

4 Envisioning and Designing the Future

The holy grail for generalized AI is to achieve humanlike intelligence. We have humanlike intelligence, it's called humans! [...] [A]nd what we really need to solve today's problems and tomorrow's problems is superhuman intelligence. So, and I believe the only way to achieve that, at least today, is with humans in the loop. (Michelucci 2019a, 11:22–11:47)

In his talk at the Microsoft Research Faculty Summit with the title “Crowd, Cloud and the Future of Work: Updates from human AI computation,” Michelucci argued for HC as an alternative to AGI. Instead of pursuing “strong AI” or AGI efforts to build “human-like intelligence” with machines, HC aims to keep “humans in the loop” to build “superhuman intelligence.” Michelucci, who I regard as an HC advocate, together with colleagues at the Human Computation Institute follow such approaches to human-in-the-loop computing by building HC-based CS systems. For this pursuit, Michelucci described going back to the “biggest pictures [he] can think of” (Jan. 21, 2021). The biggest picture for him was the understanding of information processing, which forms the basis of his HC visions and enables the development of HC systems for “human survival” (Michelucci, Jan. 14, 2021). The institute is, therefore, dedicated “to the betterment of society through novel methods leveraging the complementary strengths of networked humans and machines. We strive to engineer sustainable participatory systems that have a profound impact on health, humanitarian, and educational outcomes” (Human Computation Institute, n.d.). With this dedication, the institute places its ethical mission at the center of all its endeavors and daily working practices to develop “superhuman intelligence.”

While Michelucci's words cited above do not describe what including humans in the loop to build human–AI systems means in concrete terms, they do provide a glimpse of how HC advocates construct the narrative of HC from AGI as a counterpole. Human computation, as I will discuss in this chapter, is not only understood to be *better* and more *ethical* because humans remain in control, but also more feasible. Michelucci's words represent and carry imaginations of the “human in the loop” in HC which not only shape the discourse on HC but also the development of such systems in general, as well as in the field of HC-based CS. Imaginations are crucial and powerful forces in the formation of assemblages and their desire. Therefore, it is important to disentangle these imagi-

nations to further analyze how HC systems form, taking into account the interplay of future visions, everyday negotiations and associated human–technology relations. The latter are fundamental to HC because human-in-the-loop imaginations do not consider humans and AI in HC as separate but as relations between humans and AI. As I will show in this work, the intraversions of human–technology relations are, hence, always imagined.

This chapter, thus, analyzes the imaginaries from the perspective of HC advocates and developers as fundamental parts of the (everyday) becoming of HC-based CS assemblages and their intraverting human–technology relations. I first unpack the ambiguous term “human-in-the-loop” within HC, drawing particularly from the empirical material collected during my fieldwork at the Human Computation Institute. The analysis of this expression reveals different shifting and partly competing imaginations of concepts, such as “humans,” “AI,” “computation,” or “technology” in HC and the broader field of AI. The analysis is structured around six imaginaries related to “humans in the loop.” It should be noted that “human-in-the-loop computing” is a common phrase in computer science and technology development that goes beyond HC. I specifically focus here on its meaning in the field of HC. The first part of this chapter is structured around the individual imaginaries. First, HC forms a counter-imaginary to AGI. This describes the basis of HC endeavors and imaginations. Humans are put at the center of HC achievements in contrast to AGI, where the focus lies on the technological achievements. With humans in the loop, following the imaginary, unprecedented capabilities can be reached that exceed today’s AI capabilities. In this overall counter-imaginary of HC, I describe situated human-in-the-loop imaginations which refer to the imagined *human* in the loop—who is imagined to be in the loop and *how* are humans imagined in HC?—; GWAPs imagining humans as *players* and loops as *games*; the imagined *loop* between humans and AI as human–AI conversations of the future; and the imaginations of humans in the loop as *crowds* in the loop. Finally, HC is imagined to lead to a future *Thinking Economy*, in which human creativity is unleashed.

