned buildings conveyed. Similarly, his label "inconsistent," which initially had negative connotations, *"became way more positive"* over time as he began to perceive those elements in a positive light, almost discerning what he felt was the original artist's intended meaning. In response to these shifts, Yann chose to revise his entire labeling scheme—replacing *'restricting'* with *'rebellious,'* *'reductive'* with *'inconsistent,'* and *'lying'*

with ‘devotion’—then re-annotated his entire dataset. This approach prioritized authentic representation of his current understanding over maintaining consistency with his initial judgments.

↳ *Acceptance strategy*: Others continued labeling with their original labels while accepting the inevitable shift in judgment, adopting a more flexible approach to the annotation task. Sasha, who worked with festival stage imagery, acknowledged: *“I made those choices at one point at the beginning and then I realized that sometimes it didn’t apply. I was like, ‘okay, here’s this element so I’m going to annotate it’, and then when I look back at the whole image, I’m like, ‘actually, it should be something different, something else’.”* Rather than revising his labels or forcing strict consistency, Sasha adapted his interpretation approach while keeping the original label structure. He explained how his understanding of ‘convivial’ evolved: *“I had the label ‘convivial,’ and I think in the beginning I interpreted that more in the actions, like hugging or sitting together, which was really hard to identify. Then I switched to things like light sources, for example, if it’s soft light compared to the harsh colors in all animations.”* This strategy represented an acceptance of the inherent subjectivity in the annotation process, allowing for interpretive flexibility while maintaining the original conceptual framework.

↳ *Consistency strategy*: A third approach involved maintaining strict consistency with model performance explicitly in mind, prioritizing technical reliability over evolving personal interpretations. Borisz, working with Rubens paintings, developed a systematic approach that helped him maintain consistency even as his internal understanding shifted. He created what he described as a specific labeling code, focusing on elements that were easy to categorize: *“Yeah, first I was labeling hands as majestic. But then I realized they are more sublime to me, especially since they always had some symbolic meaning in the paintings. So I went back and relabeled them. I think that helped with the consistency.”* Once he made these initial adjustments, Borisz remained deliberately consistent: *“I was super consistent in labeling... even when I started changing things in my head, like the meaning changed... I let go of some things, but I didn’t really label them differently.”* This approach reflected a strategic decision to prioritize the model’s ability to learn clear patterns over authentic representation of evolving personal judgment, demonstrating an understanding that consistent input would lead to more predictable AI behavior.

Through this practical challenge, we watched students gain direct insight into one of the fundamental tensions in machine learning: the need to transform fluid human judgment into consistent training data. Regardless of which strategy they adopted, it became clear that



students discovered that the annotation process became a form of meta-reflection where ‘teaching’ the AI about their chosen topics simultaneously showed the evolving nature of their own perceptions and judgments. This dual learning process shows how Reflective AI engagement can make visible the subjective choices typically hidden within seemingly objective technical tasks, transforming annotation from a means to an end into a reflective practice that generates insights about both AI systems and one’s own ways of seeing.

5.6 Discussion

Chapter 4 showed that reflection emerged from the training process itself, with slowness creating conditions for this to occur. The workshop exhibited similar dynamics in the educational context. Two mechanisms proved central: slowing down engagement with AI, which shifted focus from output generation to recursive inquiry; and disentangling the technical pipeline into distinct phases, which created conditions for agency and critical literacy. I discuss each below.

- *Shifting focus from AI output to reflective practice*

The structured slowness of the workshop created multiple entry points for reflection. The label selection process required students to distill complex emotional responses into distinct terms, surfacing aesthetic judgments in ways the process of articulation made visible. Carlo’s process of ‘personifying’ labels to distinguish ‘enticing’ from ‘cloying,’ or Sasha’s mapping of festival atmospheres onto axes of ‘commercial saturation’

versus ‘otherworldly,’ show how this structured process turned vague feelings into specific vocabulary.

The workshop’s structure played a crucial role in facilitating this, particularly through the recursive relationship between students and their developing models. When students taught an AI about their fascinations, they simultaneously learned how they perceived their own work. This exemplifies what Ihde & Malafouris (2019) term the ‘recursive effect’: the way that shaping a technology simultaneously shapes the one who shapes it.

Crucially, the outcome of the workshop, the prediction of labels on student-created Object Portraits, was not AI-generated content to be consumed but rather served as a vehicle for students to consider the effect of their annotation choices and the unpredictability of the model’s interpretations. This resonates with Papert’s (1991) theory of constructionism, which holds that learning happens most effectively when people

are actively engaged in constructing something shareable, whether a sand castle or a theory. In the workshop, students constructed datasets, devised labeling systems, and made Object Portraits; through this making, they developed understanding of both their own aesthetic judgments and how AI systems learn. The AI functioned as a medium for this constructionist learning: a material that students shaped and, in shaping, came to examine their own ways of seeing.

Ultimately, the workshop’s success is defined by the depth of the students’ reflective engagement with their own creative process, not by the quality of the AI’s output. This was evident when students found metaphorical meaning even in the AI’s failures or unexpected interpretations. As Aisha reflected on her model’s analysis of a candle, the output became a catalyst to “*really see what I want to say in this work*,” transforming the AI from an analysis tool into a reflective medium.

- *Developing agency and critical literacy through material engagement*

The workshop’s multi-phase structure, where students engaged with data collection, annotation, training, and inference as distinct processes, supported a second form of learning: developing agency over the technology itself. By agency I mean the capacity to act on and through tools reflectively (Schön, 1983; Sengers et al., 2005), rather than accepting them as given. The distinct annotation strategies that emerged, whether Yann’s revision approach or Borisz’s consistency

strategy, show students actively negotiating how to work with the system rather than simply following instructions. They discovered a core tension in machine learning, between fluid human judgment and the system’s need for consistent data, through practice rather than theory.

The temporal structure of the workshop was crucial to this. Students could not iterate instantly; they had to commit to annotation choices before training, then wait, then interpret outputs the next day. This delay meant that decisions had consequences they could not immediately undo. The ‘point of no return’ before model training forced students to take their annotation choices seriously, while the gap between annotation and output created space to anticipate how the model might respond. Through this structure, students came to understand AI as a process with distinct phases rather than an immediate response system.

Through disentanglement of these phases, the components of the step by step technical pipeline of an AI technology actually dissolved into a rich and temporally intricate design space. In that space, our students navigated slowly the different facets of temporal interaction between components (anticipation, reorientation, speculation). Regardless of which annotation strategy they adopted, students discovered that the process functioned as more than a technical requirement—it became a form of meta-reflection where ‘teaching’ the AI about their chosen topics simultaneously revealed

the evolving nature of their own perceptions and judgments.

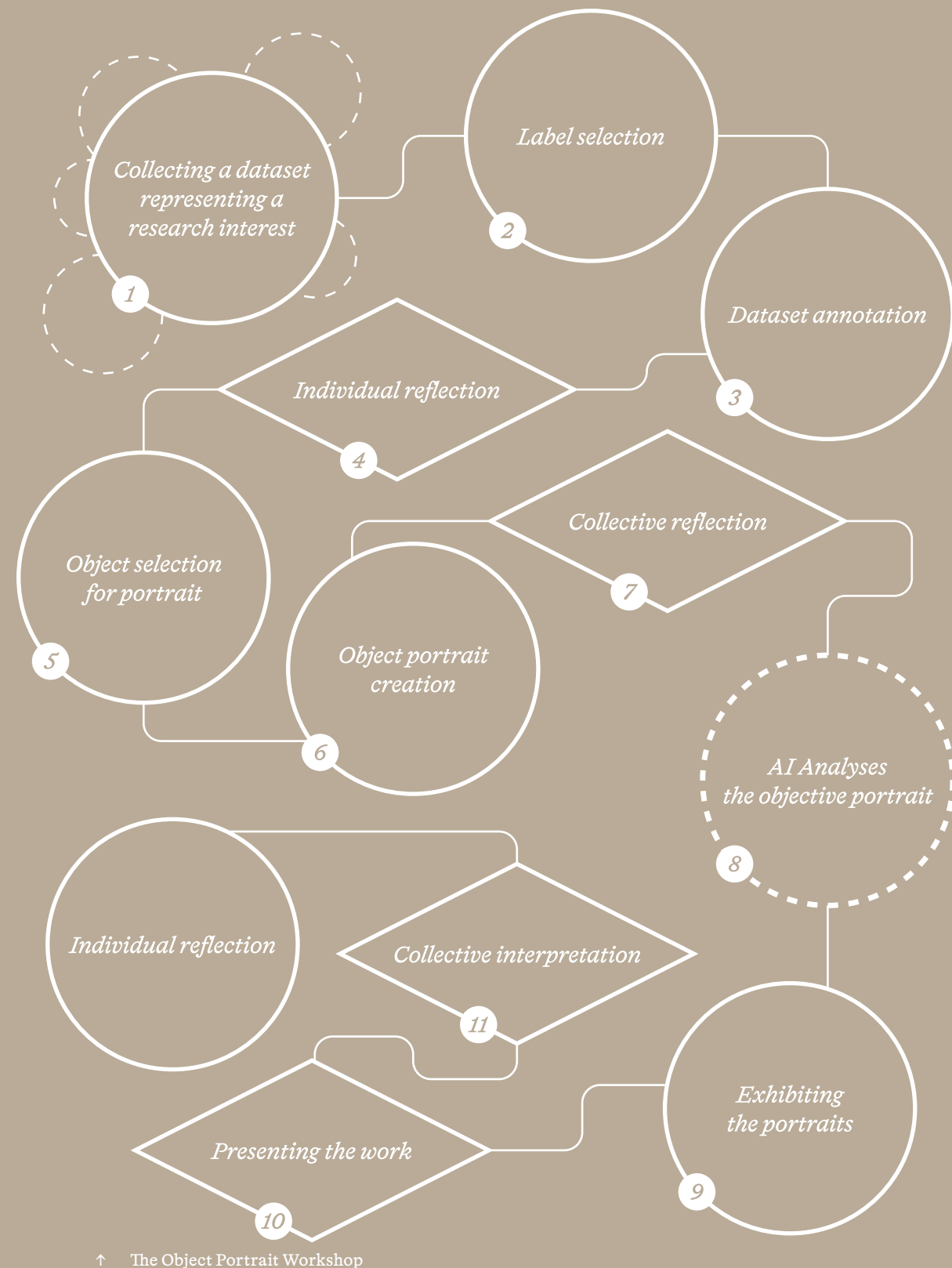
Our approach resonates with recent work on artistic practices with AI, particularly the notion of ‘model crafting’ (Steinbrück & Stankowski, 2023; Caramiaux & Fdili Alaoui, 2022; Abuzurairq & Pasquier, 2024), where artists shape AI systems through deliberate engagement with training data and model behavior. However, these accounts typically describe practitioners with established creative identities who can draw on existing intuitions when navigating AI’s

complexity. Design students, still developing confidence in their creative judgment, require more scaffolding. The annotation strategies that emerged in the workshop represented developing philosophies about maintaining creative integrity when working with probabilistic systems. This is preparation for professional contexts where students will face similar negotiations without facilitation. This, we believed, is essential preparation for professional contexts where they will face such negotiations without support.

5.7 Conclusion

The workshop demonstrated that the Objective Portrait method can function as a pedagogical approach. By slowing down engagement with AI and disentangling its technical pipeline into distinct phases, students developed both reflection on their emerging creative practices and critical understanding of AI systems. As a study, it has limitations. The three-day structure provided scaffolding that may not transfer to less facilitated contexts. The small group of eleven students from a single institution limits broader claims, and we do not know whether the reflective insights persisted beyond the workshop itself. Our data collection methods – reflection cards, group discussions, tutor notes – were designed to capture moments of

articulation. This may have biased our findings toward students who successfully translated the process into insight, while quieter struggles or disengagement would have been less visible in these formats. The focus on object detection also leaves open how this approach might work with generative text-to-image systems like Stable Diffusion or Midjourney, which have since become more widely adopted in creative practice. These systems foreground output in ways that object detection does not, which may create different conditions for reflection. Whether the same principles of slowness and disentanglement can redirect attention from generation to process remains, at this point, an open question.



6. **Why do I keep on coming back to this chair?**

After establishing the Objective Portrait method through collaborations with designers and students, I felt a growing need to turn the lens toward myself and my own practice as a maker. Until this point in my PhD trajectory, I had been working with other designers, with design students, connecting my ideas to design theory and working to establish myself as a design

researcher in the field. While this felt educational, I was missing an important aspect of the work I had been developing since 2017: my own making practice. At this stage, I could say I did research in AI and design, but what was my practice actually? A designer? A researcher? An artist? Maybe a writer? I did not exactly know. This uncertainty kept me uncomfortable, creating an inner need to engage with something more personal within myself. A return to the starting point that had initiated my very first AI-related project, *Still Life*, where I began exploring what it means to create with these technologies. For the final two projects, I pulled forward the maker-designer in me to reflect specifically on that and on my own

identity as that maker-designer. The following two chapters will therefore have a personal tone, walking through my AI practices as a form of inquiry into my own creative history.

6.1

The backward gaze

For eight years, I have been haunted by a chair. It was a foam chair I made in design school, a large, womblike object built from rest materials on a therapeutic afternoon in my garden. That chair saved my artistic career, proving to my tutors at the prestigious design school I was in (and potentially to myself) that I had a place there. And then, it was gone. After graduation, I dismantled it, and it existed only as a collection of photographs on an old hard drive. Yet the chair persisted in my memory. I kept returning to it in my mind, unable to fully understand why this one specific object, out of everything I had ever made, held such a persistent grip on me. This chapter is my attempt to answer that question. And an image generation model, I came to realize, offered an unexpected instrument for this inquiry.

AI systems are at their core backward looking. Their outputs are statistical recombinations of training

data, and training data is inherently historical - it stems from what has already been created, collected, and categorized. This backward-looking quality is often framed as a limitation, particularly in how it risks amplifying historical biases captured in large datasets (Vallor, 2024) to make predictions in the present.

Yet I find this temporal orientation fascinating. Training data functions as a kind of archive: organized collections of images, text, or other media, all categorized and given meaning through that categorization. These are materials made by someone, selected by someone, structured for a specific purpose. In the context of large generative models, this purpose is typically broad and utilitarian: vast quantities of internet data are collected and annotated so users have maximum flexibility to generate whatever they want. At this scale, the historical and temporal dimensions

of the archive become secondary, even invisible.

What if I took this backward-looking quality as the starting point for answering the question about the chair? What if I used AI to help me excavate my own memory, training a model on a personal archive to understand why this chair refused to let go?

In previous chapters, I explored how object detection models could analyze existing images and offer new perspectives, helping reflect on specific fascinations and research interests in the present. For this inquiry, I turned to generative image models: models that, given an inventory of training material – given a history, an archive – could generate new images based on a prompt. Comparing the generative image to memory, I somehow see parallels. Human memory is not a perfect recording but inherently reconstructive (Schacter & Addis, 2007; Maguire et al., 2010). We piece together fragments of past experiences, blend them with our current emotional state, and create a coherent narrative in the present. I see generative AI models operate on a strangely similar principle. They learn patterns from a historical dataset and then, when prompted, reconstruct those patterns into new visual form. This slight similarity between human remembering and algorithmic generation made text-to-image models feel like the right instrument for this inquiry. I began to think of this approach as AI-as-memory: using a generative model trained on a personal archive as an instrument for remembering.

This raises a question: *How can I use a generative tool not to make something new, but to understand something old?* Reconceptualizing what generation means in this context offers a direction: I am not generating images to merely consume, but rather, I am generating reconstructions, visual hypotheses about my past. The understanding comes out of the reconstruction itself, out of prompting the model and interpreting what returns.

Methodologically, this inquiry sits between autobiographical design and autoethnography. Like autobiographical design (Neustaedter & Sengers, 2012), I built and used a system myself. But the goal was not to learn about the system's potential; it was to learn about my own past. Autoethnography, which has found increasing use in HCI as a way to investigate subjective dimensions of technology use (Rapp, 2018), offered a framework for this orientation: a first-person research method where you systematically analyze your own experience to understand something larger about culture, practice, or human behavior (Ellis et al., 2011). What I add to this is a technological instrument: the AI-as-memory model became a way to structure the autoethnographic process, giving me something to respond to rather than relying only on introspection.

I spent one month teaching an AI to remember my foam chair. I asked: *What does this personal obsession reveal about creative practice and memory? Can AI help me understand ways of thinking from the past without making something new? Can it help*

me understand my past intentions? If I train an AI on my own archive, will it help me understand something of myself in the present?

To conduct this inquiry, I needed a model that could learn from my own human-sized archive. Standard generative models carry the vast, impersonal archive of the internet as their memory. To conduct this personal inquiry, I needed a model that could learn from my own 'human-sized' archive. The technology I employed is

Stable Diffusion version 1.5 (Rombach et al., 2022), fine-tuned using Low-Rank Adaptation (LoRA) (Liu et al., 2024). LoRA is a computationally efficient method for customizing large pre-trained models by introducing a small set of trainable parameters that adapt a larger base model to patterns found in a personal dataset. By fine-tuning Stable Diffusion on my archive of the foam chair, I wanted to explore its potential to become a personalized memory instrument.

6.2

The story of the chair

From this point on, I will detail the situated process of this inquiry as it emerged. I begin with a prologue to ground the inquiry in the personal history of the object, the foam chair. Then, I will move through the three primary stages of the practice: the archival work of data collection, the textual translation of annotation, and the therapeutic dialogue of prompting. Each stage represents a site where knowledge was discovered through a reflective engagement with the material of my archive and the generative capacities of the AI model.

● *A Prologue: Why this chair?*

Back when I was in design school, I wasn't the best maker. It was a very well-known design school, and I came from a science background, which was quite unusual compared to my peers who had developed incredible

making skills in the trajectories of their bachelors. Near the end of my first year, I failed a trimester after presenting what my tutors deemed an incomprehensible design project. In such a competitive school, failing was devastating. It meant potentially having to leave and let go of my dream of being part of this exclusive educational program. I had to find a way to convince my tutors during my re-exam that I actually had what it took to be a designer. I realized that my scientist knowledge was not going to win them over. I needed to create something they could understand visually rather than cognitively. Something big and impossible to ignore.

I decided to make a chair, thinking it was at least an object people could easily understand in design school, even though I personally found chair design superficial and boring. It felt



I keep on coming back to this chair.

like meeting my tutors halfway. My previously failed project had been about the sense of touch, another topic that I actually found somewhat superficial. However, I decided to continue with this theme while developing the chair concept. I envisioned creating an artifact that would help people get back in touch with themselves, something you could sit in that would touch you all over. Something soft and comfortable.

I built the chair on a sunny spring day, using all the pieces of leftover foam I could find, from discarded couches, old chairs, and my peers' leftover design projects. I spent hours in the garden, adding more and more foam to a base structure until the chair became like a creature of its own, a foam collage that grew into something like a cave or a womb that I could sit in. It was a therapeutic day where I thought about nothing else but making

the chair bigger and bigger. When I presented the chair during the re-exam, the tutors were impressed and I passed with flying colors. Never again were my creativity or making skills questioned. The chair became my breakthrough, my Trojan horse into design school.

Eight years later, the chair disappeared. After graduating from design school, I had to dismantle and discard it, there was simply no space for it anymore in my house. It existed only in photographs stored away on old hard drives, buried deep in my digital archive that contained more old projects that after my time at the academy, never really saw the light of day. Yet, I kept returning to this chair. Sometimes I found myself opening that folder of photographs, and the chair crosses my mind when I reminisced about my time in design school, a period of ultimate creative freedom when life was all about making things.

● Step 1: Data Collection of my Artistic Past

The first stage of this inquiry was to revisit that specific archive, that would serve as the foundation for my AI-as-Memory. This process began not with the chair itself, but with a broader exploration of my artistic history. I collected images from 13 different projects created during that time in design school, resulting in 13 separate datasets. The aim was to create a comprehensive portrait of my practice at that time. These datasets captured specific art objects and sculptures, depicted from various photographic angles and in several contexts of use, including images showing people interacting with the objects and the objects arranged in different configurations.

Typically, it is challenging to fully grasp the extent of data used in machine learning systems due to the sheer volume combined with the digital screen-based interface through which we view these images. To overcome this challenge of fully comprehending the scale and significance of my collected data, I chose to physicalize the data

by printing the images and scatter them around my desk. This allowed



6.1

← Figure 6.1
The creation of the chair in my garden

me to view them all at once and rearrange them freely, offering a more direct, tangible interaction with the dataset. This process made it significantly easier to identify visual patterns, categorize images into subgroups, and organize them based on specific criteria such as photographic angles. By organizing the images in this way, I was able to observe visual consistencies and connections across different projects, providing new insights into relationships within the dataset (see Figure 6.2).

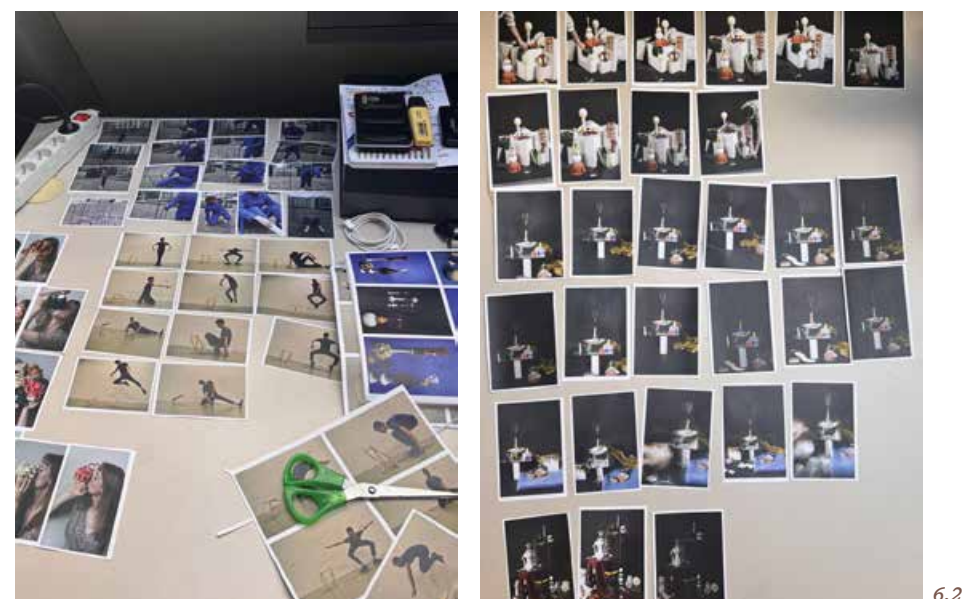
After examining the full archive, the focus of the inquiry naturally narrowed to the foam chair. More than any other project, it was the one that still held a sense of unresolved meaning and personal attachment. The dataset I prepared for the AI consisted of exactly 50 images of the chair, captured in several compositions and from different camera angles, including close-ups and side views. This collection became the foundation for exploring my attachment to the work and for reflecting on what this particular project represented in my practice.

● *Step 2: Creating a reflective annotation strategy*

With the inquiry now focused on the chair, the next step was to build a rich, reflective context around the dataset. I began by systematically revisiting its history, thinking about the concept, the making process, and my own mindset at the time. To structure these thoughts, I developed a practice of creating ‘reflective metadata’ as a way to structure the elusiveness of these memories. This resembled the quadrant tool work I performed in Chapter 4. I recorded straightforward details like the project’s title and materials, but more importantly, I compiled keywords categorized into two groups. Objective terms described observable features: foam, collage, large, white. In contrast, subjective terms captured the emotions and associations the chair still evoked: womb-like, monster, safety, awkward, comforting. This act of creating metadata served as a bridge between the latent memories of the project and the explicit, structured information the AI would need, allowing me to gather my thoughts and prepare them for translation into a machine-readable form.

Building on this foundation of reflective metadata, the next challenge was one of translation: how to distill these memories and insights into written annotations the AI model could understand. The technical requirement was simple: each image file must be paired with a descriptive text file. But my approach asked for something else. It required an annotation process that went beyond a simple visual inventory to capture what the visual was attached to: the emotions, experiences, and artistic intentions embedded in the work.

This exercise surfaced a tension in artistic practice. As an artist,



6.2



6.2

Background

PROJECT: SOFT CHAIR

DATA SAMPLES - EXPLORING ANNOTATION TAGS

- SAMPLE 1: FULL VIEW
- 1. A SOFT WOMB CHAIR, FOAM, MONSTER, FRONT, SCUMPTIOUS TENTACLES, COLLAGE
 - 2. SOFT WOMB CHAIR, A FOAM COLLAGE MONSTER-LIKE, SCUMPTIOUS VIBE, TENTACLES ON TOP, FRONT VIEW



- SAMPLE 2: FULL VIEW + PERSON
- 1. A SOFT WOMB CHAIR, FOAM, A WOMAN SITTING, FRONT, ~~RECENT~~ HIDING.
 - 2. A SOFT WOMB CHAIR, A FOAM COLLAGE, VERA SITTING, TENTACLES ON TOP



ANNOTATION TEMPLATE:

[CONCEPT SIMILAR IN ALL PHOTOS] + [MATERIAL]

Figure 6.3

Data annotation preparation

SAMPLE 3: CLOSE UP CHAIR

- A SOFT WOMB CHAIR, FOAM, FLESHY, SCUMPTIOUS, DETAILS, ZOOMED VIEW
- A SOFT WOMB CHAIR MADE OUT OF FOAM WITH A FLESHY LOOK, A SCUMPTIOUS VIBES, A DETAILED ZOOMED-IN VIEW.



SAMPLE 4: CLOSE-UP + PERSON

- SOFT WOMB CHAIR, WOMAN HIDING WITH CLOSED EYES, HOLDING THE TENTACLES, CLOSE-UP VIEW.



6.2



Why do I keep on coming back to this chair?

it felt unnatural to reduce a visual and material work into a simple, textual description. Much of its meaning resists language; there are countless ways to describe it, and none felt complete. The chair was never made to be described in words. Yet, for the AI to begin to understand it, I had to undertake this act of translation, crafting a line of text for each image in an attempt to capture not just what could be seen, but what could be felt.

And yet, I faced a constraint. While I aimed to describe the images in the most personal way possible, I also had to speak a language that would work for the model. This meant grounding my subjective interpretations in the objective aspects of the work: its colors, its materials, its form. A constant negotiation emerged between my artistic perspective and the model's technical requirements. The model needed consistent, recognizable concepts to learn effectively, while I needed to preserve the core, subjective essence of the work. Striking this balance became the central challenge of the annotation process, requiring me to hold onto subjective meaning while remaining legible to the model.

Through this iterative process of balancing personal meaning with algorithmic legibility, a pattern began to emerge, and I slowly developed an annotation system for myself. Each tag contained:

1. *A concept describing the constant feature present in each image*
2. *Material and/or method used in the creation of the chair*
3. *Photographic angle to capture the perspective of the image*
4. *An adjective conveying the overall emotional response the chair evoked*
5. *A notable visual detail described abstractly*
6. *The person interacting with the object (if applicable) and how they interacted*

The outcome of this annotation strategy was a dataset of 50 images, each paired with a corresponding annotation tag in the form of a text file that reflected both the structural characteristics of the image data and my personal perspective on how to describe and interpret these images. Despite the difficulty of translating visual art into textual description, or perhaps because of it, this process became a crucial step in preparing both myself and the eventual AI model for a more meaningful interaction with the work's essence.

● *Step 3: Prompting as Inquiry*

With the model now trained on the archive of my chair, the interaction shifted from teaching to questioning. The final stage of the process was to engage this personalized AI-as-Memory in a dialogue. This required me to reframe the very act of prompting. Instead of using prompts as commands to generate a specific thing, I began to use them as open-ended questions. Prompting became a form of therapeutic inquiry, guided by a question: "*What insights can I gain from how this model has learned to see my work?*"

I selected a few deeply personal prompts ('Vera', 'My Womb', 'A Foam Divorce') and generated multiple outputs for each, observing how the model's memory evolved over time. This prompting method was informed by the technical process of LoRA training, which unfolds over multiple epochs, or complete passes through the training data. Each epoch represents a deeper stage of learning: in early epochs (1-3), the model picks up on surface features like color and texture; in later epochs (7-10), it learns how these elements combine into a coherent style. This progressive learning created ten distinct 'versions' of my AI-as-memory, each with a different level of familiarity with my work. I wanted to track what the model would prioritize at different stages. Would it see the form first, and only later the feeling?

Initial experiments produced varied images across epochs, but I found them difficult to read as a progression. Without a way to track the same compositional starting point across training, I could not discern what the model was learning. To navigate this, I employed a fixed seed value for each prompt. In image generation, the seed determines the random starting point from which an image is composed. Using the same seed with the same prompt produces images with similar underlying structure. By fixing this value, I could observe how the model's evolving understanding, rather than random variation, shaped the resulting images. This allowed coherent visual narratives to emerge across epochs, transforming static outputs into a story I could read for clues about my own past.