By analyzing these visions and imaginations, which go hand in hand with normative principles guiding the everyday practices, it becomes clear that the development of HC systems can be understood as “ethical projects” (Ege and Moser 2021a) striving for a *better* future and human–technology relations as they *ought* to evolve (Jasanoff 2015a, 4). Moreover, it becomes apparent that there are not only overlaps between the imaginaries but also tensions between the HC future imaginaries and existing implementations, which can be explained with cultural anthropologist Genevieve Bell’s observation that AI, and in this case more precisely HC, is “both a technology, and an imagination of a technology” (2021, 4).

The analysis of the imaginaries and visions necessarily decontextualizes them to a certain extent and separates them from everyday practice. This is due to my attempt to juxtapose concrete imaginaries with the example of the Human Computation Institute and general perspectives found in literature and the development of the field of HC.

For this reason, the second part of this chapter analyzes examples of infrastructuring that go hand in hand with the imaginations at the Human Computation Institute. Infrastructuring, as I will show, not only materializes these imaginaries but also destabilizes and forms them. Within the realm of HC, these practices can be observed at dif-

ferent levels. On the one hand, it is part of the bigger goal of establishing HC and the vision of hybrid human–AI systems, which I describe as *fundamental* infrastructuring. Here, I discuss the example of Civium, an “eco-system” which is being developed to build a basis for more “sustainable” HC (Michelucci 2019b). On the other hand, infrastructuring is part of everyday practices of maintaining and developing technological systems in general. I provide examples of what everyday materializing and negotiating human-in-the-loop imaginaries looks like from the perspective of the team, using the case of Stall Catchers. However, this discussion of infrastructuring continues throughout the next chapters with a special focus on the development of the Stall Catchers’ data pipeline in Chapter 6, since exploration and infrastructuring play a major role in forming, re-defining, and materializing human-in-the-loop imaginaries in HC. Because, while the ultimate goal seems to be clear, HC-based CS systems, according to the imaginaries described in the following, must stay at the edge of AI developments and scientific research. Furthermore, the role of humans and technology in these sociotechnical assemblages is not set *a priori* but negotiated along the way, guiding the future development of HC systems and further (re-)forming the visions. Developing HC systems, therefore, depends on future imaginations of human–technology relations (building on the counter-imaginaries), practices of exploration, and path dependencies (Klausner et al. 2015; De Munck 2022) introduced in the past. At the same time, and this will be the focus of the next chapters of this work, imaginaries are negotiated and partly also contested and contingent in the everyday. Inoperable materialities and life cycles of mice unaligned with the research process form or reform these relations just as much as other actors, most prominently the participants, who bring in their own aims and values, and engage in their chosen ways in relations with technology.

The aim of this chapter is to reconstruct the imaginaries driving the development of HC, in a manner akin to Bachman’s analysis of the “work on the medium” with the example of Dynamicland (2018). These depictions are presented in a way that is aimed at closely reflecting and analyzing how they were presented by HC advocates, to then be contextualized and critically contrasted with other perspectives in later sections of this work.

My close collaboration with and ethnographic fieldwork at the Human Computation Institute allowed me to thoroughly analyze the imaginations driving the institute and provide a thick interpretation of a concrete example of narratives and normative principles that form and steer the current development of HC systems. While some of the points to be elaborated came to the fore due to questions I asked in interviews, others emerged, for example, from work meetings for brainstorming new features or project ideas or discussing current issues and problems with existing projects. Conversations, particularly during the field research periods in Ithaca, NY, in which I shared my ethnographic observations with Michelucci and other team members, sparked reflections on and comparisons of our perspectives, the questions we were asking in our respective work, and the approaches we took to arrive at answers. It was in these discussions that the imaginations underlying the development of HC became apparent.

The focus on the Human Computation Institute’s specific approach to HC, thereby, is especially interesting due to its striving for the development of HC systems generally and not primarily to the solution of one particular scientific problem, as in the exam-

ple of Foldit, which was specifically developed to solve the problem of protein structure prediction and design. Insights from other empirical material, for example, from interviews with HC advocates such as Bry, will be included in the following to contextualize the institute's imaginaries and further explore and problematize HC visions. Additionally, without claiming completeness, I include literature and HC standpoints presented in media content from the field of HC (and HI) to situate these imaginaries further.

As discussed in Chapter 2, I build upon Jasanoff and Kim's term "sociotechnical imaginary" (2009; 2015) as well as Hilgartner's notion of "vanguard visions" (2015) to think about HC as a counter-imaginary. The aim is to stress how HC advocates are forming their visions in contrast to AGI through boundary work and by explicitly distancing themselves from such AGI endeavors. The concept of sociotechnical imaginaries forms a helpful starting point for analyzing how HC-based CS projects are envisioned due to its focus on the imaginaries' normativity and the materialities that are part of sociotechnical assemblages (Jasanoff 2015a, 19). Some of the imaginaries specific to the Human Computation Institute can even be understood as normative principles that guide the team's work in maintaining existing and developing new projects. However, if we take Jasanoff and Kim's definition seriously, HC imaginations, as opposed to AGI narratives, would be more accurately described by Hilgartner's term "vanguard visions." These visions originate from "sociotechnical vanguards" who are

relatively small collectives that formulate and act intentionally to realize particular sociotechnical visions of the future that have yet to be accepted by wider collectives, such as the nation. These vanguards and their individual leaders typically assume a visionary role, performing the identity of one who possesses superior knowledge of emerging technologies and aspires to realize their desirable potential. Put otherwise, these vanguards actively position themselves as members of an avant-garde, riding and also driving a wave of change but competing with one another at the same time. (Hilgartner 2015, 34)

While Hilgartner observes that visionaries or groups also compete and support slightly different views for the field of synthetic biology, I obtained similar results in the analysis of HC visions. Vanguard visions have to move and legitimize themselves between powerful sociotechnical imaginaries and, hence, cannot be successfully established solely by argumentative means (Hilgartner 2015, 45). Instead, "hands-on activities are needed to present their futures as feasible and achievable" (Hilgartner 2015, 45).

Among HC's sociotechnical vanguards is Michelucci, whose institute is cited as a source for new HC developments (Lazar, Feng, and Hochheiser 2017, 435–436). The initiative "2014 Human Computation Roadmap Summit," led by Michelucci, earned him and his colleagues the title of "crowdsourcing pioneers, and visionaries" in the *MIT Technology Review* magazine (Emerging Technology from the arXiv 2015). Due to Michelucci's prominent position in the field, focusing on the imaginaries at the Human Computation Institute is a fruitful point of departure for the analysis of HC visions.

The narratives and visions presented are not necessarily shared by all HC and HI practitioners or visionaries. Focusing on a few sociotechnical vanguards, however, allows one to gain a deeper understanding of HC visions, imaginaries, and their directions and sci-

entific programs, because, as Koch writes following Ulf Hannerz' understanding of culture, "[i]ndividuals [...] as producers of culture play a key role in cultural analysis" (Koch 2005, 26). They "are constantly inventing culture or maintaining it, reflecting on it, experimenting with it, remembering it or forgetting it, arguing about it, and passing it on" (Hannerz 1992, 17).

Moreover, HC is becoming an increasingly "organized field of social practices" (Appadurai [1990] 2002, 50; cited in Jasanoff 2015b, 327) that can currently be situated between vanguard visions and widespread sociotechnical imaginaries. Seeing HC as a form of counter-imaginary to AGI helps me to analyze the specifics of its visions and their emergence. I show how these visions not only distinguish themselves from AGI sociotechnical imaginaries but are also related and situated in common perspectives of AI research. After all, as Hilgartner writes, "[s]ociotechnical vanguards seek to make futures, but (to paraphrase Marx) they cannot make them simply as they please; they do not make them under self-selected circumstances but do so using vocabularies and practices already given and transmitted from the past" (2015, 50). To a certain degree, these vanguard visions guide the evolution of HC systems and, therefore, of intraversions of human–technology relations (by which they are also produced), even though they are also contested and negotiated.