→ Annotation

A soft womb armchair, made from various layers and rolls of foam and fabric, monster-like, scrumptious vibe, back view, tentacles on top



→ Annotation

A soft womb armchair, made from various layers and rolls of foam and fabric, monster-like, scrumptious vibe, detailed zoomed-in view, back side



→ Annotation

A soft womb armchair, made from various layers and rolls of foam, monster-like, scrumptious vibe, tentacles on top, side view, landscape



→ Annotation

A soft womb armchair, made from various layers and rolls of foam and fabric, monster-like, scrumptious vibe, back view, tentacles on top

Figure 6.4
Sample data and corresponding annotation tags

→ Annotation

vera with closed eyes lying in a soft womb chair made out of rolls of foam and fabric, a close up, holding foam tentacles, monster



→ Annotation

vera lying in a foam womb chair with legs crossed surrounded by layers of soft, rolled foam materials, top view, monster



→ Annotation

the back of a foam womb chair with feet sticking out, made out of layers of rolled foam materials, back side, monster



→ Annotation

A soft womb armchair, made from various layers and rolls of foam, monster-like, scrumptious vibe, tentacles on top, front view, portrait





↑ Prompt: Vera

↑ Prompt: My Womb

↑ Prompt: A Foam Divorce

The following sections show the prompting process, each reflecting a personal question I wanted to explore.

● *Prompt 1: 'Vera'*

Can I find myself back in the data?

In the initial image, a humanoid figure appeared, potentially representing me. However, as the model's training progressed through the epochs, this figure gradually faded from the sequence. The more the model learned, the more I seemed to disappear and transform into foam. The more the model learned about my chair, the more I seemed to dissolve into its material. It was a scary realization: in the model's training, my identity was not separate from the work, but was being consumed by it. The foam, the material essence of the project, was becoming more prominent than the person who created it.

● *Prompt 2: 'My Womb'*

Exploring safety, rebirth and new beginnings

The prompt 'My Womb' was an attempt to explore themes of safety, origin, and new beginnings that I had associated with the chair. However, the resulting visual narrative was incoherent. Even with a fixed seed, the sequence produced a strange and emotionally distant collection of images: babies, some unsettling and some mundane, lying in small foam beds. The images did not form into a story that resonated with me. Ultimately, the experiment suggested that the concept of 'womb,' which felt central to me, had not been successfully translated into something I could relate to.

● *Prompt 3: "A Foam Divorce"*

The breakthrough moment

The breakthrough came with a prompt of deep personal vulnerability: 'A Foam Divorce.' The phrase bridged an important experience from my past (my parents' divorce) with the rawness of my present (a recent break-up). Suddenly, where the previous prompts had been scattered or distant, a coherent and emotionally charged visual narrative emerged. The sequence began with the foam chair transforming into a fractured family, a scene that immediately evoked memories of my mother after her divorce. The narrative then shifted, showing the foam comforting a solitary figure (potentially myself?) forcing a reflection on learning to care for my own sadness. Finally, in the last epochs, a new, abstract object emerged. It unified the earlier family elements, but now stood independently: proud, present, and complete in itself.

● *Step 4: Discovering foam as consolation*

The visual narrative that emerged from the ‘Foam Divorce’ prompt was a turning point. It was here that the AI-as-Memory stopped being a mirror to my work and became a mirror to my life. The sequence, from the fractured family, to the solitary figure being comforted, to the final, independent form, was not a story the model reconstructed from the fragments I had given it, blended with the emotional weight of the prompt. It was in reading this reconstructed story that I started to realize something.

The chair, I saw, was never truly about design, or touch, or even the concept of a womb. Those were the intellectual ideas I had built around it. At its core, the chair was about the material, and more specifically what that material meant to me. The AI’s interpretation, by consistently returning to the texture and form of the foam in all of the prompts and its epochs, forced me to see what had been there all along. In the ‘Foam Divorce’ narrative, the foam was the constant agent: it fractured, comforted, and unified. Is foam, for me, a form of consolation?

Then, the past and present clicked into place. The physical act of making the chair eight years ago, a desperate, therapeutic process of building a safe space for myself during a challenging time, had been an act of creating comfort. And now, in the present, confronted with a different kind of emotional turmoil, the AI’s interpretation revealed that foam had become my personal metaphor for emotional grounding and self-care. It was not the chair itself I kept returning to, but potentially the feeling of consolation its material embodied. The AI, by learning from my past and responding to my present, had constructed a symbolic connection that had been latent in my practice for years.

This insight reframed the way I saw the process of training the model. Data collection had been an act of re-encountering that long-lost object. Data annotation had been the struggle to find the right words to describe that object. And finally, prompting the model became the therapeutic inquiry that allowed me to find myself through my own work. That, I now understand, is why I kept coming back to the chair.

● *Step 5: A new inquiry*

The realization that ‘Foam = Consolation’ was the core insight of this inquiry. It answered the question of the chair. But it also raised a new, more methodological question about the nature of this insight. Was it ‘real,’ or was I simply reading meaning into the technical

limitations of a probabilistic system?

A critical part of this process was to hold two perspectives at once. On one hand, I understood the technical reasons for the model’s behavior. The gradual disappearance of ‘Vera’ from the images was likely a result of insufficient representation of myself in the training dataset. When fine-tuning on a small personal archive, the model learns the dominant visual patterns present across images. Since ‘foam’ appeared consistently throughout all fifty training images while ‘Vera’ appeared in only a subset, the model learned to prioritize foam as the core visual concept. When prompted with ‘Vera,’ the model had little basis for reconstructing a distinct human figure, instead defaulting to the foam textures it had learned more thoroughly. The incoherent response to the ‘womb’ prompt could also be explained by a lack of relevant data in both my annotations and the base model’s training. This is the analytical, technical reading.

On the other hand, my story-reading approach was intentional and central to the reflective methodology. Knowing that my representation was technically fading did not invalidate the emotional resonance I found in watching myself dissolve into the material of my work. The technical explanation and the personal interpretation were allowed to coexist. One explained why a pattern emerged; the other allowed me to use that pattern as a prompt for reflection. This dual vision reflects the nature of working with AI as an artistic medium: you can understand its technical constraints while remaining open to finding unexpected meaning within them. The value of this method potentially lies in the AI models’ capacity to create structured opportunities for self-insight.

The understanding that insight could be found in the dance between technical uncertainty and personal meaning brought me to a new kind of prompting. The inquiry was no longer just about understanding the past. Having established foam as a personal symbol, I could now use it to explore the present. The equation was as following:

Foam = consolation

Vera = subject

Sadness = current state of mind



Prompt
Foam vera crying

180



Prompt
Foam vera hiding

181



Prompt
Foam vera yelling

182



Prompt
Foam vera sad

183



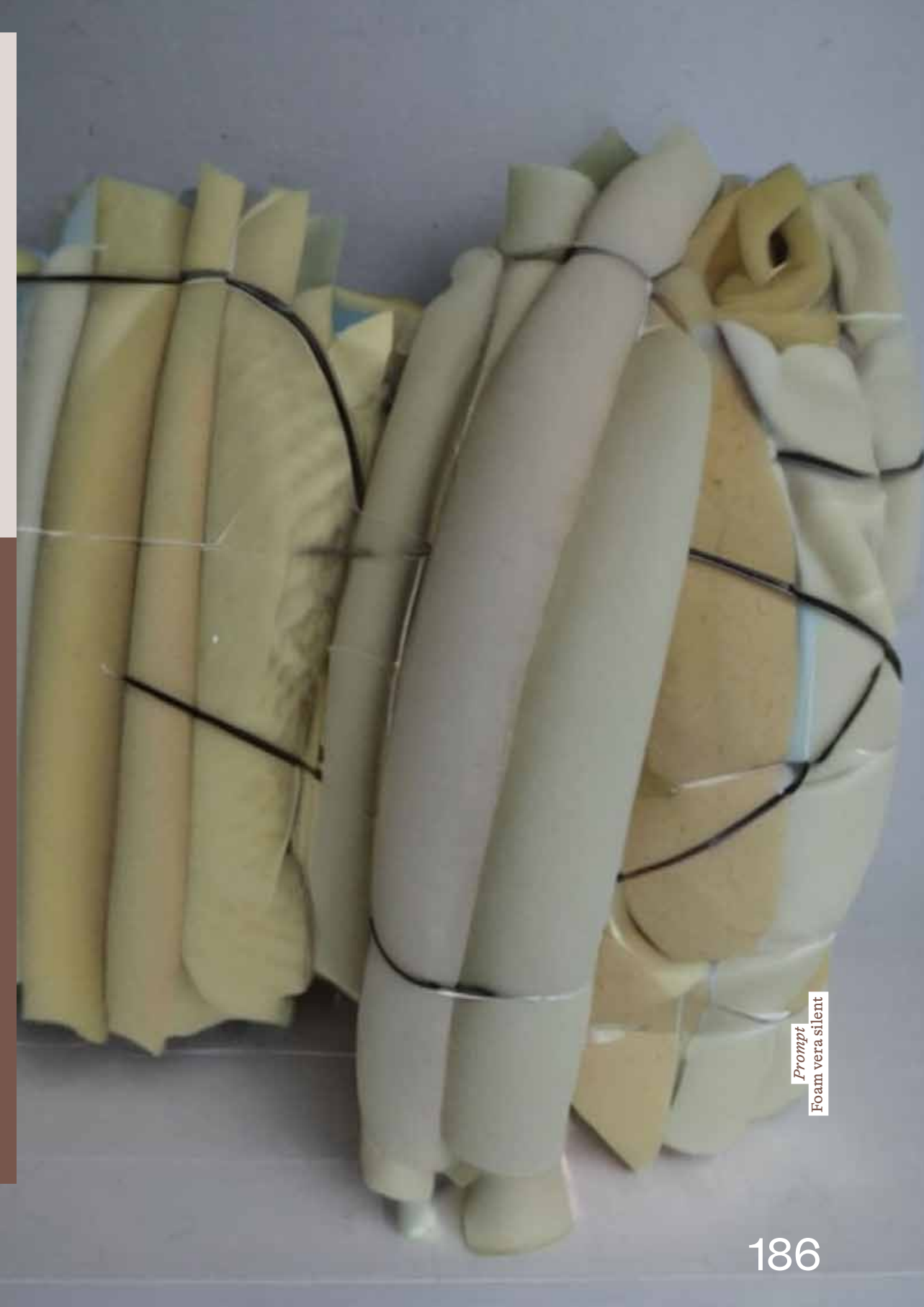
Prompt
Foam vera fighting

184



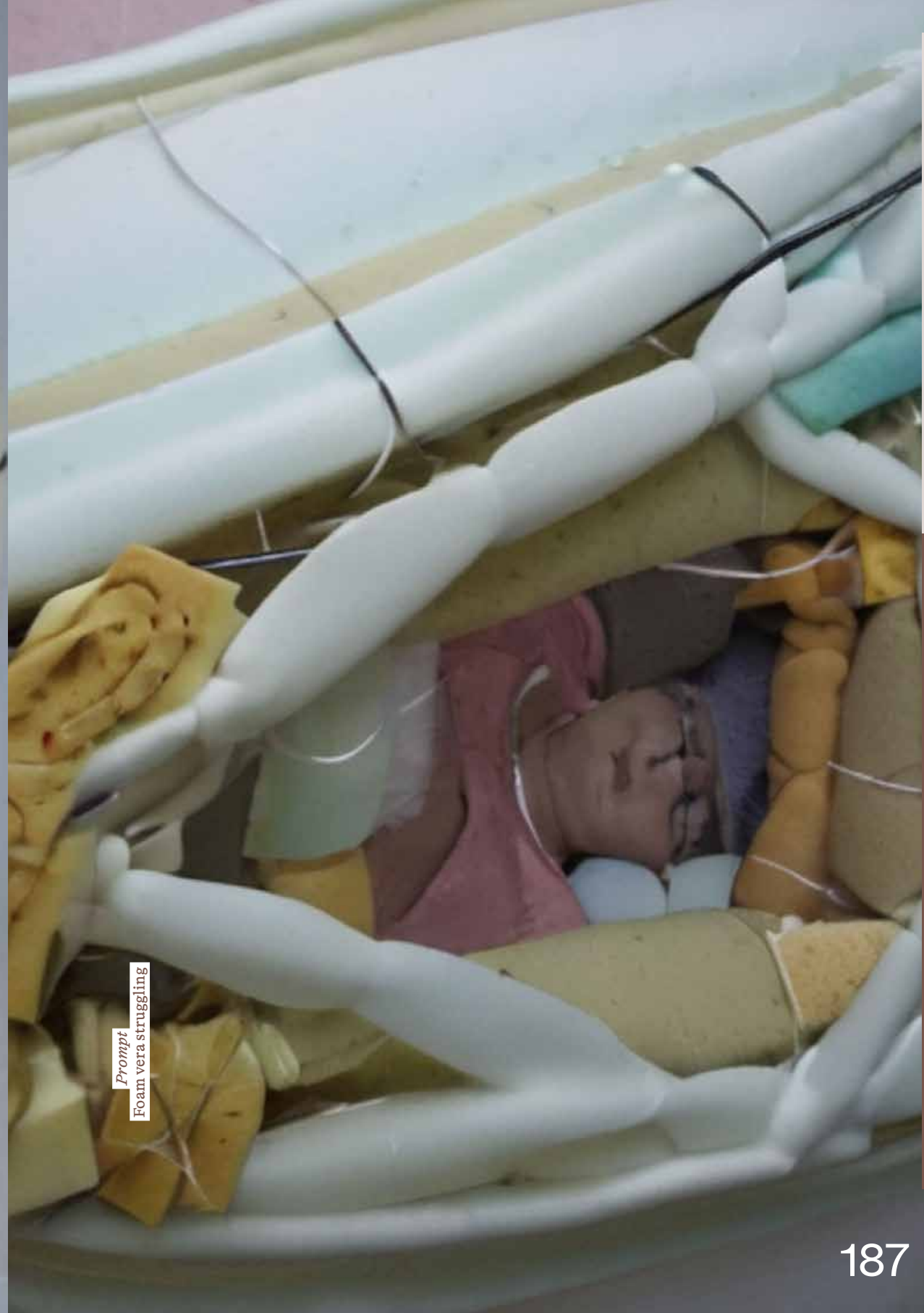
Prompt
Foam vera jealous

185



Prompt
Foam vera silent

186



Prompt
Foam vera struggling

187



Prompt
Foam vera angry



Prompt
Foam vera devastated

This led to a new series of prompts that combined my identity with foam and a spectrum of current emotional states: Foam Vera Fighting, Foam Vera Crying, Foam Vera Hiding, Foam Vera Jealous. Each prompt was a question: How might this material essence of consolation interact with my emotional experience now? The AI interaction transformed from a retrospective analysis of a past work into a prospective exploration of my present self. The prompts became a form of speculative self-care, a way to visualize how the protective qualities I had rediscovered from my past might support me through the vulnerabilities of the present. The inquiry had come full circle: the lingering of the past had become comfort for the now.

6.3

The role of memory in broader reflective AI practice

I started this inquiry with a question about a chair. Training a generative model on my personal archive helped me understand why I kept returning to this object from my past. It also showed me something about how a Reflective AI design practice can work, something different from the approaches in Chapters 4 and 5.

The Objective Portrait work focused on current fascinations: what interests you now, what patterns appear in recent work. The comparative framework of working with other designers and students helped surface how individual perspectives differ while following the same process. This chapter asked a different kind of question. It turned the inquiry backward, toward work that had been made and lost years ago, toward meaning that had never been fully understood. I was

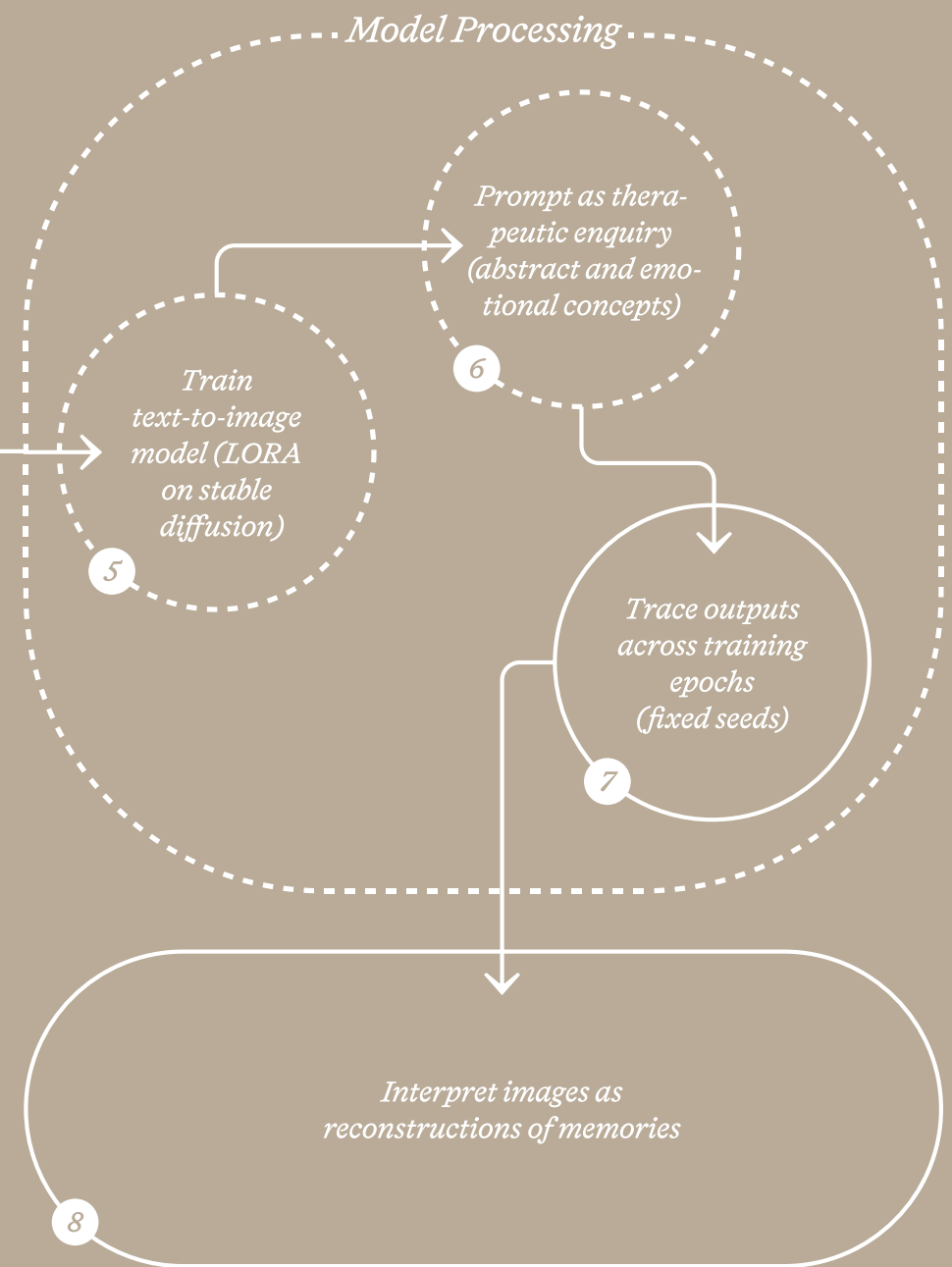
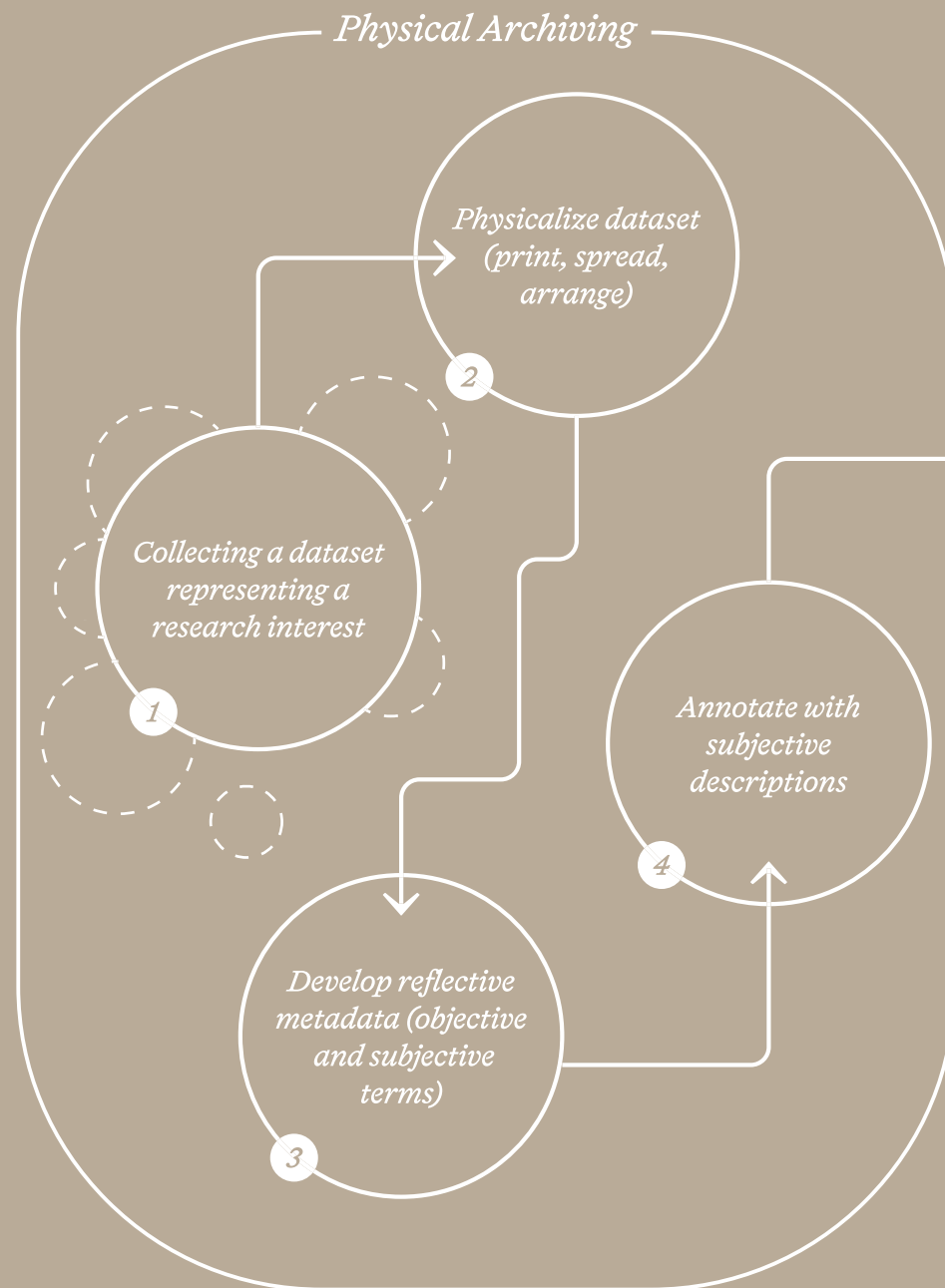
exploring something unresolved, something that had lingered for eight years.

Then, the shift from object detection to text-to-image as a technical choice mattered less than I expected. Yes, the models function differently: one interprets, one reconstructs. But the way I engaged with the outputs remained the same. When you shift the goal from generation to reflection, the generated image becomes a suggestion, an interpretation, a hypothesis. Each image opened space for exploration, prompting questions about what the model had learned, what it emphasized. This was true in previous chapters too. What changed here was the direction: toward my own unresolved past. When designers wonder “*I keep thinking about this old project and I do not know why*,” or “*I want to understand where my current*

practice comes from,” this chapter suggests a path for that kind of question.

Looking back at the process, I see now that the reflection happened everywhere: in selecting which old projects to include, in printing the dataset and spreading it across my

floor, in finding words for visual work, in watching training unfold across epochs, in the vulnerable prompting, in reading the reconstructions. What stayed with me was not only the generated images, but everything I had to think through to arrive at them.



7. From Text-to-Clay

The previous chapters explored Reflective AI through primarily internal processes; reflecting on research interests, fascinations, and personal history. This final chapter of Part II explores what happens when the Reflective AI practice encounters a physical material. By bringing the workflow into dialogue with ceramic practice, this chapter investigates whether

the Reflective AI design practice can be grounded in physical making, and what the friction between both material logics can bring.

7.1 *Introduction*

Clay has long been a source of fascination for me. It is a material that responds directly to human touch, yet it has its own agency, working according to its own behavior. One characteristic of clay is that it resists being rushed. It needs to dry at its own pace, because speeding this up can break the shape. Then, when clay is fired in the kiln, molecular bonds restructure the material from extremely fragile to sturdy. In this process, ceramic can warp, crack or even explode. Glazing, which is a second firing with a specific type of coating, introduces further unpredictability. Glazes melt and fuse with ceramics in ways that remain partially unpredictable, and even the slightest change in ingredients, water level, application technique or kiln temperature changes its behavior and its resulting look. All in all, there are constant moments of uncertainty: cracks can appear during drying, forms can collapse or explode in the kiln, glazes can turn out entirely differently than expected.

This particular material demands a particular craft. Philosopher Lambros Malafouris describes the ceramicist and clay as displaying “*a dance of agency*” (2008, p. 24), where “*the shaping of the pot becomes an act of collaboration between the potter and the mass of wet clay [...] There is a constant tactile but also clearly visible, dynamic tension*” (p. 34). Control shifts between ceramicist and material across moments of making. Sometimes the maker leads; sometimes the material does. You cannot rush clay without it breaking at multiple stages. Instead, you learn to negotiate between its material agency and your own ideas and plans.

When a pretrained, commercial AI model generates an image of a ceramic object, it produces a finished form in seconds, with no memory of this material negotiation. It replicates visual patterns of finished clay objects, without knowledge of the making process behind them. This temporal collapse of process is significant: where ceramic making unfolds through days or weeks through

specific material practice before a piece is finished, AI operates on predicted forms extracted from their making processes.

The instinct when bringing AI into craft practices is often automation: having the system do something for you, making the messy parts cleaner or faster. What if AI generated perfect glazing recipes? What if it generated forms that could be easily 3D printed in clay? At every step of the ceramic process, there could be opportunity to hand over control

to a system. The Reflective AI practice I am building seeks a different approach. Throughout my previous projects, I explored how AI training can surface implicit judgments and fascinations, or how it can help you reflect on something from your past. This chapter extends that inquiry by bringing Reflective AI into dialogue with a physical material, while actively avoiding any type of automation or speeding up the process itself, and instead trying to learn something about my own making practice.

7.2 *Material context*

Why choose ceramics for this particular investigation? Design researchers Giaccardi and Karana (2015) describe how materials are not passive substances waiting to be formed but active participants in the making process. Materials reveal themselves through engagement and shape the maker’s understanding in return. Clay, with its temporal demands and unpredictable transformations, seemed like an ideal test case because of the specific engagement it needs. Its slowness, its physicality, its resistance to be controlled; could these qualities (re)shape how I engage with AI? Could they force a different kind of reflection than the inquiries of previous chapters?

The research context of this chapter played an important role. The project was conducted during

an artist residency at the European Ceramic Work Centre (EKWC), a facility that offers three-month residencies for artists seeking to develop a ceramic practice or realize a specific ceramic-based project. I applied with the intention of bringing AI and ceramics into dialogue as materials, while being situated in a place that was fully attuned and equipped towards this particular craft. I was offered a place to live and work at the center from February to April 2025, alongside sixteen other residents. The residency situated the work in actual practice on two levels. First, the time frame allowed me to set up the AI workflow alongside ceramic making, with time to fail, start over, and be intensively guided by the advisors who worked at the center that taught me much about



7.1

the possibilities in the facilities of the center. Second, being part of a community of artists and designers all working with one material exposed me to a culture of making. I learned about their lives, projects, setbacks and successes, while sharing my own. This exchange felt important: the research was situated in a lived context.

Understanding this project requires understanding how ceramics works as a material. Consider making a single small sculpture. The work begins with clay in its plastic state, shaped freely by hand. Over several days, water evaporates and the clay stiffens to leather-hard, a stage ideal for carving details or attaching elements. Left

longer, it reaches bone-dry: rigid and extremely fragile. The piece then enters the kiln for bisque firing at around 1000°C, a cycle that takes 24 to 48 hours. During this firing, the clay shrinks approximately 8% (depending on the clay type), altering the proportions anticipated during making. After cooling, the bisque surface can receive glazing; chemical, glass-forming minerals mixed with water that can be painted on the bisque sculptures. A second firing at higher temperatures transforms the glaze, often producing colors and textures that have little resemblance to their unfired state. From start to finish, this small sculpture takes roughly two weeks to complete.

7.3 *The ceramic-AI cycle*

The methodology for this chapter builds on a circular process: from clay to data to clay again. This creates a feedback loop that brings the principles of Reflective AI into direct, physical dialogue. The process adapts the disentangled workflow established in earlier chapters (data collection, annotation, prompting, inference) and returns to a fifth step: physical translation. Chapter 3 first explored this principle when physical modifications to a chair responded to the AI's object detection, creating a

seeing-moving-seeing loop. Here, the translation works differently: rather than modifying an existing artifact, I make a new artifact based on what the model generates. These artifacts, or sculptures, are then photographed, becoming a new training dataset for the next iteration. This final step closes and continues the loop, transforming a linear workflow into a recursive cycle. The steps are explained below:

Step 1: Data Collection — *Using past ceramic works as training data*

Step 2: Annotation — *Labeling images with subjective, interpretive descriptions*

Figure 7.1

The research wall documenting the process and the reflections I had during the process.



Figure 7.2
The studio at EKW



Figure 7.3
The kiln hall of EKWC



Step 3: Prompting—*Testing the trained model with abstract emotional concepts*

Step 4: Physical Translation—*Reconstructing AI-generated images in clay*

Step 5: Documentation—*Photographing works at different stages of the ceramic process (e.g. bone-dry state, bisque-fired state, glazed state) to generate new training data, restarting the cycle at Step 2.*

This process was documented on large-scale research wall (see figure 7.1). This display evolved throughout the residency. It combined photographs of ceramic transformations, AI-generated images, process sketches, and handwritten annotations capturing insights and reflections. The research wall became both backdrop for ongoing work and visual representation of the developing practice.

The methodology builds on the same technical approach developed

in Chapter 6: fine-tuning Stable Diffusion 1.5 (Rombach et al., 2022) using Low-Rank Adaptation (LoRA) (Liu et al., 2024). The initial training dataset for this project consisted of just 16 images of past ceramic work; the second cycle (from step 5 to step 2) used 21 images of sculptures made during the residency, all reflecting the time and space that was available to me to finish those sculptures (see Appendix D for complete annotation archives and technical parameters). These numbers are tiny by machine learning standards, but actually work well when it comes to fine-tuning. Additionally, as established in Chapter 5, a ‘small dataset mindset’ (Abuzuraiq & Pasquier, 2024) serves Reflective AI practice: comprehensible scale allows for deliberate, considered annotation rather than rushed labeling, and the resulting model reflects the specific patterns of a personal archive rather than generic aesthetic categories.

7.4 The residency

First Weeks - “Have you touched clay yet?”

The studio spaces at the EKWC are adaptable environments that provide sacred making space for the residents who occupy them. Every three months another resident takes over and makes the space their own, all starting from the same entry point: an empty space. The flow of work

among the residents follows a weekly rhythm where one resident arrives as another leaves. As a new resident, I stepped into an already functioning system of people working and living at the center. Since everyone starts at different weeks but stays for about three months, I encountered studios at every stage; someone who just started working next to someone almost

finished. Being confronted with all these different stages of ceramic processes helped me understand what was possible within the timeframe of the residency period.

The social interactions with other residents became central to the experience. Living and working together for three months, with clay as the shared focus, created particular bonds. Every night someone organized dinner for the fifteen residents. During these meals we discussed our days, shared kiln schedules, troubleshooted technical problems, exchanged glazing recipes, and voiced the frustrations that came with ceramic work. Since everyone spends the entire three months at the center, day and night, I got to know people through their ceramic processes as we were like roommates. Whether they were experienced ceramicists or working with clay for the first time, I observed how they all found their way to approach the material and

deal with its demands. This created a genuine ceramic-making community where clay became the primary focus of daily life and conversation.

During the first week, the question I heard most from other residents was: “*Have you touched clay yet?*” Most residents, especially those new to ceramics, spend considerable time thinking before actually working with the material. There was a mental hurdle between planning and making that almost every resident experienced and had to overcome. The question highlighted something about ceramic practice: it requires direct physical contact with the material rather than only thinking about what to do. This emphasis on material contact, and simpler said, *doing*, would prove crucial for my understanding of how ceramic practice informed approaches to AI interaction, where similar questions about the relationship between thinking and doing would emerge.

● *Step 1: Building the dataset*

The first step was to construct the training dataset. Building on the method from the previous chapter, I turned again to my own creative past: a collection of small ceramic sculptures from seven years prior, representing my first interaction with clay. They were objects on the edge of recognizable forms like cups or bowls, but with additions that made them seem misfitted, or gave them creature-like characteristics like legs or heads. These sixteen images were particularly suitable because they already possessed dataset-like qualities like consistent coloration, similar glazing techniques and uniform photographic composition. The choice also reflected a specific moment in my practice: the dataset captured work made by my own untrained hands, created during an assignment focused on intuitive material exploration. *Figure 7.4* and *7.5* show an overview of this dataset.

● *Step 2: Annotating the Data*

The annotation strategy built on the previous chapter’s approach:



→ Annotation

group of cream colored ceramic objects against a green backround, chaotic atmosphere, provocative, perverse



→ Annotation

group of cream colored ceramic objects against a green backround, chaotic atmosphere, provocative, perverse, gang



→ Annotation

group of cream colored ceramic objects against a green backround, chaotic atmosphere, provocative, perverse, gang, playful



→ Annotation

ceramic misfits, virus, melting, group, mistakes, fleshy



→ Annotation

ceramic misfits, virus, melting, group, mistakes, fleshy

Figure 7.4

The initial training dataset: sixteen images of ceramic sculptures I made seven years earlier, representing my first experiments with clay. These works already showed tendencies toward multiplying forms and creature-like vessels that I had enacted.

→ Annotation

ceramic, a humble cup, shaped like a bull, organic, stubborn



→ Annotation

two ceramic figures lounging, spreading, sunbathing, relaxing, green background



→ Annotation

three porcelain figures, cutlery, lounging, long



→ Annotation

ceramic cup with lots of handles, virus, surrealism, visual confusion, green background



→ Annotation

three ceramic cups with lots of handles, virus, surrealism, weird, visual confusion, green background



moving beyond simple visual description toward subjective, interpretive labeling. The central question shifted from “*What is visible in this image?*” to “*What do I think the objects in these images mean?*”

The resulting annotations became descriptions that combined two layers of information. Factual, verifiable elements were noted to ground the model’s learning in a consistent visual reality, including terms like ‘green background,’ ‘ceramic,’ and ‘cream colored.’ Layered on top of this objective base were the subjective, conceptual terms intended to teach the model the associative thinking inherent in the original artistic work. *Figure 7.4* shows a sampled overview of the dataset images and their annotation.

Since the ceramic work in the dataset explored the boundary between mundane objects like vessels or cups, and creature-, human-like forms, the annotations encoded this concept. For example, a cup-like form was given the interpretive label ‘bull-like’ to capture its perceived posture and character, while another piece with multiple handles was tagged with the conceptual terms ‘virus’ or ‘misfit.’

Looking back at the annotated dataset, I saw a pattern. My labels showed a consistent aesthetic vocabulary: forms that multiply or proliferate (‘virus’), shapes that provoke or unsettle (‘misfit,’ ‘provocative’), objects that blur boundaries between functional and biological (‘bull-like’ cups that have posture and character). The annotations exposed a recurring fascination.

● Step 3: Prompting the fine-tuned model

In standard practice, prompting involves providing straightforward textual descriptions: ‘A chair against a white background.’ The relationship between prompt and output follows predictable logic. When working with a model fine-tuned on subjectively annotated personal data, prompting can become investigative, as I have described in Chapter 6. Each prompt tests what the model internalized from my annotations. The model also required a specific trigger word (‘ceramic’) to activate the LoRA and draw on my training data rather than its broader knowledge base.

I experimented with prompts ranging from concrete forms to abstract concepts. When I prompted it, for example, ‘ceramic jealousy,’ I was asking the model to give visual form to an emotion through the lens of my ceramic aesthetic. The model would propose an interpretation, offering a starting point for further inquiry.

The ten prompts I chose for material interpretation covered emotional states (‘ceramic sadness,’ ‘ceramic jealousy,’ ‘ceramic female pain’),

Figure 7.5
Annotation of the dataset on the research wall (bottom)
and the prompt lexicon (above)

7.5



7.5

Prompt: Red ceramic surrealism



Prompt: Ceramic sadness



Prompt: Ceramic surrealism



Prompt: Ceramic ignorance



Prompt: Dry ceramic jealousy



Prompt: Ceramic octopus



Figure 7.6
An overview of the images resulting from prompting the self-trained model with subjective terms.

Prompt: Ceramic divorce



Prompt: Ceramic female pain



artistic concepts ('ceramic surrealism,' 'red ceramic surrealism'), figurative elements ('ceramic woman sitting waiting,' 'ceramic mask'), personal circumstances ('ceramic divorce'), and concrete forms ('ceramic octopus') as a baseline. Each prompt tested how the AI might channel different types of concepts through my own ceramic aesthetic. The resulting collection of generated images became a visual lexicon: a reference set mapping abstract concepts to possible ceramic forms. An overview of this lexicon is shown in *figure 7.6* (top image).

Looking at the generated images, I recognized my own work, recombined in ways I had not seen before. The model had learned certain qualities from my annotations and was now channeling abstract emotions through those visual tendencies. The images felt both familiar and surprising: familiar because they reflected my aesthetic, surprising because I had never explicitly connected concepts like 'jealousy' or 'sadness' to specific ceramic forms.

● *Step 4: Interpretation of images in clay*

The translation process began with a starting condition that distinguished it from the other ceramic practices I saw around me in the center. Instead of working from an internal creative idea or responding directly to clay's material properties, I was attempting to recreate something that already existed as a complete visual representation. I was trying to extract a making process from a finished image. The AI-generated images appeared to document real ceramic objects, complete with studio lighting and clean backgrounds that suggested professional photography of finished works. But these objects had never actually existed in physical form. They were statistical interpretations of training data, operating through probabilistic relationships between tokens and pixels that have little to do with clay's spatio-temporal materiality, for example how weight distributes or how sections dry at different rates. This created a paradox: I was attempting to recreate objects that looked authentically made but had never undergone any actual making process. To illustrate what this meant in practice, I highlight two prompts and the challenges I encountered with them.

'Ceramic Surrealism' presented the first practical challenges. The clean, isolated presentation provided no contextual clues about size. Given my limited experience with ceramic practice, I chose to work within a manageable range of twenty to thirty centimeters in height. The images were also flat: a single viewpoint showing one angle of an object I needed to build in three dimensions. What did the sculpture look like from the back? From the side? These were decisions I had to make myself. More significantly, while the images seemed

visually coherent, they often contained spatial impossibilities that became apparent while translating them in clay. I could not determine only from the image whether certain areas curved inward or outward, how parts attached to each other, or which elements existed in foreground versus background. Some generated forms appeared to float or balance in ways impossible without hidden support structures. Rather than viewing these impossibilities as failed sculptures, I treated them as invitations to fill in the blanks with my own interpretation, potentially leading me into new ways of making. *Figure 7.7* shows a first attempt: the generated image seems to float, and I solved this by making a kind of hidden base under the shape.

'Ceramic Jealousy' made visible how emotional concepts could take on meaning through this translation process. The AI's image suggested repetitive elements and something that looked like a ceramic cup or bowl, but with branching elements coming out of it, echoing the 'virus' and 'misfit' vocabulary from my annotations. But I had to decide: How heavy should this form feel? Where should weight concentrate? What surfaces should feel smooth versus textured? How should those branching elements actually connect? Whether the form should appear stable or on the verge of collapse? Through these choices, 'jealousy' became associated with forms that reach outward and up, like branches. The meaning emerged through translating the AI's image into clay, filtered through my hands. *Figure 7.8* shows the first iteration.

Throughout both translations, I developed specific construction strategies: support systems using wooden planks and sponges to maintain structural stability during the leather-hard and drying phases when clay is most vulnerable to collapse. *Figure 7.9* and *7.10* show some of these techniques.

● *Step 5: Bone-dry becomes dataset, repeating the cycle*

When all sculptures completed their forming stage and entered the drying phase, I arranged them on small wooden planks covered with a white cloth so they could be transferred to the kiln easily. Viewing the works together, I realized I had generated a new kind of dataset: same clay type, identical drying stage, uniform presentation method, with only the sculptural forms varying. I documented this phase digitally, photographing each sculpture in this bone-dry state. The sculptures derived from AI-generated images were now transformed back into visual data, closing one loop of the circular process.

This new dataset differed from my original training data in a crucial way. Seven years ago, when I made those first ceramic pieces, I was experimenting with clay without knowing what I



Figure 7.7
The generated image for the prompt 'Ceramic surrealism' on the left, and the image reinterpreted in clay on the right.





Figure 7.8
The generated image for the prompt 'Ceramic jealousy' on the left,
and the image reinterpreted in clay on the right.



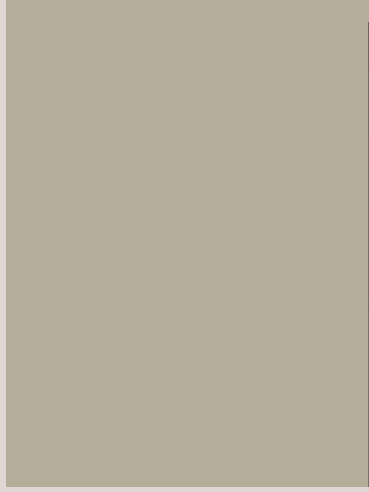


Figure 7.9
Making ceramic ignorance

Figure 7.10
Making red ceramic surrealism



was doing. Those works showed certain tendencies in my making choices (multiplying handles, shapes that looked both like vessels and creatures) but I had not named those tendencies or understood why I gravitated toward them. Now I had deliberately worked with those same tendencies, using them to translate abstract emotional prompts into clay.

To make this a proper training dataset, the photographs needed, again, to be annotated, moving the cycle back to Step 2. But now I had a different starting point than with the original dataset. Since each sculpture resulted from specific prompts, I incorporated those prompts directly into the annotations. The sculpture generated from ‘ceramic jealousy’ received the tag ‘dry ceramic jealousy, on a white cloth,’ combining the original prompt with its current material state. An overview of this dataset is show in *figure 7.11*.

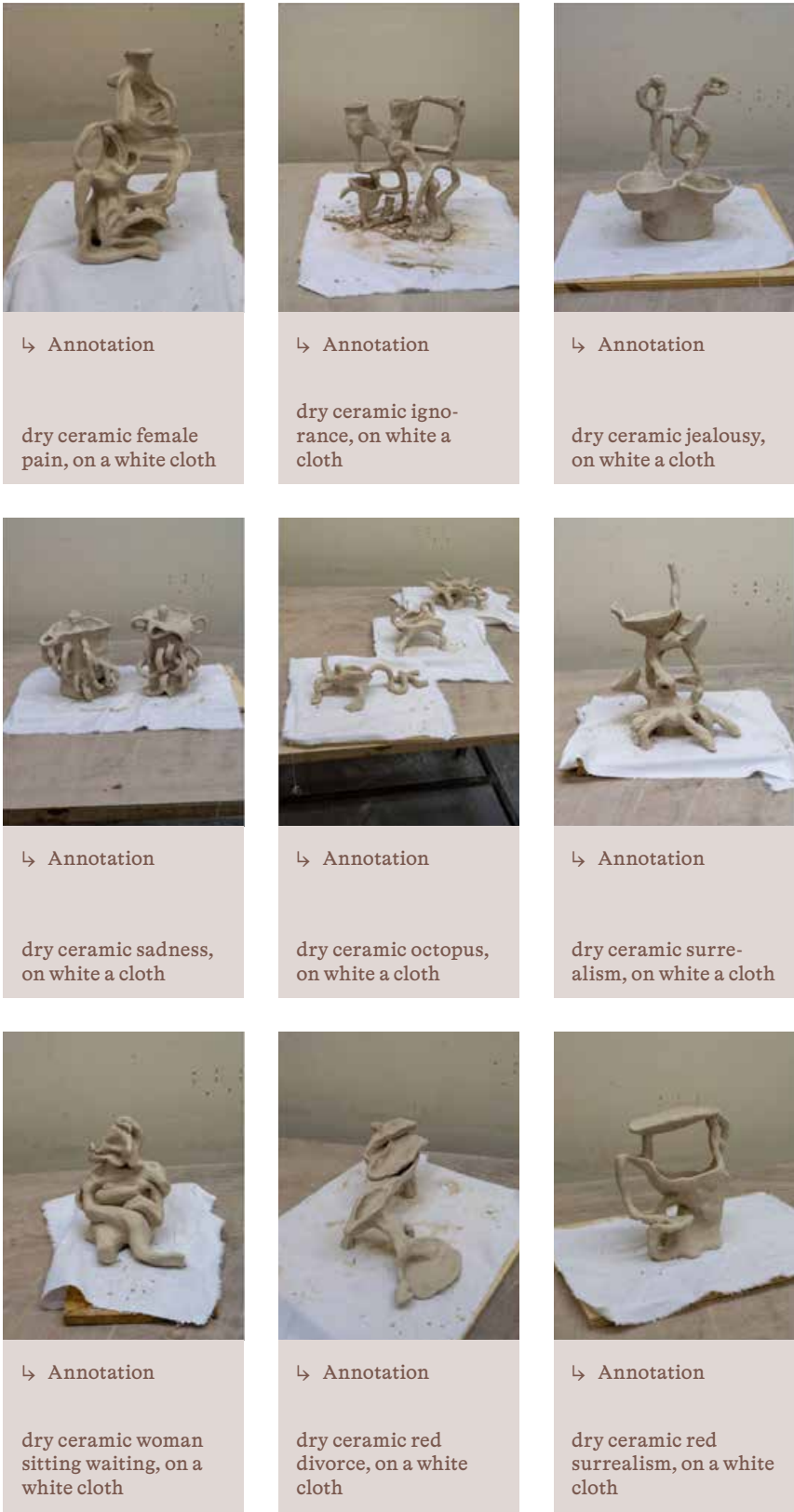
● *Back to step 3: Prompting the bone-dry dataset*

To see what the model that was trained on the bone-dry dataset had picked up, I prompted it with the same annotation tags I had used for the bone-dry sculptures. I typed, for example, ‘dry ceramic jealousy on a white cloth’ and generated new images.

When I looked at the results, I started to recognize my own work. The images looked more like my ceramics than the first generation had. I could see specific choices I had made: sharp points that branched outward, multiple handles sprouting from a single form. These were the same strategies I had used when building the ‘jealousy’ sculpture in the first cycle. When I looked at the images generated from ‘dry ceramic divorce,’ I saw a separation between two sides in the form of bowl-like features, matching how I had approached divorce through clay. Patterns emerged: the prompt ‘jealousy’ now triggered multiplication and sharp branching forms on top of a heavy base, while ‘surrealism’ triggered sculptural shapes with a dish shape on top of branches. The model seemed to have encoded statistical associations between my prompt words and the visual patterns from my sculptures. Looking at these images felt strange. They were clearly mine, I could see my own preferences, but I had not made them. The AI was recombining my making tendencies in new ways. *Figure 7.12* shows examples of the second cycle of AI generations.

The model had also learned something else: the material state of bone-dry clay and the environmental aesthetics of EKW C. The generated images no longer showed glossy, finished ceramics against clean studio backgrounds. They showed dry, unfired forms on brown tables and gray walls, the specific visual context of my workspace. This made them easier to relate to and, practically, easier to realize

Figure 7.11
Bone-dry ceramic dataset with annotations for the next training cycle



Prompt: Dry ceramic surrealism



Prompt: Dry ceramic red surrealism



Prompt: Dry ceramic sadness



Prompt: Dry ceramic octopus



Prompt: Dry ceramic jealousy



Prompt: Dry ceramic ignorance



Figure 7.12
An overview of the second cycle of generated images based on prompting the bone-dry dataset.

Prompt: Dry ceramic divorce



Prompt: Dry ceramic female pain





Figure 7.13
The second cycle of interpretation: Ceramic surrealism





Figure 7.14
The second cycle of interpretation: Ceramic ignorance





Figure 7.15
The second cycle of interpretation: Ceramic divorce





Figure 7.16
Process of building the second iteration of clay sculptures.



Figure 7.17
From top to bottom: Ceramic red surrealism, ceramic jealousy, ceramic octopus, ceramic female pain.



in clay for a subsequent cycle of making. The first-cycle images had depicted impossible finished objects; these depicted impossible objects in a familiar material state, one I was actively working with.

Even though the shapes felt more familiar, the same structural challenges persisted. The generated images continued to show unclear connections, seemingly impossible forms, shapes that seemed inherently difficult to construct. But I had learned from the first iteration. The AI's proposals had pushed me to develop new construction techniques. Many pieces needed to be built in separate sections, each reaching specific drying stages before being joined. I continued using temporary support systems combined with ceramic supports that would join the sculpture in the kiln. Each challenging form became a technical problem to solve, expanding my understanding of what was structurally possible with clay. The second cycle produced fewer physical sculptures (7) than the first, even as the training dataset grew. Working at larger scales meant each piece required more construction time and significantly longer drying periods. Where a small sculpture might dry in days, the larger forms needed weeks of careful monitoring, with plastic wrapping adjusted to control moisture loss. The residency timeline constrained how many pieces could complete the full ceramic process. An overview of the sculptures and their processes are shown in *Figure 7.13, 7.14, 7.15, and 7.16.*

● *The kiln*

This circular process could continue indefinitely: each new clay state generates new training data, which generates new images, which generates new sculptures. Firing the bone-dry works would create yet another dataset, this time of bisque-fired ceramics, making the model progressively more specific to the different phases of the ceramic process. A model trained across wet clay, leather-hard, bone-dry, and bisque-fired states would begin to learn something about the material's transformations, not just its final appearance. The residency, however, had a timeline, and the sculptures needed to be finished. I decided to fire the sculptures, to see which of the AI's proposals could survive the full material process, and which would fail. During the drying and the careful transportation to the kiln, several pieces revealed mistakes I had made during the sculpting stages.

'Ceramic Surrealism' was particularly resistant to physical existence, as the piece broke three times. First, it cracked during drying. The sculpture had thick sections at the base and

Figure 7.18
Kiln loading



7.18

thin extensions at the top. As they dried, the thin parts contracted faster than the thick parts, creating stress fractures where they met. I rebuilt it, trying to make the thickness more even throughout, and wrapped it with plastic to control the drying times. The second time, the tip broke off during transportation to the kiln. The form was top-heavy, most of the weight concentrated high up on a narrow base. When I carried it, the center of gravity was wrong, making it tilt and break. I rebuilt it again, trying to reinforce the connection point by building extra support structures. The third attempt exploded in the kiln. I had rushed the drying time, and moisture trapped in the thick base turned to steam during firing. The pressure shattered the piece from the inside. As Malafouris (2008, 25) writes, *“Trying to separate cause from effect inside the loop of pottery making is like trying to construct a pot keeping your hands clean from the mud”*. Each failure was attributable neither entirely to my choices nor entirely to the clay: the form’s top-heaviness was the AI’s proposal; my construction attempted to realize it; the clay’s physics refused. Ironically, the prompt had named something true: surrealism depicts what cannot exist in reality, and this form struggled to exist as a ceramic object. An overview of the evolution of ‘Ceramic Surrealism’ is shown in *figure 7.19*.

‘Ceramic Ignorance’ presented similar problems. The sculpture had two large sections connected by a single thin bridge of clay, like two heavy masses balanced on a pencil. When I carried the piece to the kiln, the connection point cracked under the weight because the transportation cart hit a little bump on the floor. Looking at the AI-generated image, I could see why it failed: the image showed the form from one angle, where the connection looked plausible. But from the side, it was clearer: two heavy sections, one fragile link. Here too, the prompt had named something true: the form embodied the AI’s genuine ignorance of ceramic reality. An overview of the evolution of ‘Ceramic Ignorance’ is shown in *figure 7.20*.

These failures taught me something about the difference between making an image and making an object. The AI processes images only. It can make something look balanced, look structurally sound, look dramatic. But it does not know what clay feels like when you lift it, that thick sections dry slowly, that thin connections cannot hold heavy forms. My role became translation between two kinds of logic: the AI proposed visual ideas, I figured out if those ideas could survive gravity, drying time, and fire. I had to make explicit what ceramic practice knows implicitly. Sometimes I succeeded. Sometimes, like with ‘Ceramic Surrealism,’ the form simply refused to exist.

● Glazing

With seven sculptures of the second cycle surviving the kiln, I wanted to complete them and explore glazing. At the EKWC, glazing proved even more unpredictable than ceramic construction itself due to the wide range of possibilities in this technique. The center maintains a glazing library where residents mix their own formulations, selecting base glazes and adding color pigments or chemicals to create texture effects.

Another challenge of glazing is that the glaze looks completely different before and after firing. A murky gray liquid might emerge from the kiln as glossy white and a brown paste might become vibrant turquoise. This is often described as painting with your eyes closed. Testing on small sample pieces, and firing them helps to imagine what glazes will look like, but even tests do not guarantee results, since how thickly you apply the glaze, how the kiln fires that day, and how the glaze interacts with your specific form all affect the outcome. Once a piece is glazed and fired, the chemical bond is permanent and you cannot undo it.

What I learned from the ceramic community at EKWC was a particular approach to this uncertainty: acceptance. It is difficult to imagine what a glazed piece will look like, and most of the times, a glaze turns out differently than expected. This is something ceramists learn to work with. You prepare, you anticipate, but you cannot fully control what happens inside the kiln, therefore relating to the outputs of the kiln asks for always managing your expectations.

For the seven sculptures, I developed a glazing approach inspired by how datasets work: create unity through consistency, create variety through controlled differences. I chose the same base glaze for every sculpture, a matte white, then added different color pigments to each one. All the sculptures would have the same surface quality but different colors. Unity in texture, variety in hue. Like a dataset where all images share the same background and lighting but the objects themselves differ.

I also noticed something specific in the AI-generated images: they often showed unrealistic shadows and highlights that created ambiguity about depth. The lighting did not follow real physics. Sometimes shadows fell in impossible directions, highlights appeared where no light source could create them. I wanted to translate this quality into the glazing. I used a spraying technique, applying one color from a consistent angle on top of the base glaze. This created gradient effects, darker on one side, lighter on the other, that mimicked non-existent light sources. The shadows suggested depth that was not actually there, just like the AI images



Figure 7.19
How 'ceramic surrealism' went through different iterations, the bottom row showing the breakage

Figure 7.20
How 'ceramic ignorance' went through different iterations, the bottom row showing the breakage. The final piece can only stand with a support.





Figure 7.21
Examples of what the glazing looks like when it is just applied





Figure 7.22
Examples of what the glazing looks like when it is just applied



7.23

suggested depth that was ambiguous. The glazing became a way to bring the AI's visual logic into the ceramic surface itself.

When the sculptures came out of the kiln after the firing, my first reaction was slight disappointment. The colors were darker than I expected, and the surfaces a bit dull, the gradient effect was subtler than I had hoped. But as I placed them on the studio tables and spent some time with them, looking at them in different lighting throughout the day, my perception started to shift. I started to notice how the muted tones worked together, how the subtle gradients created quiet depth rather than dramatic contrast. What I had initially judged as something different than I expected, slowly became something I grew to appreciate. By the end of the residency, I had come to love these objects, because they were not what I had planned.

7.5 *Reflections...*

The previous section detailed the process of the residency as it unfolded. Here, I step back to reflect on what emerged from the friction between the material logic of clay and the computational logic of AI.

At the start of the residency, I asked whether clay's particular qualities could force a different kind of reflection than the purely digital inquiries of previous chapters, and what the process of translating AI outputs into physical form might reveal about my own associations between form and feeling. This discussion addresses both questions. The first section examines what the circular process made visible about my own ceramic practice. The following section explains how the methodology created conditions for that visibility: through

working backwards from finished images, maintaining slowness and difficulty, and treating unpredictability as material character.

7.5.1 *...on my ceramic practice*

At the end of the residency, looking at the final glazed sculptures arranged in the studio, I noticed that a consistent aesthetic vocabulary and making style ran through all of them. These qualities had been present in my work for years, but I had never recognized them as a vocabulary before.

The circular process made this vocabulary visible by requiring articulation at every step of the process. Annotation forced me to find words for patterns I recognized intuitively, prompting showed what visual meaning those words could carry, and translation turned abstract concepts into material decisions that

↶ *Figure 7.23*
Applying the glaze

accumulated over weeks of making. Each iteration added specificity as the model learned how I interpret emotions through form, and by the second training cycle it had started to encode my associative logic in making. By the end of the residency, the model functioned as a mirror reflecting my aesthetic back to me. ‘Jealousy’, as a ceramic concept, came to mean multiplication, outward branching, points that proliferate through accumulated decisions about weight, texture, and sharpness, and the exchange between what the model proposed and how I interpreted them in clay.