Human-In-The-Loop Imaginaries

Human Computation as a Counter-Imaginary to Artificial General Intelligence¹

At the heart of the human-in-the-loop imaginations of HC are humans, who are put in the spotlight when talking about the achievements of HC systems. While this may seem trivial at first, it is one of the key differences between HC and many strong AI narratives. The latter claim technical achievement and commonly gloss over human involvement altogether or even imagine humans to be out of the picture. For example, it is still usually not mentioned that most AI models have been trained on data annotated by humans, that datasets have been manually cleaned beforehand, or that rules and rewards, in the case of reinforcement learning, had to be set initially; not to mention the adjustments and maintenance work that guide the lifetime of any sociotechnical system. In contrast, HC, one could say, takes the opposite approach in sketching their imaginaries as they tend to hide the system's technical complexity and the AI algorithms operating in the background, and focus on humans when it comes to achievements. This difference is key to the imaginary of HC as a counter-imaginary to AGI. To better illustrate how this unfolds, I first turn to imaginations of AI in general with a recent example.

On May 14, 2022, Nando de Freitas, who is an ML professor and research director at Alphabet Inc.'s AI subsidiary and research laboratory DeepMind, tweeted about Deep-

1 First ideas of this section were presented at the 8th International Working Conference of the Digital Anthropology Commission German Society for Cultural Anthropology and Folklore Studies, "Digital Futures in the Making: Imaginaries, Politics, and Materialities" on September 15, 2022, at the University of Hamburg.

Mind's new AI-agent Gato (Reed et al. 2022): "It's all about scale now! The Game is Over! It's about making these models bigger, safer, compute efficient, faster at sampling, smarter memory, more modalities, INNOVATIVE DATA, on/offline, ... 1!" (Nando de Freitas [@NandoDF] 2022). With "the game is over," de Freitas declared that the search for AGI is over with Gato AI. All that is needed now is scale. The "generalist agent" Gato is presented as the solution to a decades-long search for the path to strong AGI. While it surely presents an interesting development in ML and reinforcement learning, the purpose of mentioning de Freitas' tweet does not lie in the question of whether Gato presents the straight pathway to AGI or not. Instead, it serves as an example of a very current manifestation of the beliefs in AGI.

While the latter have been particularly pursued since the early 1950s (Fjelland 2020, 2), imaginations of future machines with "intelligence" have long existed in literature and mythology and even before the term AI itself was coined in the 1950s (Koch 2005, 9; Cave and Dihal 2019). But since then, researchers have focused on developing AI, often striving for AGI with "human-like intelligence" (Fjelland 2020, 2). Despite various attempts to define AI, it, similar to its natural model, cannot be clearly defined (Koch 2005, 48). Koch (2005) has already aptly demonstrated this in her research on the "culturality" of the technological formation of AI in Germany in the 1980s and 90s. Nevertheless, this is still true today, where the term AI is used to refer to various technologies from robotics to NLP and DL and where there are different understandings of what the ultimate goal of AI research should be. One of the most prominent endeavors is the development of AGI or "strong AI." Some advocates of the AGI thesis argue that AGI, utilizing its "human-like intelligence," will subsequently lead to Artificial Super Intelligence (Kurzweil 2005; Bostrom 2016; cited in Peeters et al. 2021, 220).²

Even though the implementation of actual AI systems has been lagging, it has been attracting lively interest in public discourse since its early years. Koch attributed the widespread interest in AI to the fact that it challenges humans to reconsider how they perceive themselves in a way that few other technologies do (2005, 48–49). The "scientific attempt to technically reproduce abilities that until then had been regarded as originally human is perceived by many people as a threat" (Koch 2005, 49), leading to an extensive social debate in which participants from various backgrounds contribute (Koch 2005, 49). Similar observations can be made today—still and to no lesser extent—for discourses in, at least, many European countries and the USA.

The journey of AI research has seen ups and downs since its uptake in the 1950s. An example of the latter was the "AI winter" of the mid to late 1980s, which emerged due to overoptimistic expectations of AI's possibilities from AI sponsors, industry, and government in the face of a "failure to deliver systems matching these unrealistic hopes, together with the accumulating critical commentary" (Nilsson 2010, 408). However, with the increase in computing power and advances in, for example, ML and NLP, as well as recent developments, such as the text-to-image model DALL-E and others, AI—and increasingly the discussion of its societal and ethical implications—remains a hot topic of public debate. The "AI winter" has been overcome, and AI is becoming more pervasive in everyday life, demonstrating capabilities in ways that were not expected of these

2 Other scholars divide AI research into cognitivist or engineering approaches (cf. Koch 2005).

approaches. But as promising as these advances are that let AI advocates dream and announce that “the game is over,” as Nando de Freitas did, predictions as to when the idea of AGI will materialize have not yet been fulfilled, and it is highly controversial as to what extent such an AI is even possible.