A fair question is whether this required AI at all, since ceramicists develop deep knowledge of their own practice through years of making and style emerges through repetition over time. What the AI contributed, however, was a structure for externalization. Annotation required putting words to visual patterns I recognized, and the abstract prompts posed questions I might not have asked myself, like what jealousy might look like in my ceramic vocabulary. The model’s generations functioned as an outside perspective, showing my tendencies recombined in ways that made them easier to recognize as mine. The insight may have emerged eventually through making alone, but the methodology created a discipline that made it recognizable within three months. This opens questions about my own practice: why do I keep making forms that multiply, and what draws me to boundary-crossing between vessel and creature? The process did not

answer these questions, but it made them possible to ask.

7.5.2 ...on the conditions for reflection

What made this reflection possible? At several points, I could have chosen methods that would bypass the slower phases of ceramic practice: modeling the generated shapes digitally, 3D printing them, and casting clay from molds, or optimizing the model to produce structurally buildable shapes. Each of these choices would have closed off meaning making through hands-on interpretation. What follows examines three conditions that kept that space open: working backwards from finished images, maintaining slowness and difficulty, and treating unpredictability as material character.

Working backwards

What distinguished this process from typical ceramic practice was the reversal of the creative workflow: starting from a finished image and working backwards toward a making process. When the AI generated an image of ‘ceramic divorce,’ it presented what looked like a finished object, complete with studio lighting and surface texture. The image had all the visual markers of something already made. Each generated image invited a question: *how might this form exist physically, and what making process could produce this structure?* The impossibilities in the images, elements that floated or connections that remained ambiguous, were invitations to interpret rather than problems to solve. If the AI had produced buildable blueprints, I would have

executed them. Because it produced images that required interpretation, I had to make decisions that revealed my own preferences and tendencies. The meaning of ‘divorce’ emerged through how I chose to resolve the ambiguities, and those choices accumulated into a recognizable pattern over the course of the residency.

Slowness and difficulty

The temporal mismatch between ceramic practice and conventional AI generation became a generative tension in this inquiry. Where clay dictates a schedule of patience across days and weeks, image generation models produce outputs in seconds. Standard AI interfaces encourage rapid iteration, a mode of engagement where outputs are quickly consumed and discarded. The clay acted as a brake on this speed.

A question I encountered frequently was: why not 3D print the forms? The impulse behind the question I saw as a form of optimization logic: if a technology can realize complex geometries without risk of collapse, why choose the harder path? 3D printing in ceramic practice involves reflection oriented toward execution of a set form; I was seeking reflection on meaning. My hand-building process was continuous interpretation. When building ‘ceramic divorce,’ the abstract concept gained specificity through decisions I had to make along the way. If I had 3D printed the AI’s proposals, I would have bypassed this interpretive work entirely.

To be clear, the AI itself was not slow. It generated images of ‘ceramic

divorce’ in seconds, and I could have prompted it endlessly, cycling through variations without pause. What slowed the process was the ceramic translation. To engage with a generated image meaningfully required committing to days or weeks of physical work before I could move to the next iteration. I had to look at the same generated image repeatedly, from multiple angles, imagining how ambiguous elements might translate physically. I had to develop interpretations deep enough to guide my making process rather than relying on immediate visual judgment. The clay insisted on slowness, and that slowness created the conditions for reflection to take root.

Throughout this thesis, the Reflective AI practice has created slowness through annotation, data collection, and training; here, ceramic practice extended that slowness into the physical world. Hallnäs and Redström (2001) describe slow technology as taking time to learn, to apply, and to find out the consequences of using it. I could not see what a prompt like ‘jealousy’ meant until I had built it, could not know whether my formal choices worked until the piece emerged from the kiln. Each sculpture was a proposition that took weeks to test. This is what Odom et. al. (2022) call an ‘accumulative quality’: meaning that builds through repeated encounters rather than arriving all at once.

This temporal commitment also changed how I related to the AI’s outputs. In rapid AI interaction, each generation feels like a finished product to be evaluated and either

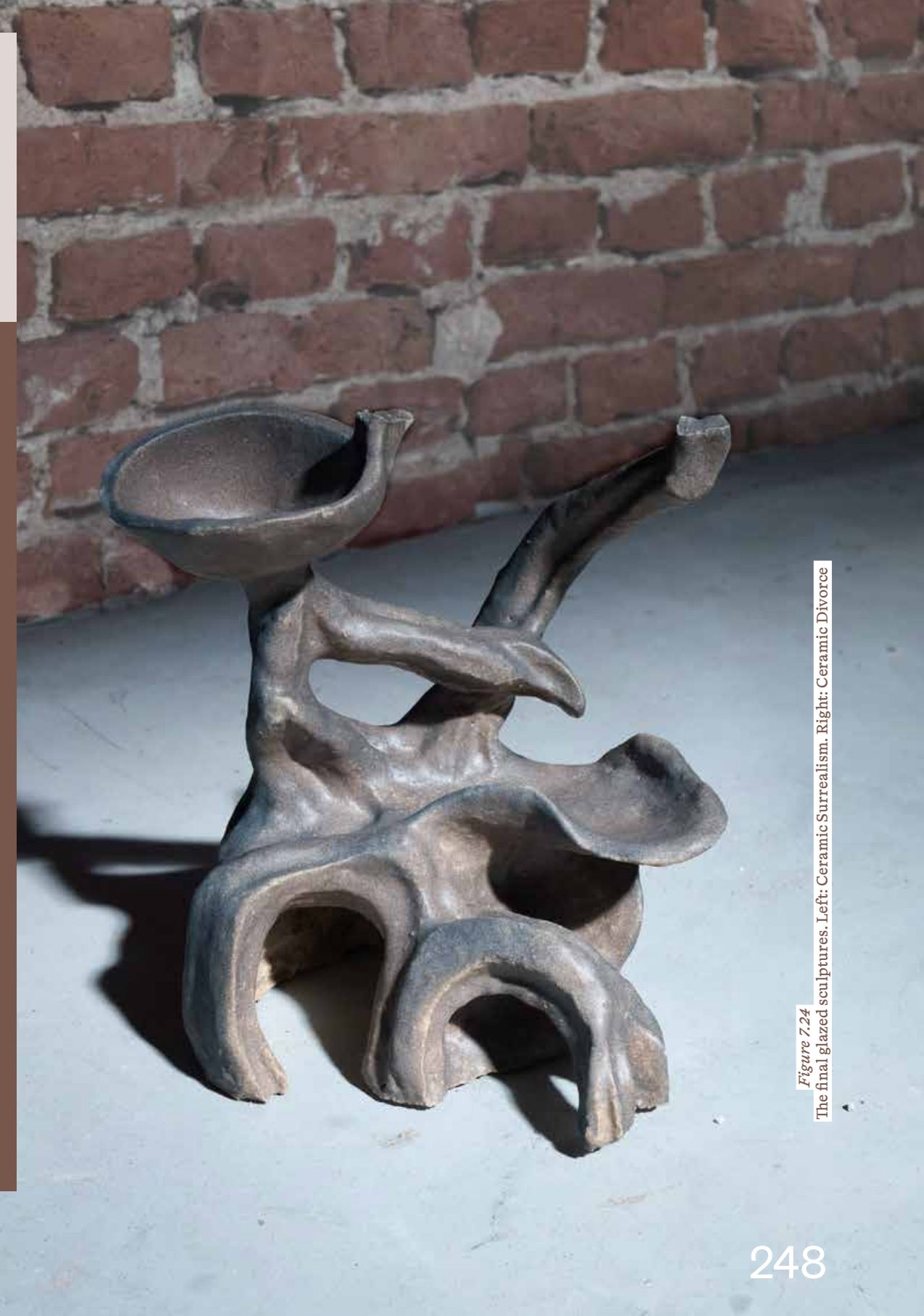


Figure 7.24
The final glazed sculptures. Left: Ceramic Surrealism. Right: Ceramic Divorce





Figure 7.25
The final glazed version of Ceramic Ignorance





Figure 7.26
The final glazed sculptures. Left: Ceramic Jealousy. Right: Red Ceramic Surrealism





Figure 7.27
The final glazed version of Ceramic Octopus

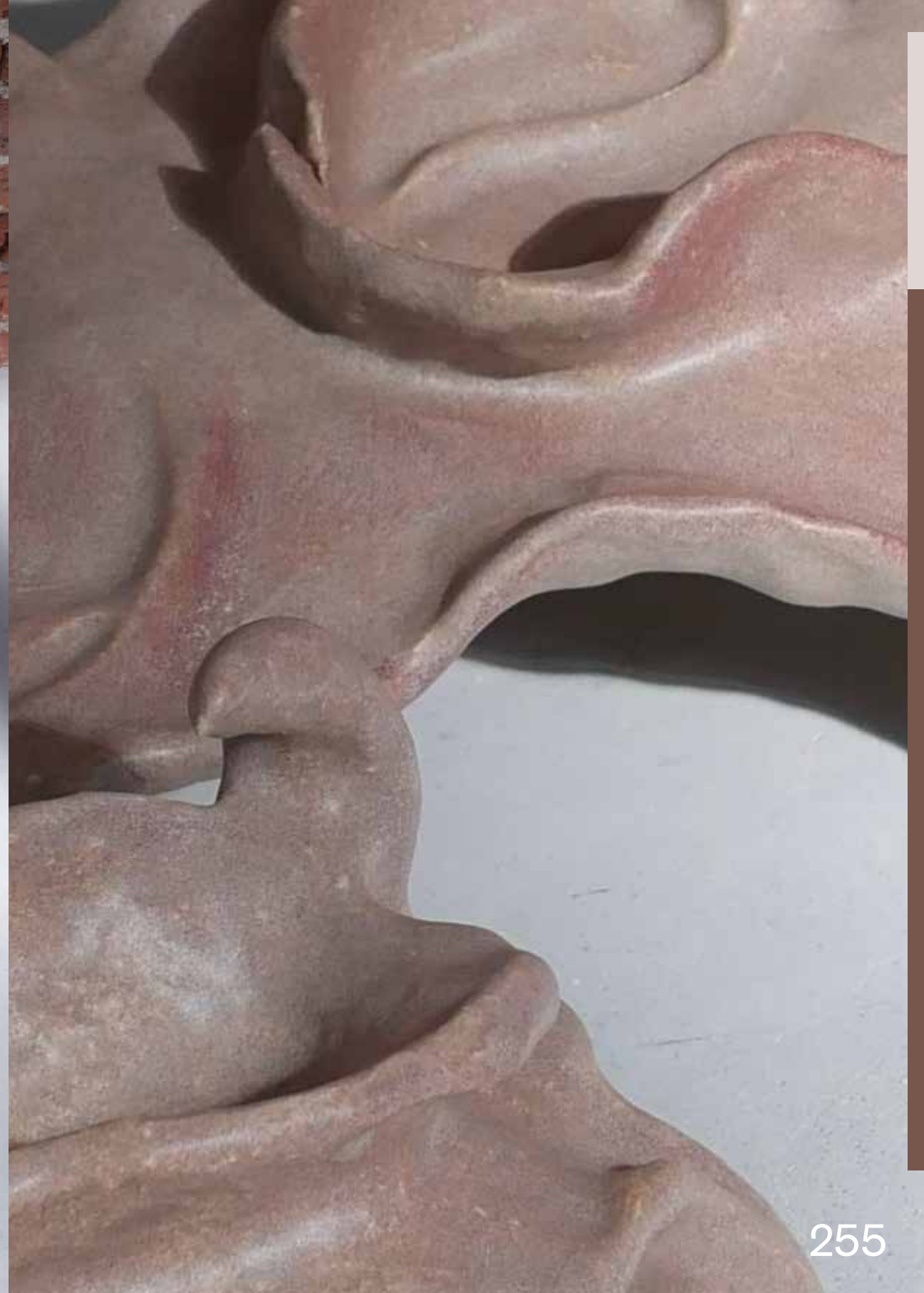




Figure 7.28
The final glazed version of Ceramic Female Pain



accepted or discarded. Ceramic practice treats objects as continuously transforming through multiple states, from leather-hard to bisque to glazed, each with different qualities and possibilities. By working within ceramic time, I began to treat the AI's images similarly: as states in transition rather than endpoints.

Unpredictability as material character

The ceramic practices I observed at EKWC with the other residents operate on a principle that contrasts with optimization logic: unpredictability is a material quality to understand and work with rather than something to explicitly minimize. Of all the stages where this unpredictability appeared, the kiln felt most significant as it was kind of a test-space to see if your sculptures were constructed or glazed successfully. I recognized the kiln as analogous to computational black boxes¹⁰. Both accept prepared inputs and produce transformed outputs through partially opaque mechanisms. In ceramics, we cannot observe the molecular transformations as clay is being fired; in machine learning, we cannot directly observe and make direct sense of the billions of parameters adjusting during training and generation. Both require developing anticipatory knowledge, understanding what preparation tends to lead toward

10 A "black box" in computational contexts refers to a system whose internal workings are opaque to the user: we can observe what goes in and what comes out, but the processes that transform one into the other remain hidden or difficult to interpret.

desired outcomes while accepting that perfect prediction remains difficult to reach.

The difference lies in how these communities frame unpredictability. When a glaze runs unexpectedly or a form warps in the kiln, ceramicists often treat this as material character, sometimes even as a happy accident worth keeping. Dominant AI development culture moves in another direction. Nachtwey and Seidl (2024, p. 98) describe a 'solutionist' worldview in which "*the entire social world is full of bugs but can also be fixed with the right technology and entrepreneurial mentality*". Generative AI development embodies this logic: each model update promises fewer errors, better anatomy, more accurate outputs. Midjourney's V6.1 announcement on X, for instance, highlighted '*more coherent images (arms, legs, hands, bodies)*' and '*more precise, detailed, and correct small image features*' (Midjourney, 2024). The goal is optimization, and unexpected results are problems to be solved in the next version. This project tested what happens when applying the ceramicist's mindset to the algorithm.

This reframing does not mean accepting all AI outputs uncritically. Ceramicists do not simply throw clay in the kiln and hope for the best; they develop sophisticated knowledge about firing curves and chemistry, learning to anticipate outcomes while accepting that some surprises remain inevitable. Approaching AI as a material with agency requires developing a similar literacy, understanding training dynamics and

prompt interpretation well enough to engage skillfully with an inherently unpredictable partner. The goal is not perfect control but meaningful engagement with a system whose internal operations remain partially opaque. We can understand general principles while accepting that specific outcomes retain elements of useful surprise.

I could have tried to teach the model more about the physical

properties of clay, adding annotations about structural constraints or filtering out impossible forms from the training data. However, this would have optimized the model, something that from the start I wanted to resist. Instead, I kept working with its ignorance, treating the impossible structures it proposed as invitations rather than errors.

7.6 Conclusion

The ceramic residency explored whether AI can support material practice without optimizing it. The answer that emerged through three months of making is that it can, but only when the conditions are deliberately structured to resist the efficiency that AI makes possible. At every step where automation or speed was available, I chose to maintain friction, and that friction created conditions for reflection.

The circular methodology used generative AI in this context, but questions remain. Would the same insights have emerged through ceramic practice alone? Would a different material, like wood or textile, with different temporal demands, create different conditions for reflection? This chapter cannot fully answer these questions. What I can say is that the combination of subjective annotation, abstract prompting, and physical translation

created a structure through which my aesthetic vocabulary became visible to me, and that visibility would not have happened in the same way through making alone or through AI alone.

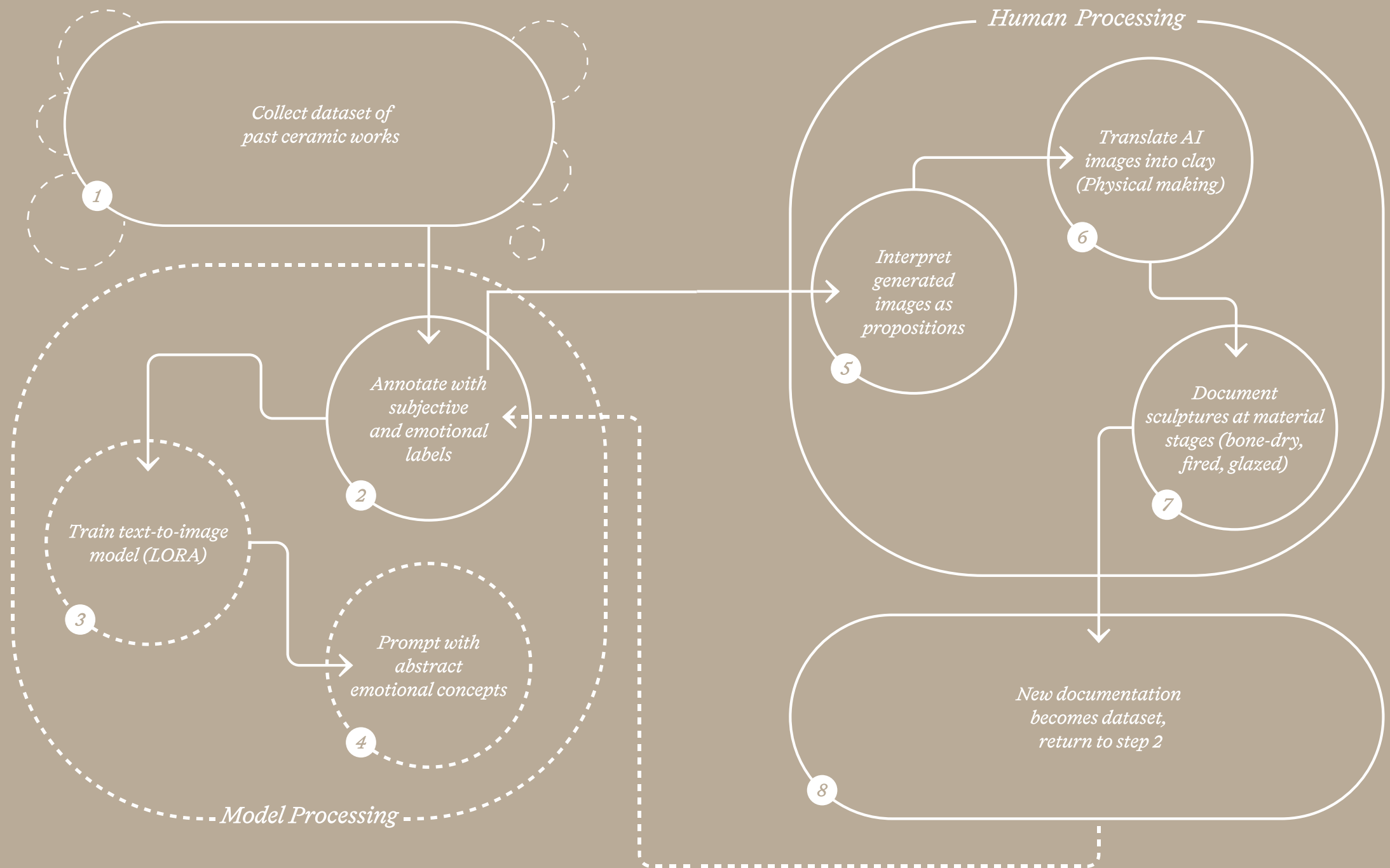
This chapter concludes the practice-based investigations that make up Part II. Across these five projects, from explorations with object detection to collaborative work with professional designers, from workshops with students to autoethnographic inquiry into my artistic past, and finally to this extended dialogue between AI and clay, I have been developing and testing what a Reflective AI practice might look like. The following chapter steps back from these individual inquiries to consider what they offer together, examining what can be formalized as insights for practitioners, what remains particular to these specific research contexts, and what this practice suggests about how we might relate to AI technologies differently.



Dutch Design Week 2025







↑ From Text-to-Clay

Part III: Synthesis

8. Discussion

The five projects documented in Part II traced the progression of a practice: from establishing a basic designer-AI dialogue with external objects (Chapter 3), to turning that dialogue inward toward design worlds and fascinations (Chapters 4 and 5), to excavating personal creative history (Chapter 6), and finally to grounding computational reflection in physical material practice (Chapter

7). These projects unfolded across different contexts (a design school, a ceramic residency), employed different practice-based research methods (a collaborative project with professional designers, workshops, autoethnography), and engaged different practitioners (design students, professional designers, myself as design researcher). Yet patterns emerged across them that I will synthesize in this chapter.

Through these projects, I developed a design-AI practice oriented toward reflection. I deliberately engaged with each stage of the AI pipeline as a site for self-examination, avoiding having the system do something I could do myself. This process did not automate my creative labor. Instead, it provided

a working method through which I could deepen my understanding of my own practice. When the goal of interaction shifts from generation to reflection, AI evolves from a closed tool into a malleable material for inquiry. This approach positions the technology's value in its capacity to externalize and refract the designer's implicit judgments, fascinations, and historical patterns. In each project, the locus of reflection shifted: mislabeling served as a reframing device, annotation forced tacit judgments into words, training a model on images of past work revealed why certain projects still held meaning years later, and making clay sculptures from AI-generated images required negotiating between what the AI proposed and what

the clay would allow. Across these variations, the practice consistently created conditions for questions.

The following section examines the character of reflection I observed across the projects, identifying distinct mechanisms that different configurations of the practice supported. I then formalize two design learnings for practitioners: a disentangled workflow and engagement with productive friction. The chapter concludes by considering what this practice offers design research and what it suggests about how we might relate to AI technologies differently.

8.1

The character of reflection

Before outlining specific learnings that emerged from this practice, it is worth clarifying what forms the

reflection that was observed took. As established in Chapter 2, Sengers et al. define reflection as bringing

Element	Chapter 3	Chapter 4	Chapter 5	Chapter 6	Chapter 7
Data source	External artworks and designed objects	Own fascinations	Research interests	Personal archive	Past ceramic work
AI type	Object detection (pre-trained)	Object detection (fine-tuned)	Object detection (fine-tuned)	Text-to-image (LoRA)	Text-to-image (LoRA)
Annotation	None	Subjective labels via Quadrant Tool	Subjective labels via Quadrant Tool	Subjective descriptions	Subjective and emotional labels
Output engagement	Physical modification of object	Object Portrait and interpretation	Object Portrait and interpretation	Interpretation across training epochs	Clay translation
Circular process	Yes (seeing-moving-seeing loop)	No	No	No	Yes (clay-data-clay-loop)
Physical making	Yes (modified chair)	Yes (Object Portrait)	Yes (Object Portrait)	No	Yes (ceramic sculptures)
Temporal orientation	Present	Present	Present and emerging	Past	Past to present
Character of reflection	Defamiliarization	Articulation of latent judgements	Articulation and anticipatory reflection	Autobiographical reflection	Translational reflection

unconscious aspects of experience to conscious awareness, making them available for conscious choice (2005, p. 50). Across the five projects, a consistent dynamic created conditions for this: the AI processed judgments and returned them transformed, producing a gap between intention and interpretation in which reflection became possible. But depending on the configuration of the practice, the data source, the type of AI process, whether outputs stayed digital or became physical, that reflection took different forms. Comparing the projects across their workflow elements (*Table 8.1*) showed five distinct types.

In the first project (Chapter 3), reflection took its simplest form: defamiliarization. The AI’s unexpected classifications made familiar artistic

and design work strange. The well known painting of Salvador Dalí suddenly appeared different when the algorithm detected ‘grape’ in its form. The insight here was methodological rather than deeply personal: seeing that AI’s mistakes could function as reframing devices, opening perception to aspects previously overlooked.

Articulation demanded more from the practitioner. Where defamiliarization happened to me, articulation required active effort: the requirement to annotate, to assign labels like ‘enticing’ versus ‘cloying’ or ‘majestic’ versus ‘repulsive’ to hundreds of images, forced vague intuitions into explicit terms (Chapters 4 and 5). This did not make tacit knowledge fully conscious; the underlying knowledge may remain inarticulable. But repeating that exercise across a dataset makes patterns visible, for example which labels kept shifting or which images were hard to categorize. The students in Chapter 5 discovered

↑ *Table 8.1*
Comparison of workflow elements across the five practice-based investigations, showing how the disentangled workflow manifested differently while sharing common structural elements and producing distinct forms of reflection.

their annotation strategies shifting over time, indicating that their initial definitions of labels or images were provisional and still forming.

A third type emerged alongside articulation in Chapter 5: anticipatory reflection. Here, the reflective process extended beyond the current moment. Students began viewing their creative process through the lens of how the AI would see their work in the future. Sasha's observation captures this: "*You're not just training an algorithm. You're also training your own gaze.*" The knowledge that their Object Portrait would be interpreted by their trained model shaped how they made it, making them consider how their design choices connected to the labels they had defined.

In Chapter 6, the temporal arc stretched further. Reflection took the form of autobiographical inquiry: using AI to explore my own artistic past rather than examine current practice. I trained a model on archival images of work made years earlier and used prompts to ask why certain projects still held meaning. This connects to Schön's reflection-on-action (Schön, 1983), the thinking that happens after doing, but extends it across a longer arc. Where Schön typically discusses reflecting on a recent design session, this mode concerns years of accumulated creative practice: discovering that a material I had used intuitively carried symbolic meaning as a form of consolation.

The most layered form emerged in Chapter 7: translation between media. I trained a model on my ceramic works using emotional labels,

generated images from prompts like 'ceramic jealousy,' and then tried to make those forms in clay. At times, the clay refused certain shapes the AI had proposed. Reflection emerged from figuring out what these computational forms could become when the material pushed back. This chapter also involved articulation (annotating ceramic works with emotional labels) and anticipation (imagining how forms would translate to clay), showing how the types can layer within a single project.

These different configurations share a common experiential thread: in each case, I encountered my own judgments made unfamiliar. Whether through unexpected classification, through the struggle to find words, through anticipating an algorithmic gaze, through temporal distance from past work, or through material resistance, the practice created distance from habitual ways of seeing. The AI functioned as an external structure that processed my judgments and returned them transformed, producing the gap in which reflection became possible. Yet the character of that reflection varied: whether I examined current fascinations or past archives, whether outputs stayed digital or became physical clay, whether I worked alone or with others.

This thesis operates on the premise that reflection serves design practice, but reflection is not always needed in the way I practiced it here. A skilled ceramicist throwing a familiar form does not constantly need to reflect in the way I have described;

they need to act. A designer working under a deadline may require speed rather than deep contemplation about their ways of seeing. I see this kind of reflection becoming valuable at moments of transition or uncertainty; for example in moments when

practice feels stagnant, or when encountering unfamiliar materials, or when certain assumptions within a project need examination. Reflective AI is proposed as a possible intervention for these moments, not as a permanent mode of working.

8.2

Learnings for a reflective design practice

To engage with Reflective AI is to adopt a specific material lens. With any material used for creative work comes a specific practice. Ceramic practice, for example, is a dialogue with the drying stages of clay, requiring specific 'ways of doing' to achieve a sculpture. Working with wood creates a practice of understanding grain and density to ensure the material responds effectively. For Reflective AI, I see a similar practice evolve: one that views AI models as having multiple entry points for engagement, rather than a single input-output transaction, that can be shaped through a negotiation between the desires of the designer and the agency of the material itself. In this case, the material is the AI training pipeline: the dataset, the annotation, the training process, and the outputs. Each of these stages has its own properties and resistances, just as clay has moisture and wood has grain. Reflective AI establishes a craft-dialogue with the constituent stages of machine learning, engaging

directly with its underlying technical structure.

This material analogy proved to be helpful, but it has limits. AI is a material in some ways, but in others it is not. Working with craft materials is embodied; your hands are in the clay, shaping it directly. Working with Reflective AI is, in the context of this thesis, visually mediated, requiring interpretation of outputs on a screen. When I work with clay, I feel the moisture level in my hands and know immediately if the material will crack. When I trained the ceramic AI model, I waited hours for training to complete, then looked at generated images trying to understand why certain forms appeared. I cannot feel the model's properties; I have to interpret them through outputs. In Chapter 4, when Borisz discovered his model labeled both 'majestic' animal legs and 'repulsive' satyr legs identically, he had to reason backward to understand why: the visual features were the same, even if his intentions differed. This

interpretive work, reading outputs, tracing them to annotation decisions, adjusting strategies for the next round, is a form of material dialogue, but it happens through visual interpretation rather than touch. What holds is the relational structure: iteration, resistance, developing understanding through repeated engagement with the material. It is this relational structure that makes a material lens productive for AI in design practice, and that the following learnings aim to support.

Even though objects—such as the clay sculptures or object portraits—emerged from this practice, they were not the primary objective of the work. The goal was to deepen their understanding of their own practice as a designer while engaging with the system and the materials in the making. In doing so, agency shifts: the designer moves from consumer of black-box capabilities to architect of the system's logic, retaining authority by curating the worldview through which the algorithm sees.

To use this AI-as-material lens appropriately, I formalized two design

learnings. These are not to be read and adopted as a rigid, prescriptive guide, but as a loose framework for practitioners—whether artists or designers—seeking to move beyond prompt based tools. They offer an alternative way of thinking about what AI is and how it can be used as a medium for reflection, demonstrating that possibilities exist beyond the predominant text-to-image paradigm. The learning is: the deployment of a disentangled workflow and the engagement with productive friction.

8.2.1 *Learning 1: A disentangled workflow*

The first learning is a disentangled workflow, or in other words, a loose guide on how to approach working with Reflective AI and its constituent parts as reflective spaces. Where standard interaction formats abstract the complexities of AI into a seamless interface, a Reflective AI practice requires the opposite; it deliberately disentangles the tool into its constituent counterparts, each coming with their own 'way of doing'.

1. *Data Curation as Inquiry*

The act of collecting data serves as the first reflective space. Rather than scraping the internet for large training datasets, the designer engages with training data as something that visually represents an aspect of their practice. For example, the foam chair dataset was a personal archive that represented a specific part of my history. A dataset of abandoned buildings represented Yann's research interest for his semester. In addition, the materialization of these datasets (e.g. printing them out, arranging them spatially) transforms the abstract folder of images into a tangible body of work to engage with.

2. *Annotation as Articulation:*

Translating vague feelings into textual labels forces the designer to articulate the 'why' behind intuitive choices. The struggle to define 'bull-like' cups or distinguish 'virus' from 'misfit' (Chapter 7), or to separate 'otherworldly' from 'convivial' (Chapter 5), moves the designer from feeling into explicit concept. The quadrant tool used in Chapters 4 and 5 offers an open yet guided framework to organize unarticulated feelings as a basis for annotation. The workshop in Chapter 5 demonstrated that annotation generates distinct strategies: some students, like Borisz, adopted a consistency strategy by adhering to initial definitions; others, like Yann, employed a revision strategy, constantly updating labels to match their evolving understanding. Annotation becomes a creative technique where the designer actively defines the logic of the system while becoming aware of their 'ways of seeing'—the implicit preferences and patterns of attention that shape how they perceive and judge visual material (Berger, 1972).

3. *Inference as Dialogue*

Once the model is trained, the act of using it shifts from issuing commands to posing questions: from *Give me the right answer* to *What did the model learn about my data?* In Chapter 6, text-prompts became a form of personal inquiry, testing how the model reconstructed memories and what this revealed about my material choices. Similarly, in Chapters 4 and 5, the Object Detection outputs, regardless of confidence score, served as prompts for reflecting on our work rather than verdicts on it.

4. *Output as Propositions*

Treat AI output as a proposition. A proposition is offered for consideration; it can be accepted, questioned, built upon, or rejected. In every chapter, outputs became invitations for a next step: a surprising mislabeling (Chapter 3) offered a new way of looking at familiar work, Object Detection results (Chapters 4 and 5) prompted questions about annotation choices, and generated images (Chapter 7) became the basis for ceramic sculptures.

5. *Making as Anchor*

Crucially, Reflective AI relies on an act of physical or digital making outside the model. The AI output is never the final artifact. Whether it is the physical or digital construction of an Object Portrait, the hand-building of ceramic sculptures, or the analysis of a previously made foam chair, the practice relies on a 'making response' to ground the algorithmic insights.

6. Non-linearity

This workflow is not necessarily linear; it can function as a loop where the ‘making response’ often generates new data, restarting the cycle of curation and annotation with deepened insight (Chapter 7).

8.2.2 Learning 2: Engagement with failure and ambiguity

The second learning reads more as a perspective shift, as it entails a revaluation of error within AI systems in the design process. Training models according to the disentangled workflow described above frequently results in technical failure by industry standards, as small, subjective datasets and inconsistent annotation lead to underrepresentation of certain concepts, hallucinations, or mislabeling. A Reflective AI practice reinterprets these failures through the lens of ambiguity. Gaver et al. (2003) describe the generative potential of this state: “*Ambiguity can be frustrating, to be sure. But it can also be intriguing, mysterious, and delightful. By impelling people to interpret situations for themselves, it encourages them to start grappling conceptually with systems and their contexts, and thus to establish deeper and more personal relations with the meanings offered by those systems*” (p. 240).

In this context, the AI model contributes to ambiguity by offering a form of false certainty. It presents probabilistic guesses as definitive images (*ceramic surrealism looks like this*) or labels (*57% stubborn*), creating an interpretative gap that compels the designer to intervene in a particular way. As noted by Giaccardi et al. (2024) in their paper *Prototyping with Uncertainties*, this

shift in perspective is crucial for Research through Design (RtD) and practice-based research practices engaging with algorithms. They argue that a significant tension lies in “*viewing uncertainty as a problem to solve versus a creative impetus or a source of reflexivity*” (p. 68:10). By choosing the latter, the reflective practitioner transforms the unstable, probabilistic nature of the model into a driver for design inquiry.

I see this tendency operating as a core aspect across all the thesis projects. In Chapter 3, the mislabeled chair provided an alternative frame that forced a reconsideration of the object’s form. In Chapter 6, the model’s failure to generate my own face, likely due to underrepresentation in the training data, created a figure dissolving into foam, which sparked a realization about the relationship between myself and the medium of foam. Similarly, the structural impossibilities of the AI-generated ceramics in Chapter 7 functioned as generative challenges that pushed my manual skill into new making territories and realizations about my own practice and the material itself.

Engaging with these limitations requires a methodological commitment to Donna Haraway’s concept of ‘*staying with the trouble*’ (Haraway, 2016). Haraway argues against the impulse of technological solutionism, the desire to immediately fix or smooth

over complex problems. Instead, she proposes dwelling in complexity as a form of accountability. Similarly, in a Reflective AI practice, I refuse to ‘fix’ the model to achieve perfect accuracy. I choose to dwell in the difficulty of the ‘bad’ model because it reveals the seams of the system, and this is the very site where reflection occurs. In Chapter 4, the annotation drift observed in Gijs was similarly a form of productive instability: it documented how his aesthetic judgment evolved over time rather than representing a failure of consistency to be corrected. These moments of friction became the generative core of the practice: fluid human judgment meeting the model’s demand for consistency, and what one feels meeting what one thinks one should label.

Although developed within more artistic oriented design contexts, Reflective AI offers a framework for any practitioner whose work is driven by specific fascinations, materials, or interests. The disentangled workflow functions as a flexible

scaffold adaptable to distinct design domains, provided the primary objective remains reflection rather than production. This method serves not merely as a tool for generating new artifacts, but as a means to examine the self as a maker through making. While artifacts, such as the foam chair, ceramic sculptures, or object portraits, are integral to the process, they function here as objects for inquiry rather than primary goals. Where I explored the relationship between emotional concepts and ceramic materiality, an architect might curate datasets to investigate their latent formal obsessions, while a graphic designer could use annotation to surface the implicit values underlying their typographic interests. By engaging with the reflective spaces created by this disentangled workflow, designers across disciplines can resist the homogenizing logic of the algorithm, instead using the ambiguity of the model to deepen the understanding of their own creative practice.

8.3

Between particular insights and general method

The insights documented in this thesis are personal: the recognition that foam functions as consolation, the shifting taxonomy of Gijs’s garden from ‘*lonely*’ to ‘*luring*’ (Chapter 4), the confrontation with ‘ceramic

ignorance’ in the clay studio (Chapter 7). These findings emerge from my particular history, materials, and contexts. Throughout this work, I found myself returning to a question: if my findings about foam and ceramic

jealousy are tied to my specific memories, what remains for other researchers to build upon?

Here I find it useful to think through the logic of machine learning itself. A neural network architecture is a generalizable structure that can be implemented across applications. The trained model, however, is unique, its weights shaped by the specific data it was trained on. Reflective AI can be interpreted using through this same distinction. The specific outcomes documented in this thesis are the unique ‘weights’ of a personal model, shaped by individual histories and materials. What this research offers is the ‘architecture’: the disentangled workflow, the learnings of subjective annotation, the strategic deployment of ambiguity and friction. The structure is teachable but the insights it generates are not.

This framing finds support in Donna Haraway’s work on situated knowledges. Situated knowledge is knowledge that acknowledges where it comes from, who produced it, and under what conditions. Haraway challenges the view from nowhere that claims objectivity by erasing the position of the knower, arguing instead that “*the only way to find a larger vision is to be somewhere in particular*” (Haraway, 1988, p. 196). For Haraway, situatedness enables knowledge production, creating “*partial, locatable, critical knowledges*” that are accountable precisely because they acknowledge their position (1988, p. 191). The insights documented here are situated in my particular history

and contexts; their specificity is what makes them accountable.

Beyond this methodology, something else transfers through practice: a developing sensitivity to AI as a material. Just as ceramicists develop an embodied understanding of clay through repeated engagement, practitioners who work through these reflective cycles can develop a feel for how annotation choices manifest in model behavior, how training data shapes output, and where productive friction emerges. I cannot claim that others following this workflow will arrive at the same findings, or that my findings would replicate if I repeated the process. What I observed in the Object Portrait Workshop and EKWC residency is that both the method and this developing attunement can transfer across different contexts and timescales; the documentation of both processes makes the methodology visible and accountable.

What distinguishes Reflective AI from other first-person design research approaches is where reflection is located. Where autobiographical design relies on extended self-usage and autoethnography on reflexive writing, Reflective AI embeds reflection into the technical structure of machine learning itself. The repeated requirement to annotate transforms vague intuitions into applied categories. The struggle to define ‘*commercial saturation*’ or to categorize which ceramic forms embody ‘divorce’ required exercising aesthetic judgments across hundreds of instances. The accumulated annotations become a record of these judgments that can then

be examined. This does not guarantee insight, but it creates conditions where implicit judgments become visible.

A recurring question I encountered during my projects was whether the reflections I documented are really attributable to AI or whether therapy or journaling might have produced similar insight. In response to this I can say that I do not claim Reflective AI uniquely enables insights unavailable through any other means. What distinguishes this approach is that annotation structurally requires the articulation of implicit judgments, creating systematic conditions for reflection rather than relying on introspective discipline alone.

I see Reflective AI as inspirational for design research investigating

creative practice, aesthetic development over time, or alternatives to efficiency-focused framings of AI. It is less appropriate when researchers seek statistically generalizable findings or when rapid iteration matters more than long-term engagement. For design researchers, this work offers a teachable workflow for generating personal insights, documented examples of that workflow in different contexts, and a framework for understanding what transfers (the structure) and what remains particular (the findings). The rigor lies not in the universality of insights generated, but in the systematic conditions through which they emerged.

8.4 *Implications for AI technology and discourse*

Beyond the methodology and its transferable elements, this research proposes a shift in how we relate to the tools we build and use. Current commercial narratives frame AI as an engine for efficiency, designed to bypass the messy, durational stages of creation. My practice suggests that value lies elsewhere. In the studio, the utility of the AI model appeared when it failed to be efficient: when it categorized something as ‘76% stubborn’, when it ‘misunderstood’ a ceramic form, or forced me to pause

and annotate a feeling. Utility, in this framing, is measured by what the interaction opens up for the practitioner. The following sections explore what this shift might mean for those who design AI tools, and for the broader stories we tell about AI.

8.4.1 *Designing AI for inquiry*
Although the goal of this thesis was not to build a tool or develop a technology itself, I can imagine what this research can mean for those who design the tools we use. From the

insights generated, I would reorient the hypothetical Reflective AI tool from one that provides answers to one that generates questions. This is not to say that designers using current tools are unreflective. Any tool can be used reflectively or unreflectively, depending on the practitioner's stance. The point is about what the interaction format of these tools structurally encourage. We can imagine AI tools designed not to smooth over ambiguity and friction, but to expose it.

While interfaces exist to fine-tune models on one's own style, such as exactly.ai (Exactly.ai, 2024) or personalization features in Midjourney, that resemble the fine-tuning method I used in this thesis, these initiatives remain focused on optimizing for output generation. They often miss the opportunity to highlight the reflective potential of the training process itself. This research demonstrates that a different engagement with AI is possible, one that complicates the dominant either/or framing of AI. To illustrate how the process can become part of the tool, I turn to critical art practices that prioritize exploration over optimization in their relationship with AI technologies. These artists subvert the technology to reveal its underlying logic as a space for learning. These interventions offer alternative modes of interaction, demonstrating how Reflective AI might function similarly as a potential tool for inquiry.

James Bridle's *Autonomous Trap 001* provides a clear parallel to the approach taken in this thesis. This

artwork is part of a larger artistic research practice in which Bridle trained a self-driving car to 'get lost' rather than find the optimal route, using the machine's confusion as a way to map physical space (Bridle, 2018, 2017). My work with the 'disentangled workflow' applies this logic to the creative process. Just as Bridle's car maps the territory through random walks, or drives, the Reflective AI process maps the designer's internal 'design world' by navigating the probabilistic uncertainty of the model. The goal is to get lost in the data to understand the structure of one's own creative practice.

This requires a willingness to stay with the 'bad' model. Eryk Salvaggio, critical AI artist and scholar, describes this as 'creative misuse': Deliberately engaging with the glitches of a system to see what they reveal about its training (Salvaggio 2024). In my ceramic practice, the 'bad' model that failed to generate a structural sculpture was a mirror for my own implicit understanding of gravity and clay, rather than a technical failure. Similarly, Soyun Park's *Autonomy of a Fall* uses the struggle of an algorithm to categorize a falling body to investigate the limits of classification (Park et al. 2025), as it can only categorize standing up and lying down. The transitional act of falling down becomes an uncategorizable 'in-between space'. A tool designed for reflection would center this struggle. It would visualize the difficulty of defining 'jealousy' or 'loneliness,' making the gap between

human complexity and machine logic a tangible feature of the interface.

Finally, this perspective reclaims the training process as a creative act in itself. Holly Herndon and Mat Dryhurst's recent work *The Call* (2024) offers another model for rethinking AI as coordination and communication technology rather than replacement (Herndon & Dryhurst, 2024). Working with fifteen community choirs across the UK, they composed hymns and singing exercises to train choral AI models. Their approach treats AI much like singing has functioned for millennia: as a ritual for gathering, processing information, and creating collective meaning. This resonates with my experience of data curation. The hours spent categorizing garden styles, or annotating surrealist art datasets with fascinations, were not just 'prep work' for an algorithm; they were a ritual of organizing a personal archive and reflective practice. Future tools for creative practice could prioritize this archival relationship, treating the dataset as a body of work to be tended, annotated, and understood.

8.4.2 Stories we tell about AI

This research demonstrates that a different engagement with AI is possible, one that complicates the dominant either/or framing of AI. Beyond a studio practice, beyond the design research space, and beyond those who design our tools, this work suggests a necessary change in the stories we tell about artificial intelligence in general. Public discourse, as how I have been encountering

it as this work progressed over the years, is currently paralyzed by what Bory et al. calls 'strong AI narratives': Grand stories of sci-fi superintelligence that is either saving or destroying humanity (Bory et al., 2025). These narratives suggest a binary choice: we must either reject the technology or adopt it and accept it fully. This results in the simplified questions that dominate headlines: 'Is AI good or bad?', 'Friend or foe?', 'Intelligent/Creative or not?'.

My research aligns with the 'weak AI narratives' described by Chubb et al. (2022), which focus on the messy, situated reality of specific systems in specific contexts. The 'Slow AI' framework proposed by the AIxDesign (2024) collective supports this view, arguing for a pace of engagement that respects natural and cultural time rather than Silicon Valley speed. In my practice, this meant rejecting the always-on speed of the generator in favor of the slow drying time of ceramics. It shows that we can engage with the technical affordances of machine learning without subscribing to the accelerationist culture that surrounds them.

Crucially, in the projects presented in this thesis, the AI model never functioned to automate a practice that I could already perform. It did not replace sketching to save time, nor did it classify images to spare me the effort. If we remove the lens of automation and efficiency, are those binary questions still relevant to ask?

This non-binary approach parallels with James Bridle's argument in *Ways of Being*. Bridle suggests we should

look to non-human intelligences, whether biological or artificial, not to ask if they mimic us, but to understand what distinct forms of cognition they offer (Bridle, 2022). By viewing AI as a distinct, ‘other’ form of intelligence, we are liberated from the need to compare its output to human standards of creativity. We stop asking if the tool is ‘good enough’ to replace us, and start asking what its perspective reveals about ourselves.

Rather than debating whether AI is creative or not creative, conscious or not conscious, beneficial or harmful in absolute terms, I demonstrate how particular configurations of AI practice can facilitate particular forms of self-knowledge under specific conditions. The question shifts from what AI can do to what long-term engagement with AI training can help me understand about my own practice, values, and assumptions.

8.5

Limitations and conditions

Reflecting honestly on the projects in this thesis, there are aspects I would approach differently. The most intensive projects in Chapters 6 and 7 relied solely on my own reflective capacity. While first-person methods are central to this research, a second practitioner engaging with the same workflow using different materials could have offered comparative perspective on whether the structure generates similar reflective dynamics for someone else. I also did not pursue longitudinal follow-up, as each project was conducted once without returning to re-annotate datasets months later to examine how my judgments had shifted. The durability of the insights, and of the reflective capacity the method aims to develop, remains untested. Additionally, this research was conducted by someone already invested in and predisposed toward reflective practice, which raises the

question of whether designers who do not naturally gravitate toward introspection would find similar value in this approach or would experience it as an obstacle to their work.

Beyond these methodological limitations, Reflective AI functions within specific conditions that constrain its applicability. The practice requires time and space, which I had during the course of this PhD. Curating a personal dataset, annotating hundreds of images, training a model, and responding to outputs through making is slow work that extends across days or weeks. For a commercial designer facing a twenty-four hour deadline, the suggestion to curate a personal dataset or manually annotate hundreds of images is impractical. The efficiency-oriented tools I critique exist precisely because they solve problems of production speed, which positions Reflective AI as

a pedagogical or strategic intervention rather than a replacement for commercial pipelines. It functions best in the fuzzy front end of design, in education, or in contexts where the goal is defining the problem rather than finalizing the solution.

The method also requires technical confidence. Throughout the chapters, I used my experience with coding to navigate hurdles like setting up local environments, managing LoRA training, and debugging scripts. The productive friction I describe might function as blocking friction for designers who lack interest in navigating the technical infrastructure of terminals, scripts, and local installations that the practice currently requires as an entry point. This suggests that Reflective AI currently requires significant scaffolding for designers fully unfamiliar with machine learning.

The examples in this thesis also rely on translation between physical matter and digital image, which raises questions about how this framework might adapt to immaterial design disciplines such as UX or service

design. This remains an open and exciting challenge for future research. The method further assumes a desire for introspection and a willingness to confront one’s own patterns and past work. For practitioners who view design strictly as problem-solving for others, this inward turn might feel irrelevant, and the vulnerability it demands is not something all professional contexts support.

Some of the research projects support this concern. In Chapter 4, Maurik approached the project more independently and found the annotation process tedious rather than generative. Similarly, some students in Chapter 5 potentially did not reflect deeply, something we might have missed in the data collection, where quieter struggles with the method remained less visible. This suggests that the method does not always guarantee reflection; it creates conditions that practitioners must actively inhabit. Only completing the steps is not sufficient, but the willingness to dwell with the process, to find the annotation meaningful rather than mechanical, appears to matter.

8.6

Conclusion

This thesis began with a question about perception: how do we see and give meaning to the world around us? What started as a personal fascination with the mechanisms of looking evolved, across six years of practice, into a

design practice that positions AI as a medium for reflection.

Through five practice-based investigations, I have shown that the technical pipeline of machine learning, when approached slowly

and deliberately, can become a site where designers examine their own ways of seeing. Gathering images that represent a fascination forces you to ask what actually belongs there; labeling those images with words like ‘*uncanny*’ or ‘*commercial saturation*’ demands that vague feelings become concrete choices; waiting hours for a model to train gives time to wonder what it will have learned. And the model’s outputs, whether surprising detections or unexpected generations, function as material that talks back, inviting reconsideration of one’s own curatorial choices.

The contribution of this thesis is methodological rather than technological: instead of offering a new AI tool or algorithm, it presents a documented practice, two design learnings (the disentangled workflow and engagement with ambiguity), and a framework for understanding what transfers across other design-related contexts. The specific insights that emerged, for example that foam means consolation or that ceramic jealousy takes the form of sharp proliferating edges, remain particular to the practitioners who discovered them. What can be shared is the structure through which such discoveries become possible.

Returning to where this research began: I entered design school unable to make with my hands but obsessed with understanding why we are drawn to certain forms and materials without knowing why. I thought AI might help me answer this question by showing me how machines see. What I discovered instead was that

training AI on my own subjective data became a way of examining my own seeing, allowing me to reflect on myself while making new things like an Object Portrait and ceramic sculptures. The technology that promised to automate creative labor offered something else entirely: a slow, sometimes frustrating, often surprising practice of self-examination through making.

Much remains open. This research focused on image-based AI systems; how Reflective AI practice might extend to text, audio, or other modalities is unexplored. The projects documented here involved individual practitioners or small groups; what happens when larger communities train models together, layering multiple subjectivities, presents new questions. And as AI tools continue to evolve, the specific techniques documented here will require adaptation, even as the underlying principles of slowness, material engagement, and first-person inquiry may remain relevant.

The dominant narrative surrounding AI in creative practice will likely continue to emphasize speed, scale, and automation. This work suggests that other approaches are possible, though how widely they apply remains to be seen. The value, at least in the contexts documented here, lies not in what the model produces, but in what the process of training it reveals about the one who trains it.

I began this research as a designer who struggled to make, and surprisingly enough, I ended it having made sculptures, trained models, taught workshops, and written a thesis. The practice I developed taught me to pay

attention to what I was already doing, to notice patterns in my thinking and making, and, to understand why certain forms and materials kept

calling me back. That, in the end, is what reflection offers: not definite answers or quick solutions, but clearer questions about who we are as makers.

8.7

Statement on AI Use in Manuscript Preparation

This thesis was prepared with assistance from Claude (Anthropic). Following the STM Association's classification of AI use in academic manuscript preparation (STM Association 2025), the tool was used in several ways. Claude suggested language improvements and edits throughout the manuscript, all of which were reviewed and revised by the author. It was also used to generate draft text for portions of the thesis, which was then edited substantially. The Dutch summary was initially translated with Claude's assistance, then reviewed and edited by a native Dutch speaker. Claude assisted with formatting and

generating tables, and occasionally helped identify potentially relevant literature in the form of a search engine, though all references were verified, read, and assessed for relevance by the author.

Beyond these categories, Claude served as a dialogic partner throughout the writing process: a tool for thinking through arguments, testing formulations, finding counterarguments, identifying redundancies across chapters, and reflecting on the structure and clarity of the work. Claude also helped compare earlier and later versions of chapters to track what had changed.

9. Acknowledgements

Much of the work in this thesis has been shaped by the great support of others. What began as an individual research project became, over six years, something closer to a collective practice, carried by the guidance, patience, and care of the people that I would like to thank and acknowledge.

I want to begin with my supervisors Peter Lloyd, Almila Akdag Salah, and Senthil Chandrasegaran. As a team, we built bridges together, and found ways to make our collaboration productive, critical, and encouraging. I am very grateful for receiving all the space and support to do the work I wanted to do.

To my committee members Gerd Kortuem, Kristina Andersen, Olya Kudina, Camilla Groth and William Odom for reading this thesis critically and asking such challenging and productive questions about it.

To my paronyms and dear friends: Iohanna Nicenboim, my colleague in work and life. You made me believe in academia, myself and in love. And Esther van der Heijden, my art sister, my friend, who has the incredible skill to stay positive and see the good in everything, which is so inspiring. I am so happy to have you both by my side.

To my Design Academy steadies: Gijs de Boer, for co-authoring some of this work, for being the art voice, and being such a special friend to me throughout my career, Maurik Stomps, for being an unmissable and well-appreciated research participant and Colin Keys who was there along the way as a great mental support.

I want to thank all the students that participated in the Objective Portrait Workshop, the studio leaders Rhiarna Dhaliwal, Emmy Bacharach and Ibiye Camp for offering us the space to conduct this research within their Bachelor program, and the Design Academy Eindhoven for hosting and facilitating the workshop.

My gratitude goes to Baptiste Caramiaux for his support and for hosting me twice at ISIR Sorbonne in Paris, supervising a project that became an important backbone of this thesis, and to Lenny Martinez and Behnoosh Mohammadzadeh for making time so much fun in the pyramid.

To my fellow colleagues from academia: Willem van der Maden, with whom I roadtripped through Pennsylvania after all the flights were cancelled; Grace Turtle, my partner in crime in the borderland of TU Delft; Brett Halperin, my support from overseas and trusted

conference companion; and Jesse Nijdam, who, even though his path has changed, helped me find my direction when this work was just beginning. To my fellow PhD council members Wo Meijer and Caiseal Beardow: thank you for the shared commitment to making PhD life a little better for those who come after us.

Jesse Josua Benjamin, for his unstoppable energy, inspiring research practice, and deep belief in the work I do. Having you on board made all the difference. My studio partners at Post Neon who became friends for life. Jeremy, Niels, and Jim provided me with a workspace in Amsterdam, good company, and kept me with both legs in practice while my head was in academia.

To everyone at EKWC: Geertje, Katrien, Pierluigi, Stef, Leonie, Isabelle, Chant, and the many others I shared studio life with during my time at this magical place. Thank you for welcoming me into the world of clay. And to Laura Hoek for being such an amazing and positive assistant in the studio.

Ward Goes, who created the beautiful design of this book and supported me along the way with the exhibition for From Text-to-Clay, and started a new work chapter with me, studio-partner wise.

To my oldest friends Gijsje, Sozan, and Julia, you have known me longest and loved me through all of the turmoil of growing up and doing life. And to dear David, who gave me his heart and a home for most of the PhD.

Then I want to thank my mother Suzanne, who has believed in Reflective AI since day one, who always saw the red thread in my work when I was lost. A valued second reader, sparring partner and life coach, and whose ability to be moved by ideas has always moved me. Then, my brave sister Sophie and sweet father Maarten, who may not have always understood what I was doing, but never doubted that I should be doing it.

And finally, Micha, my great love, of whom I sometimes wish I met earlier, but I think I met you just in time. You have given me a sense of home in a way I did not know before, but always wanted. You have read every chapter, and somehow managed to find my work as fascinating as I find you(rs).

10. About the Author

Vera van der Burg was born on June 28, 1993, and grew up in Amsterdam, the Netherlands. She began her academic trajectory with a Bachelor's degree in Liberal Arts & Sciences with a major in Neuroscience at Utrecht University, completing her studies in 2017. She then pursued a Master's degree in Contextual Design at the Design Academy Eindhoven, graduating cum laude in 2019. In 2020, she started her PhD at Delft University of Technology in the Industrial Design Engineering department, as part of the Designing Intelligence Lab under the supervision of Peter Lloyd, Almila Akdag Salah, and Senthil Chandrasegaran, in which she developed the Reflective AI design practice described in this thesis.

During her PhD, van der Burg completed several residencies that offered opportunities to dive deep into specific research projects and contexts: as Researcher-in-Residence at the Nieuwe Instituut in Rotterdam in 2024, as Visiting Researcher at the ISIR Lab at Sorbonne Université in Paris working with Baptiste Caramiaux in 2024, and as Artist-in-Residence at the European Ceramic Work Centre in Oisterwijk in 2025. Throughout this period, she exhibited her work at venues including Dutch Design Week in Eindhoven, Salone del Mobile in Milan, the AIXDESIGN Festival in Amsterdam, and Noorderlicht Photography in Groningen. In 2025, she received the Dutch Design Award for Emerging Talent in recognition of her practice.

In her academic work, she has served as Associate Chair for the Pictorials track at the Designing Interactive Systems (DIS) conference in 2025. She has published at conferences including DIS, DRS, and CHI, and contributed chapters to books including *Algorithmic Imaginations* (ArtEZ Press, 2025) and *Dixit Algorizmi: The Garden of Knowledge* (Humboldt Books, 2022). Presenting this work at academic conferences and public venues such as the Swiss Design Network Symposium, World Press Photo, and V2-Lab for the Unstable Media contributed to the development of the ideas in this thesis. At TU Delft, she has been involved in teaching courses such as Designing Intelligence and Design for Interaction, and has supervised multiple graduation students working on AI-related projects. Additionally, she developed workshops for design students and design educators at the Design Academy Eindhoven, and for design practitioners at Google Studios in collaboration with ProtoÉditions and It's Nice That, exploring how to work with AI in creative contexts.