Artificial intelligence researchers Marieke Peeters et al. (2021) identify three main perspectives in the current broader debate about the future of AI. While the first, the *technology-centric perspective* resembles the endeavors described above, they identify the *human-centric perspective* as the second perspective, “which holds that true intelligence can ultimately be found only in human beings and (potentially) other intelligent living creatures” (Peeters et al. 2021, 219). While AI can provide assistance to humans, it cannot develop “intelligence” itself. The third recognized perspective is the *collective intelligence perspective*, which overlaps with HC and is discussed further below. The field of AI has been accompanied from its inception by narratives and imaginaries that commonly refer to visions of strong AI. Regardless of whether they ever become “true” or not, such narratives—and this is the key point I would like to make here—shape both the expectations and beliefs according to which AI is evaluated, designed, and developed (see also Cave and Dihal 2019, 74). These both utopian and dystopian AI imaginaries are strongly shaped by fiction and politics (Bareis and Katzenbach 2022). Philosopher Stephen Cave and science communication scholar Kanta Dihal (2019) identify four pairs of the most widespread hopes and related fears of AI from a corpus of fictional and nonfictional narratives: these pairs are immortality vs. inhumanity, ease vs. obsolescence, gratification vs. alienation, and dominance vs. uprising (Cave and Dihal 2019). Each hope or fear is connected to a number of either optimistic and utopian or pessimistic and dystopian narratives (Cave and Dihal 2019, 75). These narratives are, of course, not exhaustive but form composing parts of the sociotechnical imaginary of AI that currently prevails in Western countries. Formulating these hopes and fears as pairs illustrates how sociotechnical imaginaries, no matter how widespread, also always face resistance and conflict. This is particularly true in the field of AI, where the promises and their long-term narratives and visions have not been sustained and where there is a mismatch between AI’s imaginaries and current implementations (Sartori and Bocca 2022).

Examples for such alternative imaginaries of AI are the vanguard visions of HC. The advocates of HC build their narrative in an active distinction to dystopic AI narratives and AGI visions with “boundary work” (Gieryn 1983). These HC narratives and visions build upon the limitations of today’s mere machine-based AI efforts and the disappointments due to unfulfilled promises of such strong AI visions. Sociologist Jens Beckert describes this pattern regarding new promissory stories in general, which always relate to preceding ones and create their own credibility by removing themselves from those “disappointed hopes” (2018, 284).

In a promotional video by the Carlsberg Foundation about HI at the Center for Hybrid Intelligence at Aarhus University (Elmann 2022), physicist and head of the Center for Hybrid Intelligence Jacob Friis Sherson explains:

When you watch the news, then you get the experience that artificial intelligence is a tidal wave that is coming upon us, so every week there is a new instance of how artificial intelligence has entered and transformed a new domain. And you get the

experience that it will overcome us, that there will be no place in the near future for humanity. And that's what I want to fight, because the reality is the more we try to build into artificial intelligence, the more we see that there is a core of us that cannot be replaced. (Carlsbergfondet 2019, 00:19–00:53)

As Sherson speaks, the camera alternates between him sitting in front of a large screen and different scenes lending a dramatic visual shape to his words, including an angry-looking robot army, a hand taking a queen in a chess game, a fast-running clock, and a skull surrounded by white fog. To construct the narrative, HI takes off from the fears of current AI sociotechnical imaginaries that are circulating in the media and builds on the dystopian images of these imaginaries.