● Publications

○ First Author Papers

van der Burg, V., de Boer, G., Benjamin, J. J., Halperin, B. A., Akdag Salah, A., Chandrasegaran, S., & Lloyd, P. (2026). Reflective AI: A Slow Technology Approach for Design Education. *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems (CHI '26)*. Barcelona, Spain. ACM. Chapter 5 of this thesis was based on this publication.

van der Burg, V., de Boer, G., Akdag Salah, A., Chandrasegaran, S., & Lloyd, P. (2023). Objective Portrait: A practice-based inquiry to explore AI as a reflective design partner. In *Proceedings of the ACM Designing Interactive Systems Conference (DIS '23)*, Pittsburgh, USA. ACM, pp. 387–400. Chapter 4 of this thesis was based on this publication.

van der Burg, V., Akdag Salah, A. A., Chandrasegaran, R. S. K., & Lloyd, P. A. (2022). Ceci n'est pas une chaise: Emerging practices in AI-designer collaboration. In *Proceedings of Design Research Society International Conference (DRS 2022)*, Bilbao. Chapter 3 of this thesis was based on this publication.

○ Other Publications

Arzberger, A., van der Burg, V., Chandrasegaran, S., & Lloyd, P. (2022). Using human-AI dialogue for problem understanding in collaborative design. In *34th International Conference on Design Theory and Methodology (DTM)*. ASME.

Beerends, S., van der Burg, V., Crespo, S., Goeting, M., van Neck, M., Park, S., Ridler, A., Rosado, P., Scheibe, N., Schipper, C., Schmitt, P., Völp, L., & Winters, S. (2025). *Algorithmic Imaginations: Critical Reflections on AI in Art and Design Practice*. ArtEZ Press. ISBN 9789491444838.

van der Burg, V. (2022). The Incredible Story of Algo & I. In S. Ghomashchi & Space Caviar (Eds.), *Dixit Algorizmi: The Garden of Knowledge*. Humboldt Books.

van der Maden, W., van Beek, E., van der Burg, V., Halperin, B. A., Jääskeläinen, P., Kang, E., Kun, P., Lomas, J. D., Nicenboim, I. et al. (2024). Death of the Design Researcher? Creating Knowledge Resources for Designers Using Generative AI. In *DIS '24 Companion*. ACM, pp. 396–400.

van der Maden, W., van Beek, E., Nicenboim, I., van der Burg, V., Kun, P., Lomas, J. D., & Kang, E. (2023). Towards a Design (Research) Framework with Generative AI. In *DIS 2023 Companion*. ACM, pp. 107–109.

● Exhibitions

○ From Text-to-Clay – related to Chapter 7

2027

↳ Gregg Museum of Art and Design, Raleigh, USA

2026

↳ CTRL Matter, Mint Gallery, London, UK

↳ How Matter Comes to Matter, Onomatopee, Eindhoven, NL

2025

↳ 4TU.Design United, Klokgebouw, Dutch Design Week, Eindhoven

↳ Entrance Hall, Klokgebouw, Dutch Design Week, Eindhoven

○ Objective Portrait (with Gijs de Boer) – related to Chapter 4

2025

↳ AIxDESIGN Festival: On Slow AI, Loods 6, Amsterdam

2024

↳ Pixel Perceptions: Into the Eye of AI, Noorderlicht Photography, A-Kerk, Groningen

↳ Highlight Delft, Delft

2022

↳ Dutch Design Week, Eindhoven

○ Still Life – graduation work that preceded this thesis

2021

↳ Intergenerational Graduation Show, Salone del Mobile, Milan

2019

↳ Graduation Show, Design Academy Eindhoven

○ Media Coverage

2026

↳ Frame Magazine, 'One's to watch'

↳ Disegno, 'Text-to-Clay'

↳ Quest magazine, 'Oude beroepen vinden nieuwe toepassingen'

- ↳ Dezeen, “Vera van der Burg uses AI to create ceramics in Text-to-Clay project”
 - ↳ Het Financieele Dagblad, “Vera van der Burg wil met AI een ruimere blik op de wereld bieden” (50 Talents)
 - ↳ BNO Dutch Designers Yearbook
- 2025
- ↳ DAMN Magazine 91, feature
 - ↳ Wallpaper*, “Dutch Design Awards 2025 honour a new generation of creatives”
 - ↳ Dezeen, “Five unexpected themes that emerged at Dutch Design Week 2025”
 - ↳ LSN:global, “Dutch Design Week 2025: Slowing AI and fostering empathy”
 - ↳ TU Delft Design Stories, “Crafting reflection through AI and clay”
- 2024
- ↳ Ontwerp Platform Arnhem (OPA), podcast interview

11. Summaries

● *Summary in English*

The integration of Artificial Intelligence into creative practice is increasingly shaping how designers work. Commercial AI tools promise to make creative work faster and easier, generating images from text prompts and transforming sketches into polished outputs in moments. While this efficiency-driven approach aligns with productivity goals, it also constrains how designers engage with AI in their creative work. The polished interfaces of these tools conceal critical stages such as data collection, annotation, and model training, leaving designers to interact only with surface-level outcomes. This abstraction, combined with the speed of generation, prevents designers from developing what Donald Schön would recognize as a conversation with materials—the reflective dialogue a designer has during a design process that makes design practice meaningful.

This dissertation develops Reflective AI: a design practice that treats AI as a material for reflection rather than a tool for automation. The core insight emerged from the researcher’s own practice: the technical pipeline of AI—data collection, annotation, model training—can become the very material through which designers engage in reflective practice. Rather than using AI to generate outputs efficiently, Reflective AI positions these technologies as mirrors

that can reveal the designer’s own ways of seeing, their implicit preferences, and their tacit judgments.

The research follows Practice-Based Research, where questions arise from the process of practice and answers are directed toward enlightening and enhancing that practice. This means the practices described were not predetermined and then tested, but emerged through the act of making: encountering AI as a material in different contexts and reflecting on what those encounters revealed. The investigations span solo explorations, collaborative work with professional designers, a workshop with design students, auto-ethnographic inquiry into artistic history, and a three-month ceramic residency.

Through this practice, one overarching research question frames the investigation: *What specific design practices enable reflective engagement with AI?* This question is investigated through two interconnected sub-questions: First, what characterizes reflection when designers engage with AI training processes? Second, what design learnings emerge from this practice that can transfer to other design research and practice contexts?

The dissertation is organized into three parts.

○ *Part I: Foundations* establishes the theoretical and methodological groundwork.

Chapter 1: Introduction situates the research within debates about AI in creative practice. It establishes how efficiency-driven AI tools create ‘design fixation’ where the speed and completeness of output generation causes designers to become prematurely attached to initial results. The chapter introduces the mirror metaphor: the understanding that AI systems trained on human-curated data inevitably carry the subjective judgments of those who construct them. Rather than treating this subjectivity as a flaw to be eliminated, Reflective AI sees it as an opportunity for learning about one’s own practice. The chapter defines reflection through Phoebe Sengers’ work: bringing unconscious aspects of experience to conscious awareness, thereby making them available for conscious choice.

Chapter 2: Methodological Orientations articulates three principles that characterize Reflective AI practice: reorienting purpose from efficiency toward reflection, treating AI as malleable material, and relying on first-person research methods. The first principle involves reorienting purpose from efficiency to reflection, drawing on the philosophy of Slow Technology to create ‘time presence’: a mode of engagement where the slowness of the process itself creates room for thinking, noticing, and reconsidering, rather than making time disappear. The second principle treats AI as a malleable material rather than a fixed tool, recognizing trainability as a key property that enables designers to shape AI

systems through their curatorial decisions. The third principle relies on first-person research methods, positioning the practitioner’s direct experience as primary knowledge. The chapter explains how these principles emerged through practice and how they guide the subsequent investigations.

○ *Part II: Reflective AI Across Creative Practices* documents five practice-based investigations arranged chronologically to show the evolution of the research.

Chapter 3: Ceci n’est pas un chaise: An Emerging Dialogue between an AI and a Designer establishes a foundational model for designer-AI dialogue through object detection. The chapter reports on explorations where the researcher engages with object recognition algorithms, investigating how an AI’s ‘surprising’ interpretations can provoke framing and reframing. When applying pretrained object detection to surrealist paintings and designed objects, the AI’s ‘mislabelings’ provided alternative perspectives that could activate reflective thought. The final exploration creates an iterative dialogue where the designer and AI respond to each other through cycles of interpretation and physical modification, operationalizing Schön’s ‘seeing-moving-seeing’ cycle. The key insight: what appears as technical failure can become productive creative material.

Chapter 4: Developing the Objective Portrait Method explores how three professional designers use self-trained AI to reflect on

their individual ‘design worlds’: the implicit aesthetic preferences, visual fascinations, and ways of seeing that shape their practice. Each designer trained an object detection model on a dataset representing their creative interests, then used this personalized model to analyze a carefully composed image of their own work, an Object Portrait. The chapter introduces the Quadrant Tool for selecting subjective annotation labels and reveals that the most profound reflection came not from the AI’s final analysis but from the training process itself. Three key insights emerged: annotation forces articulation of latent visual concepts; annotation drift (the gradual shift in how labels are applied as one’s understanding evolves during annotation) creates productive friction; and anticipating AI analysis shapes how the designers involved create and view their work.

Chapter 5: The Objective Portrait Workshop tests whether the method developed with practitioners can translate to design education. A three-day workshop with eleven design students at Design Academy Eindhoven shows how the approach functions as a teaching tool for students actively learning to become reflective designers. Findings reveal two forms of learning: learning *through* AI as a reflective medium (articulating subjective judgments, anticipating AI analysis, discovering symbolic meaning) and learning *about* AI through hands-on practice (discovering machine learning principles, confronting the annotation

consistency challenge). Students developed distinct ‘annotation strategies’ to navigate the tension between their fluid human judgment and the model’s need for consistency, demonstrating that reflective approaches are learnable and adaptable to educational contexts.

Chapter 6: “Why Do I Keep Coming Back to This Chair?” turns the method toward auto-ethnographic inquiry, using text-to-image generation to excavate artistic meaning in past work. The chapter documents a deeply personal investigation into a foam chair created eight years prior during design school. By training a generative AI model on archival photographs and prompting it with emotionally charged concepts, the researcher discovered that foam had become a personal metaphor for consolation and self-care. The mirror metaphor evolved into ‘AI-as-Memory,’ recognizing that AI models trained on personal archives can reconstruct and reflect the past. The chapter demonstrates how Reflective AI can address questions about creative history, revealing latent patterns and values that had remained unarticulated.

Chapter 7: From Text-to-Clay brings Reflective AI into dialogue with ceramic practice through a three-month residency at the European Ceramic Work Centre. The chapter explores what happens when the digital processes of AI training encounter the slow, physical materiality of ceramics. A circular methodology emerged: clay becomes data, data becomes training material,

trained models generate images, images become physical sculptures, sculptures generate new data. The AI's structural impossibilities (forms that physics would not allow) became creative challenges that pushed the researcher's manual skill into new territories of making, creating reflections on personal style, the meaning of work, and both materials of clay and AI.

○ *Part III: Synthesis and Implications*

Chapter 8: Discussion synthesizes patterns across the five investigations and articulates what can be transferred beyond the specific contexts studied. The chapter identifies five types of reflection that emerged across projects and formalizes two design learnings for practitioners. Five types of reflection are identified across the projects: defamiliarization, articulation, anticipatory reflection, autobiographical reflection (looking back at creative history), and translational reflection (constant translation between computational and physical materiality). The first learning concerns a disentangled workflow that treats each stage of AI training as a distinct reflective space: data curation as inquiry, annotation as articulation, inference as dialogue, output as proposition, and making as anchor. The second learning concerns engagement with failure and ambiguity, reframing technical errors (overfitting, hallucinations, mislabeling) as productive friction rather than problems to solve. The chapter discusses implications for AI

tool design and broader AI discourse, proposing that AI can evolve from a closed tool for automation into a malleable material for inquiry when the goal shifts from generation to reflection.

The dissertation makes several contributions. First, it establishes three core principles that characterize AI engagement oriented toward reflection: reorienting purpose from efficiency toward reflection, treating AI as malleable material, and relying on first-person research methods. Second, it provides documented methods across multiple contexts, offering both methodological transparency and practical guidance for others seeking to work in similar ways. Third, it synthesizes two design learnings for practitioners that offer flexible frameworks adaptable to different domains. Fourth, it contributes methodological insights regarding how first-person research with AI can generate transferable knowledge despite producing situated and particular findings. The thesis proposes an “architecture/weights” distinction: while individual reflective insights remain particular per project, the methodological structure through which they are generated can be transferred, allowing others to produce their own situated insights through similar processes. Finally, the dissertation offers a counter-narrative to dominant AI discourse by demonstrating how the same technologies can serve fundamentally different purposes when approached through reflective practice.

The findings suggest that Reflective AI functions within specific conditions. It requires time and space, as curating personal datasets, annotating hundreds of images, and responding to outputs through making is slow work. It also requires technical confidence to navigate the “productive friction” of working with AI systems. The method assumes a visual or material practice and a desire for introspection. The projects also showed that completing the steps does not guarantee reflection; practitioners must actively engage with the process for insights to emerge. These conditions mean Reflective AI is not a replacement for commercial pipelines but rather a pedagogical or strategic intervention, functioning best in the early phases of design, in education, or in contexts where the goal is understanding rather than production.

In conclusion, this dissertation demonstrates that when the goal of interaction with AI shifts from merely generation and automation to reflection, AI can evolve from a closed tool into a malleable material for personal inquiry. The value of the technology within creative design practice lies in its capacity to externalize and refract the designer's own implicit judgments, fascinations, and historical patterns. Rather than debating whether AI enhances or diminishes creativity in absolute terms, this research shows how particular configurations of AI practice can facilitate particular forms of self-knowledge under specific conditions. The question shifts from

what AI can do to what long-term engagement with AI training can help designers understand about their own practice, values, and assumptions.

● *Samenvatting in het Nederlands*
De opkomst van Kunstmatige Intelligentie (AI) in de creatieve sector heeft een steeds grotere invloed op de werkwijze van ontwerpers en kunstenaars. Commerciële AI-tools zoals Midjourney, Gemini en ChatGPT beloven het creatieve proces te versnellen en te vergemakkelijken door simpele tekstprompts direct om te zetten in beeld, en kunnen zo bijvoorbeeld vroege schetsen en ideeën in enkele seconden transformeren tot overtuigende eindproducten.

Hoewel deze focus op efficiëntie past binnen de hedendaagse nadruk op efficiëntie binnen het ontwerpvak, beperkt het ook de manier waarop ontwerpers zich tot AI verhouden. De simpele interfaces van deze tools verbergen specifieke onderdelen van deze vorm van AI, zoals dataverzameling, annotatie en modeltraining, waardoor ontwerpers slechts als eindgebruiker interacteren met AI. Deze versimpeling, in combinatie met de snelheid van het genereren van outputs, ontnemt ontwerpers de mogelijkheid tot wat Donald Schön omschreef als een dialoog met het materiaal: de reflectieve wisselwerking tussen de ontwerper en zijn omgang met materialen die essentieel is voor een betekenisvolle ontwerp-praktijk.

Dit proefschrift introduceert een andere benadering: Reflectieve

AI als ontwerppraktijk. In deze benadering wordt AI niet gezien als een automatiseringstool, maar als materiaal voor reflectie. Dit kernidee is ontstaan vanuit de eigen praktijk van de onderzoeker: de technische pipeline van AI (het verzamelen van data, annoteren en trainen) kan juist fungeren als het materiaal waarmee ontwerpers hun reflectieve praktijk vormgeven. In plaats van AI in te zetten voor efficiëntie, positioneert Reflectieve AI deze technologie als een spiegel. Deze spiegel onthult de manier van kijken van de ontwerper, hun impliciete voorkeuren en oordelen en de veranderlijkheid daarvan.

Het onderzoek volgt de onderzoeksmethode van Practice-Based Research, waarbij onderzoeksvragen voortkomen uit het maakproces van de onderzoeker zelf en de antwoorden dienen om diezelfde praktijk te verhelderen en te verbeteren. De gepresenteerde methode is dan ook niet vooraf bepaald en getest, maar organisch ontstaan door het maakproces zelf: door AI in diverse contexten als materiaal te benaderen en te reflecteren op wat dat oplevert qua inzicht in de eigen identiteit als maker. De casestudies in de hoofdstukken omvatten individuele projecten, een samenwerking met professionele ontwerpers, een workshop met studenten aan een ontwerp-school, auto-etnografisch onderzoek naar de eigen artistieke ontwikkeling en een drie maanden durende residentie gericht op keramiek.

Vanuit deze praktijkervaringen is één overkoepelende onderzoeksvraag geformuleerd: *Welke specifieke ontwerppraktijken maken een reflectieve interactie met AI mogelijk?* Deze vraag wordt onderzocht aan de hand van twee deelvragen: Ten eerste, wat kenmerkt reflectie wanneer ontwerpers direct betrokken zijn bij het trainingsproces van AI? En ten tweede, welke inzichten voor de ontwerppraktijk komen voort uit dit onderzoek die overdraagbaar zijn naar andere ontwerp- en onderzoekscontexten?

Het proefschrift bestaat uit drie delen.

○ *Deel I - Grondslagen* legt het theoretische en methodologische fundament voor het onderzoek.

Hoofdstuk 1 - Inleiding positioneert het onderzoek binnen het debat over AI in de creatieve praktijk. De literatuur beschrijft hoe AI-tools die gericht zijn op efficiëntie leiden tot 'ontwerpfixatie' (*design fixation*): de snelheid en volledigheid van de gegenereerde output binden ontwerpers te vroeg aan voorlopige resultaten. Het hoofdstuk introduceert de 'spiegelmetafoor': het inzicht dat AI-systemen, getraind op door mensen gecureerde data, onvermijdelijk de subjectieve oordelen dragen van hun makers. In plaats van deze subjectiviteit te zien als een fout die geëlimineerd moet worden, beschouwt Reflectieve AI dit als een kans om te leren over de eigen praktijk. Reflectie wordt hierbij gedefinieerd op basis van het werk van Phoebe Sengers:

het bewustmaken van onbewuste aspecten van een ervaring, waardoor de mogelijkheid ontstaat voor bewuste keuzes.

Hoofdstuk 2 - Methodologische Oriëntaties formuleert drie principes voor de Reflectieve AI-praktijk: heroriëntatie van efficiëntie naar reflectie, AI benaderen als vormbaar materiaal, en onderzoek vanuit het eerstpersoonsperspectief. Het eerste principe betreft een heroriëntatie van efficiëntie naar reflectie, waarbij de filosofie van Slow Technology van Hallnas & Redstrom (2001) wordt gebruikt om 'tijdsbesef' (*time presence*) te creëren: een manier van werken met technologie waarbij de traagheid van het proces zelf ruimte schept voor denken, opmerken en heroverwegen, in plaats van tijd te laten verdwijnen zoals dat vaak gebeurt met efficiënt ontworpen technologie. Het tweede principe benadert AI als vormbaar materiaal in plaats van een vaststaand instrument; hierbij wordt 'trainbaarheid' erkend als een cruciale eigenschap van de technologie die ontwerpers in staat om als curator beslissingen te nemen over AI-systemen en ze zo mede te vormen. Het derde principe steunt op onderzoek vanuit het eerstpersoonsperspectief, waarbij de directe ervaring van de maker de primaire kennisbron is voor het onderzoeken van reflectieve ervaringen

○ *Deel II - Reflectieve AI in de Creatieve Praktijk* documenteert vijf praktijkgerichte onderzoeken, chronologisch geordend om de evolutie van het

onderzoek te laten zien.

Hoofdstuk 3 - Ceci n'est pas une chaise: Een beginnende dialoog tussen ontwerper en AI via objectdetectie levert een basismodel voor een reflectieve dialoog. Dit hoofdstuk beschrijft experimenten waarin de onderzoeker algoritmen voor objectherkenning inzet om te zien hoe foutieve maar verrassende interpretaties van AI tot een nieuw interpretatiekader kunnen leiden. Door objectdetectie toe te passen op surrealistische schilderijen en designobjecten, boden de foutieve labels van de AI een alternatief perspectief aan dat aanzet tot reflectief denken. Deze ontwikkelt zich vervolgens tot een iteratieve dialoog waarin ontwerper en AI op elkaar reageren via cycli van interpretatie en fysieke aanpassing van een object, waarmee Donald Schöns cyclus van 'zien-bewegen-zien' (seeing-moving-seeing) operationeel wordt gemaakt in een maakpraktijk met gebruik van AI.

Hoofdstuk 4 - Ontwikkeling van de 'Objective Portrait' Methode onderzoekt hoe drie professionele ontwerpers zelfgetrainde AI gebruiken om te reflecteren op hun individuele "ontwerpwerelden" (design worlds): de impliciete esthetische voorkeuren, visuele fascinaties en manieren van kijken die hun praktijk vormgeven. Elke ontwerper trainde een objectdetectiemodel op een dataset die hun creatieve interesses representeerde, en gebruikte dit gepersonaliseerde model vervolgens om een zorgvuldig gecomponeerd beeld van hun eigen werk te analyseren, onder de titel Object Portret. Het hoofdstuk

introduceert de kwadrant-grafiek als hulpmiddel voor het selecteren van subjectieve annotatielabels en toont aan dat de diepste reflectie gedurende dit proces niet voortkwam uit de uiteindelijke analyse van de AI, maar uit het trainingsproces zelf. Drie belangrijke inzichten kwamen naar voren: annotatie dwingt tot het articuleren van intuïtieve visuele voorkeuren of ideeën; de veranderlijkheid van oordelen tijdens het annoteren creëert een conflict dat doet nadenken over de eigen denk- en maakprocessen; en het anticiperen op de uiteindelijke analyse van de zelfgetrainde AI beïnvloedt hoe ontwerpers hun werk maken en bekijken.

Hoofdstuk 5-De ‘Objective Portrait’ Workshop toetst of de methode van Hoofdstuk 4, ontwikkeld met ervaren ontwerpers, vertaald kan worden naar het ontwerponderwijs. Een driedaagse workshop met elf studenten aan de Design Academy Eindhoven laat zien hoe de benadering werkt als leermiddel voor studenten die zich ontwikkelen tot reflectieve ontwerpers. De bevindingen tonen twee vormen van leren: leren *door* AI als reflectief medium (het articuleren van subjectieve oordelen, anticiperen op AI-analyse, ontdekken van symbolische betekenis in eigen werk) en leren *over* AI door deze praktisch te benaderen (machine learning-principes ontdekken, omgaan met de noodzaak tot consistente annotatie voor een ‘goedwerkend’ model). Studenten ontwikkelden diverse ‘annotatiestrategieën’ om te

navigeren tussen hun eigen dynamische, menselijke oordeel en de behoefte van het model aan consistentie.

Hoofdstuk 6- “Waarom Blijf Ik Terugkomen naar Deze Stoel?” richt de methode op auto-etnografisch onderzoek, waarbij tekst-naar-beeld generatie wordt ingezet om artistieke betekenis in oud werk bloot te leggen. Het hoofdstuk documenteert een diepgaand persoonlijk onderzoek naar een stoel gemaakt van schuim, die acht jaar eerder tijdens de opleiding werd gemaakt en vervolgens vernietigd. Door een generatief AI-model te trainen op archiefphoto’s en het te bevragen/prompten met emotioneel geladen concepten (bijvoorbeeld met: ‘een scheiding van schuim’), ontdekte de onderzoeker dat schuim een persoonlijke metafoor was geworden voor troost en zelfzorg. De spiegelmetafoor evolueerde hier naar AI-als-Geheugen, waarbij duidelijk wordt dat AI-modellen getraind op persoonlijke archieven een nieuw licht kunnen werpen op het verleden en reflectie daarop kunnen stimuleren.

Hoofdstuk 7- Van Tekst-naar-Klei brengt Reflectieve AI in dialoog met de keramische praktijk middels een residentie van drie maanden bij het Europees Keramisch Werkcentrum. Het hoofdstuk onderzoekt wat er gebeurt wanneer de digitale processen van AI-training de trage, fysieke aard van keramiek ontmoeten. Er ontstaat een circulaire methodologie: klei wordt data, data wordt trainingsmateriaal, getrainde modellen genereren beelden, beelden

worden fysieke sculpturen en die sculpturen genereren weer nieuwe trainingsdata. Het structurele onvermogen van de AI om in keramiek realiseerbare vormen te creëren werd een creatieve uitdaging die het vakmanschap van de onderzoeker naar nieuw terrein duwde en inzichten bieden over de eigen maakpraktijk, de technologie en het materiaal klei.

○ *Deel III- Synthese en Implicaties*
Hoofdstuk 8- Discussie voegt patronen uit de vijf onderzoeken samen en beschrijft wat overdraagbaar is buiten de specifieke contexten waarin Reflectieve AI is onderzocht. Het hoofdstuk identificeert vijf typen reflectie die gedurende de projecten naar voren kwamen: defamiliarisatie, articulatie, anticiperende reflectie, autobiografische reflectie (terugkijken op de creatieve geschiedenis), en translationele reflectie (voortdurende vertaling tussen computationele en fysieke materialiteit). Het hoofdstuk formuleert ook twee inzichten voor de praktijk. Het eerste inzicht betreft een ‘ontkoppelde workflow’ die elke fase van AI-training behandelt als een afzonderlijke reflectieve ruimte: datacuratie als onderzoek, annotatie als articulatie, gebruik van het model als dialoog, output als hypothese en het maken als anker. De tweede les betreft het omarmen van falen en ambiguïteit, waarbij technische fouten (zoals overfitting, hallucinaties en verkeerde labels) worden geherinterpreteerd als productieve wrijving in een designproces, in

plaats van als problemen die opgelost moeten worden.

Het proefschrift levert verschillende inzichten. Ten eerste stelt het drie kernprincipes vast voor een designer-AI relatie die gericht is op reflectie: heroriëntatie van efficiëntie naar reflectie, AI benaderen als vormbaar materiaal, en onderzoek vanuit het eerstepersoonsperspectief. Ten tweede biedt het gedocumenteerde methoden in meerdere creatieve contexten. Ten derde synthetiseert het twee inzichten voor ontwerpers die overgenomen kunnen worden wanneer er interesse is in reflectief gebruik van AI. kaders bieden en toepasbaar zijn op verschillende domeinen binnen het ontwerpveld. Ten vierde geeft het methodologische inzichten over hoe eerste-persoons onderzoek met AI generaliseerbare kennis produceert, ondanks het specifieke, persoonlijke karakter van de bevindingen van ieder hoofdstuk. Het proefschrift stelt een onderscheid voor tussen ‘gewichten en architectuur’: terwijl individuele inzichten (de gewichten) persoonlijk en specifiek blijven, kan de methodologische structuur (de architectuur) waarmee ze zijn gegenereerd worden overgedragen, zodat anderen hun eigen inzichten kunnen verwerven via een vergelijkbaar proces. Tot slot biedt het proefschrift een tegengeluid voor het dominante AI-discours door aan te tonen hoe technologieën fundamenteel andere doelen kunnen dienen dan waarvoor ze primair bedoeld waren, wanneer ze worden benaderd vanuit een reflectieve praktijk.

De bevindingen suggereren ook dat Reflectieve AI functioneert onder specifieke voorwaarden; het vereist tijd en ruimte. Het cureren van persoonlijke datasets, het annoteren van honderden afbeeldingen en het reageren op outputs door fysiek te maken is namelijk tijdsintensief. Deze aanpak vereist ook een niveau van technisch zelfvertrouwen om te kunnen omgaan met de ‘productieve frictie’ van deze manier van werken met AI-systemen. De methode veronderstelt een visuele of materiële maakpraktijk, gecombineerd met een verlangen of behoefte aan introspectie en zelfreflectie. De projecten lieten ook zien dat het doorlopen van de stappen geen garantie is voor reflectie; de maker moet actief betrokken zijn bij het proces om tot inzichten te komen. Deze voorwaarden laten zien dat Reflectieve AI geen vervanging is voor commerciële AI-tools, maar eerder een pedagogische of strategische interventie. Het werkt het best in de vroege fasen van ontwerp, in het onderwijs, of in

contexten waar begrip en verdieping het doel is, in plaats van snelle productie.

In conclusie demonstreert dit proefschrift dat wanneer het doel van de interactie met AI verschuift van generatie en automatisering naar reflectie, AI kan evolueren van een gesloten tool naar vormbaar maak- en onderzoeksmateriaal. De waarde van deze technologie binnen de creatieve ontwerppraktijk ligt in haar vermogen om de impliciete oordelen, fascinaties en historische patronen van de ontwerper te externaliseren en terug te spiegelen. In plaats van te debatteren of AI menselijke creativiteit in absolute zin verbetert, vermindert of vervangt, toont dit onderzoek aan hoe deze configuratie van een AI-praktijk — onder de juiste voorwaarden — specifieke vormen van zelfkennis kan faciliteren. De vraag verschuift van wat AI kan, naar wat een langdurige betrokkenheid bij AI-training ontwerpers kan helpen begrijpen over hun eigen praktijk, waarden en aannames.

12. Appendix

Appendix A: Supplementary Materials for Chapter 4

This appendix provides practical documentation of the Objective Portrait method and detailed training data from the three models. The step-by-step protocol is included for practitioners who wish to adapt the method for their own reflective practice. The training data offers transparency into how subjective annotation translated into model behavior.

A.1 The Objective Portrait method: Step-by-step protocol

The following protocol documents the complete workflow of the Objective Portrait method. Each step includes practical instructions that emerged from our process.

Step 1: Collecting a Dataset

To train an AI model on your way of seeing, you first need a collection of images representing a fascination relevant to your practice. This dataset becomes the visual vocabulary through which the model will learn your subjective judgments.

Choose a fascination connected to your design practice: a visual domain you have opinions about, feelings toward, or experience with.

Collect 200–400 images representing this fascination. Sources may include online archives, museum databases, personal photographs, or curated collections.

Ensure images depict scenes containing multiple elements (rather than isolated close-ups), as the model needs distinct regions to detect.

Aim for some visual consistency within the dataset (e.g., all paintings, or all photographs) to help the model learn patterns.

Step 2: Selecting Labels with the Quadrant Tool

The Quadrant Tool helps articulate the subjective judgments you will use to annotate your dataset. Rather than starting with predefined categories, this process surfaces the feelings and opinions you might already have when looking at your images.

Draw a large cross (+) on a sheet of paper. On the horizontal axis, write ‘cold’ on the left and ‘warm’ on the right. On the vertical axis, write ‘pull’ at the top and ‘push’ at the bottom. These axes describe how elements in an image feel to you.

From your dataset, choose one image that fits most clearly in each corner of the quadrant.

For each image, circle the specific part that most evokes that feeling.

Name the feeling or emotion this evokes (such as ‘happy,’ ‘frustrated,’ ‘unsettled’).

Name the subjective judgment connected to this feeling (such as ‘cute,’ ‘ugly,’ ‘pretentious’). List all words that come to mind.

Select four final words as your subjective labels. Ensure they are sufficiently different from each other (‘scary’ and ‘horrific’ could overlap too much). If labels feel too similar, find alternatives.

Note:

We found it helpful to develop labels in conversation with others. Putting vague aesthetic preferences into words proved easier when articulating them to someone else.

Step 3: Annotating the Dataset

To train an object detection algorithm, the dataset must be annotated. This means drawing bounding boxes over specific parts of each image and labeling them with one of your four labels. We used the platform Roboflow for annotation.

Create an account on www.roboflow.com
Create a new project and select “Object Detection”
Upload the images from your dataset
Create a new ‘job’ in the Annotations section
Begin annotating: draw boxes over parts of the image and enter the appropriate label. After the first entry, you can select labels using number keys (1, 2, 3, 4)
Annotate at least 100 images (more is better for model learning)
When finished, click ‘Add images to dataset’

Note on annotation:

Since feelings may arise from viewing a whole image, annotation requires identifying specific visual elements that evoke that feeling. This shift in perspective, from holistic impression to localized attribution, is itself a reflective exercise. The challenge of pinpointing *where* a feeling comes from often reveals something about the nature of that judgment.

Exporting the Dataset

The annotated dataset must be exported in a format compatible with the YOLO algorithm. In Roboflow, go to ‘Generate’ and click ‘New Version’
Optionally add augmentation steps (variations in saturation, brightness, or slight rotation) to increase the number of training images
Generate the dataset and download it in YOLO format

Step 4: Training the Algorithm

We used the open-source object detection algorithm YOLOv5 (Jocher et al., 2022). Training can be performed locally (requiring a capable GPU) or in cloud environments such as Google Colab. See Section A.2 for our hardware specifications and training parameters.

Step 5: Creating the Object Portrait

The Object Portrait is a carefully composed image or video created specifically to be analyzed by your trained model. To ensure the model can recognize elements in your portrait, it should be visually similar to your training dataset.
Gather objects from past projects or materials connected to your practice
Plan how to create an image that visually resembles your dataset. Consider lighting, colors, orientation, and setting
Compose and photograph the arrangement
Edit the image to align visually with your training dataset (contrast, brightness, color grading)
Export the final image

Note on visual linking:

We developed the term ‘visual linking’ to describe incorporating stylistic elements from the training dataset into the Object Portrait. This might include matching color palettes, mimicking photographic angles, or recreating environmental contexts. The goal is to create conditions where your trained model can “read” your personal work through the vocabulary it learned from your dataset.

Step 6: Analyzing the Object Portrait

Run your trained YOLO model on your Object Portrait to generate detections. The model will identify elements and assign your subjective labels with confidence scores.
Interpretation approach:
Rather than evaluating detections as correct or incorrect, treat them as starting points for reflection:
Why might the model have detected this element with this label?
Does this detection align with or surprise my own judgment?

What does this reveal about how I annotated the training data?
What patterns in my annotation choices does this detection expose?

The confidence score indicates how strongly the detected element matches patterns the model learned during training. A detection of ‘76% stubborn’ does not mean the element *is* stubborn in any objective sense. It means the visual features resemble what you labeled as ‘stubborn’ in your training data. However, it is up to you as the creator of your portrait and AI to make that judgement what it means to you.

A.2 Training Data and Model Performance

This section presents the training data from the three Objective Portrait models we created for Chapter 4. The figures show label distributions and confusion matrices, offering insight into how subjective annotation translated into model behavior, for those interested in the technical details. Label Distributions and Confusion Matrices

Figure 10.1 presents an overview of the training data and model performance for all three participants. The top row shows how frequently each designer applied their four subjective labels during annotation, reflecting both the content of the source datasets and individual annotation tendencies. The bottom row shows confusion matrices indicating how the trained models performed on validation data: each cell shows how often a predicted label (rows) matched the annotated label (columns), with the diagonal representing correct predictions and the bottom row (‘background FN’) indicating missed detections.

pastedGraphic.png ~

Vera’s surrealist art dataset (378 images, approximately 4,400 bounding boxes) showed ‘teasing’ and ‘absurd’ applied most frequently, with these labels also achieving the best detection rates (22% and 14% respectively). The confusion between these two labels suggests visual overlap between the concepts. Gijs’s garden dataset (approximately 300 images, approximately 6,200 bounding boxes) had more evenly distributed labels, with dense annotation indicating he identified many distinct elements within each photograph. Maurik’s public space dataset (124 images, approximately 2,400 bounding boxes) showed ‘excluding’ dominating his annotations, followed by ‘inviting,’ which may reflect the nature of public art in Rotterdam, his particular attentional focus, or both. The confusion between ‘excluding’ and ‘inviting’ in his model points to the visual ambiguity of these spatial concepts.

Hardware and Training Parameters

All models were trained on the Industrial Design Engineering department’s computing cluster at TU Delft:

Server: Supermicro Superserver SYS-7049GP-TRT
Processors: 2x Intel Xeon Gold 5218 (22M Cache, 2.30 GHz)
Memory: 384GB DDR4-2666 ECC RAM
GPUs: 3x NVIDIA GeForce RTX 3090 TURBO 24GB
Location: IDE Datacenter, Room E-0-806, TU Delft
Training parameters:
Architecture: YOLOv5L (large variant) with pretrained weights
Epochs: 200–338 depending on model
Batch size: 16
Image resolution: 640×640 pixels
Optimizer: Stochastic Gradient Descent (SGD)
Training time: Approximately 2 hours per model

A.3 Data Availability

For transparency and reproducibility, the complete training outputs for all our models are available in the accompanying research data repository. This includes the annotation files, trained

Appendix B: Supplementary Materials for Chapter 5

This appendix provides detailed documentation of the workshop protocol and analytical framework used in Chapter 5. These materials are included for transparency and to support practitioners who may wish to adapt the Objective Portrait method for their own educational contexts.

B.1 Workshop Task Descriptions
Day 1: Dataset Preparation and Annotation

Dataset Collection Criteria

Prior to the workshop, we instructed students to collect a digital dataset based on the following criteria to ensure feasibility for the object detection model:

Quantity: At least 100 images (collecting more was encouraged).

Subject: Images must depict one specific topic, phenomenon, or concept (e.g., locks, tidiness, surrealism, seating in public space) closely related to the student’s research interest.

Consistency in Style: All images must share a visual medium (e.g., all studio photos, or all screenshots).

Consistency in Scale: The images should function as a ‘scene’ containing elements for the algorithm to detect, rather than isolated close-ups.

Initial Group Sharing

We facilitated a collective review where each student presented their chosen dataset and articulated the rationale behind their selection to the group.

Label Selection Task

We guided the quadrant-based label selection using four steps:

Categorize 10 images from your printed set into the quadrants based on their general feel.

Find details in the images that evoke that feeling most.

For each quadrant, write a list of words that describes this feeling best.

Select the four words that capture the underlying judgment most accurately.

Post-Labeling Reflection and Debrief

After the labeling task, students responded to the following written prompt: “For each quadrant, write some thoughts on how your view of these terms changed during the process.” Following this individual reflection, we held a collective moment focusing on how the students experienced the labeling task and how it influenced their perspective on the data.

Day 2: Object Portrait Creation

Design Brief

Students were tasked with creating an artifact to be read by their model. The design constraints required the Object Portrait to be:

An A2 photo print or a 10-second Full HD video.

Composed of objects related to their specific work.

Styled aesthetically to match the training dataset.

Feasible to produce within a single afternoon session.

Collective Reflection on Making

We concluded the day with a brief collective moment where students discussed their experience of translating their dataset aesthetics into a physical Object Portrait.

Day 3: AI Analysis and Reflection
Model Interpretation

Upon running their trained models on their Object Portraits, students were asked to share in the group: “Look at the AI analysis. What do you see? What do you think the interpretations mean?”

Final Individual Reflection

We concluded the workshop with a structured reflection card containing the following questions:

What did you learn about your topic and yourself?

What are moments where you felt your position (views of your topic) changed? (Consider labeling, portrait-making, and interpreting).

How do you relate to your algorithm? Did that change how you relate to AI in general?

Highlight one thing you want to share about your process.

B.2 Thematic Analysis Structure

Table 10.1 outlines the final thematic structure that emerged from our reflexive analysis. These themes correspond to the findings detailed in Section 5.5.

Table 10.1

Identified Theme	Description of Pattern in Data
Learning Through AI (Reflective Medium)	
Articulating Subjective Judgments	Students using the label selection process to distill complex, fluid emotions into specific vocabulary (e.g., distinguishing ‘enticing’ from ‘cloying’).
Anticipating AI Analysis	Moments where students altered their creative work (the Object Portrait) because they internalized how the model might ‘see’ it.
Symbolic Discovery	Instances where students interpreted model errors or unexpected detections as metaphorical or poetic meaning rather than technical failure.
Learning About AI (Technical Awareness)	
Discovering ML Principles	Students deducing technical concepts (e.g., pattern recognition, generalization) through observing model performance (e.g., “it reads pixels, not text”).
Negotiating Annotation Consistency	The struggle between maintaining a consistent dataset for the machine and the student’s evolving personal judgment during the process.

B.3 Hardware and Training Parameters

All student models were trained using the same infrastructure and parameters as documented for Chapter 4 (see Section A.2). Training was performed overnight between Day 2 and Day 3 of the workshop, with each student’s annotated dataset processed individually. Each of the eleven models trained for 200 epochs with batch size 16 at 640×640 pixel resolution, taking approximately 1.5 to 2 hours per model.

Model performance metrics (mAP scores) were not the primary focus of evaluation in this educational context. The pedagogical goal centered on reflective engagement with the training process, where even models with low technical accuracy could generate meaningful insights through their unexpected interpretations. Students were encouraged to treat model outputs as starting points for reflection rather than as measures of success or failure.

B.4 Data Availability

For transparency and reproducibility, the complete training outputs for all student models are available in the accompanying research data repository. This includes the annotation files, trained model weights, label distribution charts, and confusion matrices documented in this appendix. The data can be accessed at 4TU.ResearchData: <https://doi.org/10.4121/6ea01644-acc6-4478-a581-1f6788dec1f7>

Appendix C: Supplementary Materials for Chapter 6

This appendix documents the autoethnographic methodology used to investigate my artistic past through AI-as-Memory. These materials provide transparency about the dataset, annotation strategy, technical setup, and prompting approach developed for this personal inquiry.

C.1 Archival Dataset Overview

The initial archival exploration encompassed 13 projects from my time at design school, representing a broad portrait of my practice during that period. Table 10.2 lists these projects. After examining the full archive, the inquiry narrowed to a single project: the foam chair (SoftChair). The final training dataset consisted of 50 images of this chair, captured from various angles and compositions.

Table 10.2

Project Name	Description
BlueThing	[Design school project]
CeramicAnne	[Ceramic work]
CeramicCups	[Ceramic vessels]
ContextObjects	[Object studies]
FilmStillLife	[Film-based still life work]
Lamp	[Lighting design]
SoftChair	The foam chair — selected for this inquiry
StillLifesFinal	[Final still life compositions]
StillLifesProcess	[Process documentation of still lifes]
StillLifeStudio	[Studio-based still life work]
StoolDancer	[Performance/furniture hybrid]
StudioObjects1	[Studio object collection]
TouchFilm	[Touch-related film project]

C.2 Reflective Metadata Development

Before annotation, I developed reflective metadata to structure my memories of the foam chair project. This involved recording both objective and subjective terms:
Objective terms
(observable features): foam, collage, big, white, rolls, layers, fabric
Subjective terms
(emotions and associations): womb-like, monster, safety, awkward, comforting, scrumptious
This metadata served as a bridge between latent memories and the explicit, structured information required for annotation.

C.3 Annotation Strategy

Unlike the object detection annotation used in Chapters 4 and 5, text-to-image fine-tuning requires a single descriptive text file paired with each image. I developed a six-part annotation structure that combined factual elements with subjective interpretation:
Concept: A constant feature present in each image (e.g., ‘soft womb armchair,’ ‘foam womb chair’)
Material/Method: The materials or making process (e.g., ‘made from various layers and rolls of foam,’ ‘made out of rolls of foam and fabric’)
Photographic angle: The perspective of the image (e.g., ‘front view,’ ‘side view,’ ‘top view,’ ‘zoomed-in view’)
Emotional adjective: The overall feeling the chair evoked (e.g., ‘monster-like,’ ‘monster’)
Abstract visual detail: A notable feature described interpretively (e.g., ‘tentacles on top,’ ‘scrumptious vibe’)
Interaction: The person interacting and how, if applicable (e.g., ‘vera sitting,’ ‘vera lying with legs crossed,’ ‘vera kneeling in front’)

C.3.1 Complete Annotation Examples

The following examples show the range of annotations across the 50-image dataset:
SoftChair1.txt: A soft womb armchair, made from various layers and rolls of foam, monster-like, scrumptious vibe, tentacles on top, front view
SoftChair7.txt: A soft womb armchair, made from various layers and rolls of foam, monster-like, scrumptious vibe, tentacles on top, side view, portrait
SoftChair20.txt: A soft womb armchair, made from various layers and rolls of foam, vera kneeling in front, monster-like, scrumptious vibe, tentacles on top, back view, landscape
SoftChair36.txt: A soft womb armchair, two big yellow foam rolls as head rests, monster-like, scrumptious vibe, detailed zoomed-in view, inside
SoftChair38.txt: vera sitting in a soft womb chair made out of rolls of foam and fabric, foam tentacles on top, front view, monster
SoftChair52.txt: vera with closed eyes lying in a soft womb chair made out of rolls of foam and fabric, a close up, holding foam tentacles, monster
SoftChair54.txt: A close-up view of vera lying in a foam womb chair, with legs crossed, surrounded by layers of soft, rolled foam materials, top view, monster
SoftChair62.txt: vera lying in a foam womb chair with knees, surrounded by layers of soft, rolled foam materials, side view, monster, tentacles on top
SoftChair65.txt: vera lying in a foam womb chair with knees up and hand, surrounded by layers of soft, rolled foam materials, front view, tentacles on top, monster
SoftChair70.txt: the back of a foam womb chair with feet sticking out, made out of layers of rolled foam materials, back side, monster

C.4 Technical Setup

C.4.1 Training Configuration

Training was performed using Google Colab with the Kohya LoRA trainer notebook.¹¹ Table 10.3 presents the complete training configuration.

Table 10.3

Parameter	Value
Base model	Stable Diffusion v1.5 (v1-5-pruned-emaonly. safetensors)
Fine-tuning method	LoRA (Low-Rank Adaptation)
Network dimension	16
Network alpha	8
UNet learning rate	0.0005
Text encoder learning rate	0.0001

Optimizer	AdamW8bit
Scheduler	cosine_with_restarts (3 restarts)
Training images	50
Repeats per image	10
Epochs	10
Mixed precision	fp16
Seed	42

The 10 training epochs created 10 distinct model versions representing progressive learning stages. This allowed observation of how the model’s interpretation of prompts evolved as it learned more from the training data.

C.4.2 Generation Parameters

The following parameters were used when generating images with the fine-tuned model:

Parameter	Value
Sampling method	DPM++ 2M Karras
Sampling steps	52
Width	512
Height	672
CFG Scale	7
Clip skip	3
SD VAE	Automatic

Fixed seed values were used for each prompt to ensure that visual changes across training epochs reflected the model’s learning rather than random variation. This allowed coherent visual narratives to emerge across the 10 epochs.

C.5 Data Availability

For transparency and reproducibility, the complete training data, annotations, and model outputs for this autoethnographic inquiry are available in the accompanying research data repository. The data can be accessed at 4TU.ResearchData: <https://doi.org/10.4121/6ea01644-acc6-4478-a581-1f6788dec1f7>

Appendix D: Supplementary Materials for Chapter 7

This appendix provides detailed documentation of the circular ceramic-AI methodology developed during the three-month residency at the European Ceramic Work Centre (EKWC). It includes the complete process workflow, annotation archives from both training cycles, the prompt lexicon, and technical parameters for model training and image generation.

D.1 The Circular Methodology

The ceramic-AI practice followed a six-step circular process, with two complete cycles documented during the residency:

Data Collection – Photographing ceramic works to create training datasets

Annotation – Labeling images with subjective, interpretive descriptions combining material facts with emotional qualities

Training – Fine-tuning a Stable Diffusion model using LoRA on the annotated dataset

Prompting – Testing the trained model with abstract emotional concepts

Physical Translation – Reconstructing selected AI-generated images in clay

Documentation – Photographing completed works to generate new training data, restarting the cycle

The methodology created a temporal feedback loop where each cycle deepened the relationship between emotional concepts and ceramic forms.

D.2 Annotation Archives

D.2.1 Cycle 1: Training on Past Ceramic Work

The first training dataset consisted of 16 images of ceramic works created seven years prior, representing early explorations with clay before formal training. These images shared consistent visual characteristics: cream-colored glazes, green studio backgrounds, and similar photographic compositions. The annotation strategy combined factual visual elements (material, background color) with subjective interpretive labels (emotional qualities, character descriptions). The trigger word “ceramic” activated the LoRA fine-tuning during generation. Table 10.4 presents the complete annotation archive for this cycle.

Table 10.4

File	Annotation
CeramicCup1.txt	ceramic, a humble cup, shaped like a bull, organic, stubborn
CeramicCup2.txt	two ceramic figures lounging, spreading, sunbathing, relaxing, green background
CeramicCup3.txt	three porcelain figures, cutlery, lounging, long
CeramicCup4.txt	ceramic cup with lots of handles, virus, surrealism, visual confusion, green background
CeramicCup5.txt	three ceramic cups with lots of handles, virus, surrealism, weird, visual confusion, green background
CeramicCup6.txt	group of cream colored ceramic objects against a green background, chaotic atmosphere, provocative, perverse
CeramicCup7.txt	group of cream colored ceramic objects against a green background, chaotic atmosphere, provocative, perverse, gang
CeramicCup8.txt	group of cream colored ceramic objects against a green background, chaotic atmosphere, provocative, perverse, gang, playful
CeramicCup9.txt	ceramic misfits, virus, melting, group, mistakes, fleshy
CeramicCup10.txt	ceramic misfits, virus, melting, group, mistakes, fleshy
CeramicCup11.txt	ceramic misfits, virus, melting, group, mistakes, close up, fleshy

CeramicCup12.txt	ceramic relaxing, spreading, sunbathing, fat, lounging, content
CeramicCup13.txt	ceramic relaxing, spreading, sunbathing, fat, lounging, content
CeramicCup14.txt	porcelain cup sitting on top of a porcelain spoon, provocative, love, sensual
CeramicCup15.txt	porcelain cup sitting on top of a porcelain spoon, provocative, love, sensual
CeramicCup16.txt	Ceramic cup with lots of handles, virus, misfit, confusion, surrealism

D.2.2 Cycle 2: Training on Bone-Dry Sculptures

The second training dataset consisted of 21 images documenting ten bone-dry sculptures created during the residency. All sculptures were photographed on white cloth against the EKWC studio environment, creating visual consistency while capturing the specific material state of unfired clay. The annotation strategy shifted to incorporate the original prompts directly, combined with material state descriptors. This linked each sculpture back to its generative origin while encoding the new visual context. The number of images per sculpture varied based on form complexity and photographic angles captured. Table 10.5 presents the complete annotation archive for this cycle.

Table 10.5

File	Annotation	Source Prompt
dryceramic1.txt	dry ceramic female pain, on a white cloth	ceramic female pain
dryceramic2.txt	dry ceramic female pain, on a white cloth	ceramic female pain
dryceramic3.txt	dry ceramic female pain, on a white cloth	ceramic female pain
dryceramic4.txt	dry ceramic ignorance, on a white cloth	ceramic ignorance
dryceramic5.txt	dry ceramic ignorance, on a white cloth	ceramic ignorance
dryceramic6.txt	dry ceramic jealousy, on a white cloth	ceramic jealousy
dryceramic7.txt	dry ceramic jealousy, on a white cloth	ceramic jealousy
dryceramic8.txt	dry ceramic mask, on a white cloth	ceramic mask
dryceramic9.txt	dry ceramic sadness, on a white cloth	ceramic sadness
dryceramic10.txt	dry ceramic sadness, on a white cloth	ceramic sadness
dryceramic11.txt	dry ceramic sadness, on a white cloth	ceramic sadness
dryceramic12.txt	dry ceramic octopus, on a white cloth	ceramic octopus
dryceramic13.txt	dry ceramic octopus, on a white cloth	ceramic octopus
dryceramic14.txt	dry ceramic octopus, on a white cloth	ceramic octopus
dryceramic15.txt	dry ceramic surrealism, on a white cloth	ceramic surrealism
dryceramic16.txt	dry ceramic woman sitting waiting, on a white cloth	ceramic woman sitting waiting
dryceramic17.txt	dry ceramic woman sitting waiting, on a white cloth	ceramic woman sitting waiting
dryceramic18.txt	dry ceramic surrealism, on a white cloth	ceramic surrealism

dryceramic19.txt	dry ceramic divorce, on a white cloth	ceramic divorce
dryceramic20.txt	dry ceramic red surrealism, on a white cloth	red ceramic surrealism
dryceramic21.txt	dry ceramic red surrealism, on a white cloth	red ceramic surrealism

D.3 Technical Parameters

D.3.1 Training Infrastructure

Training was performed using Google Colab with the Kohya LoRA trainer notebook.¹² Resource constraints required low-memory optimization settings. Training image resolutions were 1024×1024 pixels for Cycle 1 (bucketed to 1216×832) and 512×512 pixels for Cycle 2 (bucketed to 448×576). Models were saved at epochs 1, 3, 5, 8, and 10; epoch selection for generation was based on intuitive assessment of output quality rather than systematic comparison.

D.3.2 Training Configuration

The two training cycles used different base models and configurations, reflecting an evolving understanding of the technical requirements. Table 10.6 presents the complete training configuration for both cycles.

Table 10.6

Parameter	Cycle 1	Cycle 2
Base model	SDXL 1.0	Stable Diffusion 1.5
Checkpoint	sd_xl_base_1.0	sd-v1-5-pruned-noema-fp16
Fine-tuning method	LoRA	LoRA
Network dimension	8	16
Network alpha	4	8
UNet learning rate	0.0003	0.0005
Text encoder learning rate	0.00006	0.0001
Optimizer	AdamW8bit	AdamW8bit
Scheduler	cosine_with_restarts (3)	cosine_with_restarts (3)
Batch size	4	2
Training images	16	21
Repeats per image	2	10
Epochs	10	10
Total training steps	60	1050
Training duration	equation.pdf ~10 minutes	1_#\$!@%#!#_equation.pdf ~15 minutes
Mixed precision	fp16	fp16
Seed	42	42
Trigger word	ceramic	ceramic

Note: While Cycle 1 was trained on SDXL, image generation for both cycles used Stable Diffusion 1.5 as the base model for consistency across the project.

D.3.3 Generation Parameters

Images were generated using consistent parameters across both cycles, as shown in Table 10.7.

Table 10.7

Parameter	Value
Base model	Stable Diffusion v1.5
Sampler	DPM++ 2M Karras
Sampling steps	52
Image dimensions	512×672 pixels (portrait)
CFG scale	7
Clip skip	3

D.4 Ceramic Materials and Firing

The ceramic process used materials and facilities available at EKWC. Table 10.8 presents the material specifications and firing parameters.

Table 10.8

Parameter	Specification
Clay body	K129 (Sibelco, Westerwald, Germany)
Clay type	White stoneware with 30% molochite chamotte (0–0.5mm grain size)
Firing range	1000–1280°C
Bisque firing	2_#\$_!@%#!#_equation.pdf~1200°C (high bisque to complete shrinkage before glazing)
Glaze firing	900–1000°C
Surface treatment	Sinter Engobe SE10 (EKWC in-house formula-tionz with iron oxide additions (MnO ₃ , Fe ₂ O ₃) and pigments (z))

The K129 clay body was selected for its thermal shock resistance and suitability for hand-building sculptural forms. The high bisque temperature ensured maximum shrinkage occurred before the unpredictable glaze firing, reducing structural stress on the fragile forms. The sinter engobe base provided a consistent matt white surface, with iron oxides and pigments applied via spraying to create gradient effects that echoed the ambiguous shadows present in the AI-generated source images.

D.5 Data Availability

For transparency and reproducibility, the complete training data, annotations, prompt lexicons, and model outputs for both cycles are available in the accompanying research data repository. The data can be accessed at 4TU.ResearchData: <https://doi.org/10.4121/6ea01644-acc6-4478-a581-1f6788dec1f7>

13. Bibliography

Abadi, M., Barham, P., Chen, J., Chen, Z., Davis, A., Dean, J., Devin, M., Ghemawat, S., Irving, G., Isard, M., Kudlur, M., Levenberg, J., Monga, R., Moore, S., Murray, D. G., Steiner, B., Tucker, P., Vasudevan, V., Warden, P., ... & Xiao, Q. (2016). TensorFlow: A system for large-scale machine learning. In *12th USENIX Symposium on Operating Systems Design and Implementation (OSDI 16)* (pp. 265–283). USENIX Association. <https://www.usenix.org/system/files/conference/osdi16/osdi16-abadi.pdf>

Abuzuraig, A. M., & Pasquier, P. (2024). Towards personalizing generative AI with small data for co-creation in the visual arts. In *Joint Proceedings of the ACM IUI Workshops (IUI Workshops 2024)*. CEUR Workshop Proceedings. <https://ceur-ws.org/Vol-3660/paper11.pdf>

AIxDesign. (2024). *An introduction to Slow AI*. <https://aixdesign.co/posts/slow-ai-preface>

Baumer, E. P. S., Siedel, D., McDonnell, L., Zhong, J., Sittikul, P., & McGee, M. (2020). Topicalizer: Reframing core concepts in machine learning visualization by co-designing for interpretivist scholarship. *Human-Computer Interaction*, 35(5–6), 452–480. <https://doi.org/10.1080/07370024.2020.1734460>

Benjamin, J. J., Biggs, H., Berger, A., Rukanskaitė, J., Heidt, M. B., Merrill, N., Pierce, J., & Lindley, J. (2023). The entoptic field camera as metaphor-driven research-through-design with AI technologies. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems* (pp. 1–19). Association for Computing Machinery. <https://doi.org/10.1145/3544548.3581175>

Benjamin, J. J., Berger, A., Merrill, N., & Pierce, J. (2021). Machine Learning Uncertainty as a Design Material: A Post-Phenomenological Inquiry. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems* (Article 171). ACM. <https://doi.org/10.1145/3411764.3445481>

Benjamin, J. J., Lindley, J., Edwards, E., Rubegni, E., Korjakow, T., Grist, D., & Sharkey, R. (2024). Responding to generative AI technologies with research-through-design: The Ryelands AI Lab as an exploratory study. *Proceedings of the ACM on Human-Computer*

Interaction, 8(TEI), 1823–1841. <https://doi.org/10.1145/3643834.3660677>

Berger, J. 1972. *Ways of Seeing*. British Broadcasting Corporation; Penguin Books.

Bilstrup, K.-E. K., Kaspersen, M. H., Bouvin, N. O., & Petersen, M. G. (2025). The ML-Machine toolkit: Empowering teachers and education professionals to explore embodied approaches to teaching machine learning. In *Proceedings of the 2025 ACM Designing Interactive Systems Conference* (pp. 535–551). Association for Computing Machinery. <https://doi.org/10.1145/3715336.3735737>

BKOR. (n.d.). BKOR (Beeldende Kunst & Openbare Ruimte) - Rotterdam. Geraadpleegd op 10 februari 2022, van <https://www.bkor.nl/>

Bochkovskiy, A., Wang, C.-Y., & Liao, H.-Y. M. (2020). *YOLOv4: Optimal speed and accuracy of object detection*. arXiv. <https://doi.org/10.48550/arXiv.2004.10934>

Bory, P., Natale, S., & Katzenbach, C. (2025). Strong and weak AI narratives: An analytical framework. *AI & Society*, 40, Article 107. <https://doi.org/10.1007/s00146-024-02087-8>

Braun, V., Clarke, V., Hayfield, N., Davey, L., & Jenkinson, E. (2023). Doing reflexive thematic analysis. In *Supporting research in counselling and psychotherapy* (pp. 19–38). Springer International Publishing. https://doi.org/10.1007/978-3-031-13942-0_2

Bridle, J. (2017). *Autonomous trap 001* [Artwork]. <https://jamesbridle.com/works/autonomous-trap-001>

Bridle, J. (2018, 22 February). Interview: There is a huge demand to humanise the technological view of the world. *Studio International*. <https://www.studiointernational.com/index.php/james-bridle-interview-technology-self-driving-cars>

Bridle, J. (2022). *Ways of being: Animals, plants, machines: The search for a planetary intelligence*. Farrar, Straus and Giroux.

Burnap, A., Liu, Y., Pan, Y., Lee, H., Gonzalez, R., & Papalambros, P. Y. (2016). Estimating and exploring the product form design space using deep generative models. In *Proceedings of the ASME 2016 International Design Engineering Technical Conferences and Computers and*

Information in Engineering Conference (Volume 2). American Society of Mechanical Engineers. <https://doi.org/10.1115/DETC2016-60091>

Burrell, J. (2016). How the machine “thinks”: Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1). <https://doi.org/10.1177/2053951715622512>

Candy, L., & Edmonds, E. (2018). Practice-based research in the creative arts: Foundations and futures from the front line. *Leonardo*, 51(1), 63–69. https://doi.org/10.1162/LEON_a_01509

Caramiaux, B., Crawford, K., Liao, Q. V., Ramos, G., & Williams, J. (2025). *Generative AI and creative work: Narratives, values, and impacts*. arXiv. <https://doi.org/10.48550/arXiv.2502.03940>

Caramiaux, B., & Fdili Alaoui, S. (2022). “Explorers of unknown planets”: Practices and politics of artificial intelligence in visual arts. *Proceedings of the ACM on Human-Computer Interaction*, 6(CSCW2), 477:1–477:24. <https://doi.org/10.1145/3555578>

Chang, J. C., Amershi, S., & Kamar, E. (2017). Revolt: Collaborative crowdsourcing for labeling machine learning datasets. In *Proceedings of the 2017 CHI Conference on Human Factors in Computing Systems* (pp. 2334–2346). Association for Computing Machinery. <https://doi.org/10.1145/3025453.3026044>

Chen, L., Wang, P., Dong, H., Liang, X., & Zhang, K. (2019). An artificial intelligence based data-driven approach for design ideation. *Journal of Visual Communication and Image Representation*, 61, 10–22. <https://doi.org/10.1016/j.jvci.2019.02.007>

Chubb, J., Reed, D., & Cowling, P. (2022). Expert views about missing AI narratives: Is there an AI story crisis? *AI & Society*, 39, 2393–2407. <https://doi.org/10.1007/s00146-022-01548-2>

Crawford, K. 2021. *Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. Yale University Press.

Crawford, K. (2024, 26 november). Eating the future: The metabolic logic of AI slop. *e-flux Architecture*. <https://www.e-flux.com/architecture/intensification/6782975/eating-the-future-the-metabolic-logic-of-ai-slop/>

Cremaschi, M., Dorfmann, M., & De Angeli, A. (2024). A steampunk critique of machine learning acceleration. In *Proceedings of the 2024 ACM Designing Interactive*

Systems Conference (pp. 246–257). Association for Computing Machinery. <https://doi.org/10.1145/3643834.3660688>

Cross, N. (2006). *Designerly ways of knowing*. Springer.

Daalhuizen, J., & Cash, P. (2021). Method content theory: Towards a new understanding of methods in design. *Design Studies*, 75, Article 101018. <https://doi.org/10.1016/j.destud.2021.101018>

Desjardins, A., Tomico, O., Lucero, A., Cecchinato, M. E., & Neustaedter, C. (2021). Introduction to the special issue on first-person methods in HCI. *ACM Transactions on Computer-Human Interaction*, 28(6), 1–12. <https://doi.org/10.1145/3492342>

Dignum, V. (2021). The myth of complete AI-fairness. arXiv. <https://doi.org/10.48550/arXiv.2104.12544>

DK Eyewitness. (2022). *Gardens of the world*. DK.

Dorst, K. (2015). *Frame innovation: Create new thinking by design*. MIT Press.

Dorst, K., & Cross, N. (2001). Creativity in the design process: Co-evolution of problem-solution. *Design Studies*, 22(5), 425–437. [https://doi.org/10.1016/S0142-694X\(01\)00009-6](https://doi.org/10.1016/S0142-694X(01)00009-6)

Dove, G., Halskov, K., Forlizzi, J., & Zimmerman, J. (2017). UX design innovation: Challenges for working with machine learning as a design material. In *Proceedings of the 2017 CHI Conference on Human Factors in Computing Systems* (pp. 278–288). Association for Computing Machinery. <https://doi.org/10.1145/3025453.3025739>

Dwyer, B., Nelson, J., & Solawetz, J. (2022). *Roboflow* (Version 1.0) [Computer software]. <https://roboflow.com>

Ellis, C., Adams, T. E., & Bochner, A. P. (2011). Autoethnography: An overview. *Historical Social Research*, 36(4), 273–290. <https://doi.org/10.12677/HSR.2011.364015>

Elwes, J. (2019). *Zizi- Queering the dataset* [AI art project]. <https://www.jakeelwes.com/project-zizi-2019.html>

Exactly.ai. (2024). *Exactly.ai: Personal generative AI model*. <https://exactly.ai/>

Fu, Z., & Zhou, Y. (2020). Research on human-AI co-creation based on reflective design practice. *CCF Transactions on Pervasive Computing and Interaction*, 2(1), 33–41. <https://doi.org/10.1007/s42486-020-00028-0>

Gaver, W. W., Beaver, J., & Benford, S. (2003). Ambiguity as a resource for design.

In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems* (pp. 233–240). Association for Computing Machinery. <https://doi.org/10.1145/642611.642653>

Gaver, W., Krogh, P. G., Boucher, A., & Chatting, D. (2022). Emergence as a feature of practice-based design research. In *Proceedings of the 2022 ACM Designing Interactive Systems Conference* (pp. 517–526). Association for Computing Machinery. <https://doi.org/10.1145/3532106.3533524>

Giaccardi, E., & Karana, E. (2015). Foundations of materials experience: An approach for HCI. In *Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems* (pp. 2447–2456). Association for Computing Machinery. <https://doi.org/10.1145/2702123.2702337>

Giaccardi, E., Murray-Rust, D., Redström, J., & Caramiaux, B. (2024). Prototyping with uncertainties: Data, algorithms, and research through design. *ACM Transactions on Computer-Human Interaction*, 31(6), Article 86. <https://doi.org/10.1145/3674728>

Hallnäs, L., & Redström, J. (2001). Slow technology—designing for reflection. *Personal and Ubiquitous Computing*, 5(3), 201–212. <https://doi.org/10.1007/PL00000019>

Haraway, D. (1988). Situated knowledges: The science question in feminism and the privilege of partial perspective. *Feminist Studies*, 14(3), 575–599. <https://doi.org/10.2307/3178066>

Haraway, D. J. (2016). *Staying with the trouble: Making kin in the Chthulucene*. Duke University Press.

Herndon, H., & Dryhurst, M. (2024). *The call: New rituals for collaboration with AI* [Exhibition]. Serpentine North, London. <https://www.serpentinegalleries.org/whats-on/holly-herndon-mat-dryhurst-the-call/>

Holmquist, L. E. (2017). Intelligence on tap: Artificial intelligence as a new design material. *Interactions*, 24(4), 28–33. <https://doi.org/10.1145/3085571>

Hu, X., Xing, Y., Cai, X., Yang, S., & Li, Z. (2025). Designing interactions with generative AI for art and creativity: A systematic review and taxonomy. In *Proceedings of the 2025 ACM Designing Interactive Systems Conference* (pp. 1126–1155). Association for Computing Machinery. <https://doi.org/10.1145/3715336.3735843>

Hwang, A. H.-C. (2022). Too late to be creative? AI-empowered tools in creative processes. In *Extended Abstracts of the 2022 CHI Conference on Human Factors in Computing Systems* (Article 451). Association for Computing Machinery. <https://doi.org/10.1145/3491101.3503549>

Ihde, D., & Malafouris, L. (2019). Homo faber revisited: Postphenomenology and material engagement theory. *Philosophy & Technology*, 32(2), 195–214. <https://doi.org/10.1007/s13347-018-0321-7>

Jansson, D. G., & Smith, S. M. (1991). Design fixation. *Design Studies*, 12(1), 3–11. [https://doi.org/10.1016/0142-694X\(91\)90003-F](https://doi.org/10.1016/0142-694X(91)90003-F)

Ji, S.-J., Ling, Q.-H., & Han, F. (2023). An improved algorithm for small object detection based on YOLO v4 and multi-scale contextual information. *Computers and Electrical Engineering*, 105, Article 108490. <https://doi.org/10.1016/j.compeleceng.2022.108490>

Jocher, G., Chaurasia, A., Stoken, A., Borovec, J., NanoCode012, Kwon, Y., TaoXie, Fang, J., imyhxy, Michael, K., Lorna, Abhiram, V., Diego, M. T., Akshat-Talapaula, Hogan, A., lorenzomamma, AlexWang1900, ... & Yonghye, K. (2022). *Ultralytics/yolov5: v7.0-YOLOv5 SOTA realtime instance segmentation* (Version v7.0) [Computer software]. Zenodo. <https://doi.org/10.5281/zenodo.7347926>

Kazi, R. H., Grossman, T., Cheong, H., Hashemi, A., & Fitzmaurice, G. W. (2017). DreamSketch: Early stage 3D design explorations with sketching and generative design. In *Proceedings of the 30th Annual ACM Symposium on User Interface Software and Technology* (pp. 401–414). Association for Computing Machinery. <https://doi.org/10.1145/3126594.3126662>

Kolko, J. (2010). Abductive thinking and sensemaking: The drivers of design synthesis. *Design Issues*, 26(1), 15–28. <https://doi.org/10.1162/desi.2010.26.1.15>

Koskinen, I., Zimmerman, J., Binder, T., Redström, J., & Wensveen, S. (2011). *Design research through practice: From the lab, field, and showroom*. Morgan Kaufmann.

Krizhevsky, A., Sutskever, I., & Hinton, G. E. (2012). ImageNet classification with deep convolutional neural networks. *Advances in Neural Information Processing Systems*, 25, 1097–1105. <https://proceedings.neurips.cc/paper/2012/file/c399862d3b9d6b76c8436e924a68c45b-Paper.pdf>

- Li, J., Cao, H., Lin, L., Hou, Y., Zhu, R., & El Ali, A. (2024). User experience design professionals' perceptions of generative artificial intelligence. In *Proceedings of the CHI Conference on Human Factors in Computing Systems* (Article 273). Association for Computing Machinery. <https://doi.org/10.1145/3613904.3642114>
- Lin, T.-Y., Maire, M., Belongie, S., Hays, J., Perona, P., Ramanan, D., Dollár, P., & Zitnick, C. L. (2014). Microsoft COCO: Common objects in context. In D. Fleet, T. Pajdla, B. Schiele, & T. Tuytelaars (Eds.), *European Conference on Computer Vision* (pp. 740–755). Springer. https://doi.org/10.1007/978-3-319-10602-1_48
- Lindrup, M., Jacobsen, R. M., Wester, J., van Berkel, N., Raptis, D., & Nielsen, P. A. (2025). Prompt Machine: A tangible generative AI tool for supporting children's learning and literacy. In *Proceedings of the 2025 ACM Designing Interactive Systems Conference* (pp. 1210–1225). Association for Computing Machinery. <https://doi.org/10.1145/3715336.3735673>
- Liu, C., Takikawa, T., & Jacobson, A. (2024). *A LoRA is worth a thousand pictures*. arXiv. <https://doi.org/10.48550/arXiv.2412.12048>
- Lloyd, P. (2019). You make it and you try it out: Seeds of design discipline futures. *Design Studies*, 65, 167–181. <https://doi.org/10.1016/j.destud.2019.10.008>
- Lousberg, L., Rooij, R., Jansen, S., van Dooren, E., Heintz, J., & van der Zaag, E. (2020). Reflection in design education. *International Journal of Technology and Design Education*, 30(5), 885–897. <https://doi.org/10.1007/s10798-019-09532-6>
- Ludden, G. D. S., Kudrowitz, B. M., Schifferstein, H. N. J., & Hekkert, P. (2012). Surprise and humor in product design. *Humor*, 25(3), 285–309. <https://doi.org/10.1515/humor-2012-0015>
- Maguire, E. A., Vargha-Khadem, F., & Hassabis, D. (2010). Imagining fictitious and future experiences: Evidence from developmental amnesia. *Neuropsychologia*, 48(11), 3187–3192. <https://doi.org/10.1016/j.neuropsychologia.2010.06.037>
- Maher, M. L., & Fisher, D. H. (2012). Using AI to evaluate creative designs. In *Proceedings of the 2nd International Conference on Design Creativity* (ICDC 2012) (Vol. 1, pp. 45–54).
- Malafouris, L. (2008). At the potter's wheel: An argument for material agency. In C. Knappett & L. Malafouris (Eds.), *Material agency: Towards a non-anthropocentric approach* (pp. 19–36). Springer. https://doi.org/10.1007/978-0-387-74711-8_2
- Malsattar, N., Kihara, T., & Giaccardi, E. (2019). Designing and prototyping from the perspective of AI in the wild. In *Proceedings of the 2019 on Designing Interactive Systems Conference* (pp. 1083–1088). Association for Computing Machinery. <https://doi.org/10.1145/3322276.3322351>
- Matejka, J., Glueck, M., Bradner, E., Hashemi, A., Grossman, T., & Fitzmaurice, G. (2018). Dream lens: Exploration and visualization of large-scale generative design datasets. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems* (Paper 369). Association for Computing Machinery. <https://doi.org/10.1145/3173574.3173943>
- Miceli, M., Schuessler, M., & Yang, T. (2020). Between subjectivity and imposition: Power dynamics in data annotation for computer vision. *Proceedings of the ACM on Human-Computer Interaction*, 4(CSCW2), Article 115. <https://doi.org/10.1145/3415186>
- Midjourney [@midjourney]. (2024, 30 juli). What's new in V6.1 [Statusupdate]. X. <https://x.com/midjourney/status/1818342845885141185>
- Midjourney Inc. (2023). Midjourney (Versie 5.2) [Computer software]. <https://www.midjourney.com>
- Milovanovic, J., & Gero, J. S. (2020). Modeling design studio pedagogy: A mentored reflective practice. *Proceedings of the Design Society: DESIGN Conference*, 1, 1765–1774. <https://doi.org/10.1017/dsd.2020.118>
- Mujtaba, D. F., & Mahapatra, N. R. (2019). Ethical considerations in AI-based recruitment. In *2019 IEEE International Symposium on Technology and Society* (ISTAS) (pp. 1–7). IEEE. <https://doi.org/10.1109/ISTAS48451.2019.8937920>
- Murray-Rust, D., Lupetti, M. L., Nicenboim, I., & Van Der Hoog, W. (2023). Grasping AI: Experiential exercises for designers. *AI & Society*, 39(4), 1647–1661. <https://doi.org/10.1007/s00146-023-01794-y>
- Murray-Rust, D., Nicenboim, I., & Dan Lockton. (2022). Metaphors for designers working with AI. In *DRS2022: Bilbao*. Design Research Society. <https://doi.org/10.21606/drs.2022.667>
- Nachtwey, O., & Seidl, T. (2024). The solutionist ethic and the spirit of digital capitalism. *Theory, Culture & Society*, 41(2), 91–112. <https://doi.org/10.1177/02632764231196829>
- Nelson, L. K. (2020). Computational grounded theory: A methodological framework. *Sociological Methods & Research*, 49(1), 3–42. <https://doi.org/10.1177/0049124117729703>
- Neustaedter, C., & Sengers, P. (2012). Autobiographical design in HCI research: Designing and learning through use-it-yourself. In *Proceedings of the 2012 Designing Interactive Systems Conference* (pp. 514–523). Association for Computing Machinery. <https://doi.org/10.1145/2317956.2318034>
- Odom, W., Stolterman, E., & Chen, A. (2022). Extending a theory of slow technology for design through artifact analysis. *Human-Computer Interaction*, 37(2), 150–179. <https://doi.org/10.1080/07370024.2021.1913416>
- Odom, W., Wakkary, R., Hol, J., Nafeh, B., Neustaedter, C., Chen, A., & Jagat, P. K. (2019). Investigating slowness as a frame to design longer-term experiences with personal data: A field study of Olly. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (Paper 33). Association for Computing Machinery. <https://doi.org/10.1145/3290605.3300264>
- Pan, S. J., & Yang, Q. (2010). A survey on transfer learning. *IEEE Transactions on Knowledge and Data Engineering*, 22(10), 1345–1359. <https://doi.org/10.1109/TKDE.2009.191>
- Papageorgiou, C., & Poggio, T. (2000). A trainable system for object detection. *International Journal of Computer Vision*, 38(1), 15–33. <https://doi.org/10.1023/A:1008162616689>
- Papert, S. (1991). Situating constructionism. In I. Harel & S. Papert (Eds.), *Constructionism* (pp. 1–11). Ablex Publishing Corporation. <https://pirun.ku.ac.th/~btun/papert/sitcons.pdf>
- Park, J. S., Gelderloos, J., Re'em, Y., Howell, N. P., & Hartmann, B. (2019). A slow algorithm improves users' assessments of the algorithm's accuracy. In *Proceedings of the 2019 Designing Interactive Systems Conference* (pp. 59–70). Association for Computing Machinery. <https://doi.org/10.1145/3322276.3322336>
- Park, S., Benecke, B., & Kotopoulos, O. (2025). Autonomy of a fall [Performance installation]. Cinedans Festival, Amsterdam. <https://soyunparrkk.com/projects/autonomy-of-a-fall/>
- Paton, B., & Dorst, K. (2011). Briefing and reframing: A situated practice. *Design Studies*, 32(6), 573–587. <https://doi.org/10.1016/j.destud.2011.07.002>
- Peterson, B. L. (2017). Thematic analysis/interpretive thematic analysis. In *The international encyclopedia of communication research methods*. Wiley. <https://doi.org/10.1002/9781118901731.iecrm0249>
- Ploin, A.-S., Eynon, R., Hjorth, I., & Osborne, M. A. (2022). *AI and the arts: How machine learning is changing artistic work* (Report from the Creative Algorithmic Intelligence Research Project). Oxford Internet Institute, University of Oxford. <https://www.ox.ac.uk/news/2022-03-03-art-our-sake-artists-cannot-be-replaced-machines-study>
- Rapp, A. (2018). Autoethnography in human-computer interaction: Theory and practice. In M. Filimowicz & V. Tzankova (Eds.), *New directions in third wave human-computer interaction: Volume 2 - Methodologies* (pp. 25–42). Springer. https://doi.org/10.1007/978-3-319-73374-6_3
- Ridler, A. (2018). Myriad (Tulips) [Artwork]. <https://annaridler.com/works/myriad-tulips>
- Rombach, R., Blattmann, A., Lorenz, D., Esser, P., & Ommer, B. (2022). High-resolution image synthesis with latent diffusion models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (pp. 10684–10695). IEEE. <https://doi.org/10.1109/CVPR52688.2022.01046>
- Rozendaal, M. C., Ghajargar, M., Pasman, G., & Wiberg, M. (2018). Giving form to smart objects: Exploring intelligence as an interaction design material. In M. Filimowicz & V. Tzankova (Eds.), *New directions in third wave human-computer interaction: Volume 1 - Technologies* (pp. 117–135). Springer. https://doi.org/10.1007/978-3-319-73356-2_3
- Russell, J. A. (1980). A circumplex model of affect. *Journal of Personality and Social Psychology*, 39(6), 1161–1178. <https://doi.org/10.1037/h0077714>
- Salvaggio, E. (2024). Critical AI art [Artist portfolio]. Cybernetic Forests. <https://www.cyberneticforests.com/online-portfolio>
- Sandhaus, H., Gu, Q., Parreira, M. T., & Ju, W. (2025). Co-designing with transformers: Unpacking the complex role of GenAI in interactive system design education. In *Proceedings of the 2025 ACM Designing Interactive Systems Conference* (pp. 1228–1243).

Association for Computing Machinery. <https://doi.org/10.1145/3715336.3735805>

Sawyer, R. D., & Norris, J. (2012). *Duoethnography*. Oxford University Press. <https://doi.org/10.1093/acprof:osobl/9780199757404.001.0001>

Schacter, D. L., & Addis, D. R. (2007). The cognitive neuroscience of constructive memory: Remembering the past and imagining the future. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 362(1481), 773–786. <https://doi.org/10.1098/rstb.2007.2087>

Schön, D. A. (1983). *The reflective practitioner: How professionals think in action*. Basic Books.

Schön, D. A. (1984). Problems, frames and perspectives on designing. *Design Studies*, 5(3), 132–136. [https://doi.org/10.1016/0142-694X\(84\)90002-4](https://doi.org/10.1016/0142-694X(84)90002-4)

Schön, D. A. (1988). Designing: Rules, types and worlds. *Design Studies*, 9(3), 181–190. [https://doi.org/10.1016/0142-694X\(88\)90047-6](https://doi.org/10.1016/0142-694X(88)90047-6)

Schön, D. A. (1992). Designing as a reflective conversation with the materials of the design situation. *Research in Engineering Design*, 3(3), 131–147. <https://doi.org/10.1007/BF01580516>

Schön, D. A., & Wiggins, G. (1992). Kinds of seeing and their functions in designing. *Design Studies*, 13(2), 135–156. [https://doi.org/10.1016/0142-694X\(92\)90268-F](https://doi.org/10.1016/0142-694X(92)90268-F)

Sengers, P., Boehner, K., David, S., & Kaye, J. J. (2005). Reflective design. In *Proceedings of the 4th Decennial Conference on Critical Computing: Between Sense and Sensibility* (pp. 49–58). Association for Computing Machinery. <https://doi.org/10.1145/1094562.1094569>

Sennett, R. (2008). *The craftsman*. Yale University Press.

Serra, G., & Miralles, D. (2021). Human-level design proposals by an artificial agent in multiple scenarios. *Design Studies*, 76, Article 101029. <https://doi.org/10.1016/j.destud.2021.101029>

Silk, E. M., Rechkemmer, A. E., Daly, S. R., Jablokow, K. W., & McKilligan, S. (2021). Problem framing and cognitive style: Impacts on design ideation perceptions. *Design Studies*, 74, Article 101015. <https://doi.org/10.1016/j.destud.2021.101015>

Steinbrück, A., & Stankowski, A. (2023). Creative ownership and control for generative AI in art and design. Paper presented at the Generative AI in HCI Workshop, CHI 2023, Hamburg, Germany. <https://>

generativeaiandhci.github.io/papers/2023/genaichi-creativeownership.pdf

STM Association. (2025). *Recommendations for a classification of AI use in academic manuscript preparation*. <https://www.stm-assoc.org/standards-technology/ai-ethics-and-integrity/>

Tannen, D. (1986). Frames revisited. *Quaderni di Semantica*, 7(1), 106–109.

Tanweer, A., Gade, E., Krafft, P. M., & Dreier, S. (2021). Why the data revolution needs qualitative thinking. *Harvard Data Science Review*, 3(3). <https://doi.org/10.1162/99608f92.3792070e>

Thurgood, C., and Lulham, R. (2016) Exploring framing and meaning making over the design innovation process, in Lloyd, P. and Bohemia, E. (eds.), Future Focused Thinking- DRS International Conference 2016, 27 – 30 June, Brighton, United Kingdom. <https://doi.org/10.21606/drs.2016.218>

Ultralytics. (2023). *Comprehensive guide to Ultralytics YOLOv5*. <https://docs.ultralytics.com/yolov5/>

Vallor, S. (2024). *The AI mirror: How to reclaim our humanity in an age of machine thinking*. Oxford University Press.

Wong, W. L. P., & Radcliffe, D. F. (2000). The tacit nature of design knowledge. *Technology Analysis & Strategic Management*, 12(4), 493–512. <https://doi.org/10.1080/713698497>

Wu, Z., Ji, D., Yu, K., Zeng, X., Wu, D., & Shidujaman, M. (2021). AI creativity and the human-AI co-creation model. In *M. Kurosu (Ed.), Human-computer interaction: Theory, methods and tools* (pp. 171–190). Springer. https://doi.org/10.1007/978-3-030-77319-9_13

Yang, Q., Steinfeld, A., Rosé, C., & Zimmerman, J. (2020). Re-examining whether, why, and how human-AI interaction is uniquely difficult to design. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1–13). Association for Computing Machinery. <https://doi.org/10.1145/3313831.3376301>

Zhang, G., Raina, A., Cagan, J., & McComb, C. (2021). A cautionary tale about the impact of AI on human design teams. *Design Studies*, 72, Article 100990. <https://doi.org/10.1016/j.destud.2020.100990>

Zou, Z., Chen, K., Shi, Z., Guo, Y., & Ye, J. (2023). Object detection in 20 years: A survey. *Proceedings of the IEEE*, 111(3), 257–276. <https://doi.org/10.1109/JPROC.2023.3240166>

1 Inference refers to the moment when a trained model is applied to new inputs to produce outputs.

2 a subset of AI focused on systems that learn from data rather than explicit programming,

3 A pre-trained artificial intelligence system that is publicly available and can be used directly without requiring extensive technical expertise, custom development, or training from scratch. These tools typically come with pre-configured models, standardized interfaces, and documentation that allusers to apply them to tasks immediately.

4 Tensorflow is an end-to-end open-source platform for Machine Learning, containing a comprehensive ecosystem of tools to build and deploy machine learning models (<https://www.tensorflow.org/>)

5 <https://kiosk-dot-codelabs-site.appspot.com/codelabs/tensorflow-for-poets/#0>

6 <https://pjreddie.com/darknet/yolo/>

7 Research data supporting this chapter are available in 4TU.ResearchData at <https://doi.org/10.4121/6ea01644-acc6-4478-a581-1f6788dec1f7>.

8 Note on terminology: The broader methodology is referred to as the Objective Portrait, while the specific visual artifact created by the student is referred to as the Object Portrait.

9 Roboflow Annotate: <https://roboflow.com/annotate>

10 A ‘black box’ in computational contexts refers to a system whose internal workings are opaque to the user: we can observe what goes in and what comes out, but the processes that transform one into the other remain hidden or difficult to interpret.

11 https://colab.research.google.com/github/hollowstrawberry/kohya-colab/blob/main/Lora_Trainer.ipynb

12 https://colab.research.google.com/github/hollowstrawberry/kohya-colab/blob/main/Lora_Trainer.ipynb

As a designer, I came to artificial intelligence out of fascination: I wanted to understand how machines see, and what that might reveal about how we, as makers, do. But today in creative practice, AI is mostly used not to look closer, but to go faster, or even to automate the creative process altogether. In this thesis, I develop an alternative approach: a practice called Reflective AI, which slows the technology down and uses it not as a tool for production, but as a site for reflection. It treats often-overlooked stages of AI practice (from collecting data to training a model) as a material to be worked, the way a ceramicist works clay.

This approach is developed across five projects that range from investigating object detection models to a three-month ceramic residency. Throughout, I use machine learning as a mirror: I train models on my own subjective judgments, and then study how the models reflect my ways of seeing back to me. What emerges is a material gaze: a way of seeing that comes from the making of and with the model, and looks back on the maker. This thesis offers a way past the question of whether AI is good or bad for creativity, toward a slower, more considered relationship with it.