While not all HC advocates aim to combat dystopian strong AI narratives in the same way as Sherson, their arguments share the common notion that the promises of AGI to “supersede humans” (Dellermann, Ebel, et al. 2019, 637) have not been fulfilled to date, and it remains unclear whether these promises will be achieved in the not-too-distant future. Law and Von Ahn write in the preface to their book *Human Computation* (2011) that HC moves beyond the extent of current AI algorithms. Despite advances in automating mathematical problems humans used to solve, the narrative goes, “there are still many problems that are easy for humans to solve, but difficult for even the most sophisticated computer algorithms” (Law and Von Ahn 2011, 2). Human computation, therefore, starts here, at so-called AI problems and by combining humans and machines they can “achieve futuristic AI capabilities today” (Michelucci and Dickinson 2016; cited in Michelucci 2019b).

One of the key differences between HC and strong AI narratives was mentioned at the beginning of this section: the attribution of claimed achievements. While Strong AI narratives claim technical achievement, HC puts humans in the spotlight when talking about the achievements. In the promotional video, Sherson argues that HI is not about “understanding what AI can do and cannot do, it is understanding what we as humans are best at” (Carlsbergfondet 2019, 01:03–01:08), and the challenge is to find the “human place [...] alongside AI” (Carlsbergfondet 2019, 00:58–01:00). By distancing itself from AGI sociotechnical imaginaries and visions, HC builds its own credibility and aims to legitimize the approach as the *right* way to develop “intelligence” that goes beyond humans and computers’ capabilities by putting humans into the focus. This allows HC advocates to circumvent the dystopian narratives and fears in public discourse associated with strong AI pursuits and to present HC as “ethical projects” (Ege and Moser 2021a). They are “future-oriented undertakings” (Ege and Moser 2021a, 7) that are permeated by values and norms of the roles that humans should take and what “artificial intelligence” should look like. This is evident in the choice of words, such as “conscientious development,” chosen by Bowser and colleagues in the report “Artificial Intelligence: A policy-oriented introduction,” published by the US research center and nonpartisan political forum:

Human computation approaches to AI [...] advanc[e] the design of AI systems with human stakeholders in the loop who drive the societally-relevant decisions and behaviors of the system. The *conscientious* development of AI systems that *carefully* con-

siders the coevolution of humans and technology in hybrid thinking systems will help ensure that humans remain ultimately in control, individually or collectively, as systems achieve superhuman capabilities. (Bowser et al. 2017, 11, emphasis LHV)

It is stressed in this statement that humans ought to ultimately remain in control of these hybrid systems which can be accomplished with conscientious development that takes into account the “coevolution” of humans and technology. The understanding of human–technology relations put forward in this quote is reminiscent of a postphenomenological approach that considers humans and technologies existing together and being mutually interdependent (Dorrestijn 2012a, 63).

The Human Computation Institute ethically frames its endeavors in a similar way, placing its ethical mission at the center of all its efforts. I would like to illustrate this mission in more detail here before situating it within other HC and HI endeavors. Michelucci explained during one of our more formal interview sessions, when asked if there was anything else important to him that he wanted to share, how his work was rooted in the “biggest ideas,” as cited in the introduction, and how HC could bring humanity closer to the “good” life. “We can use these new capabilities to bootstrap better [...], more humanistic augmentation methods that allow us to solve world problems [...] by doing things that [...] entertain us, they give us a sense of purpose and they just help us thrive as human beings.” (Michelucci, Jan. 21, 2021)

As Michelucci’s words quoted at the beginning of this chapter make clear, HC is understood to achieve unprecedented capabilities today that exceed mere computational capabilities and current AI technologies and ensure humanistic approaches.

Michelucci and computer scientist and HC researcher Elena Simperl, as editors, open the first issue of the *Human Computation Journal* with the following quote that is attributed to the economist and public servant Leo Cherne: “The computer is incredibly fast, accurate, and stupid. Man is incredibly slow, inaccurate, and brilliant. The marriage of the two is a force beyond calculation” (Michelucci and Simperl 2014, 1). Building on the understanding of computers as “fast but stupid” and a human as “slow but brilliant,” the authors set the agenda for the transdisciplinary approach that taps and combines the respective strength of humans and computers to achieve unsurpassed capabilities (Michelucci and Simperl 2014, 1).

However, despite this huge potential, there is a gap between what could be possible with HC and current HC research and implementations. According to Michelucci, the field of HC research today is not innovative enough due to the organization of academia and its inherent dependencies and power relations (Jan. 14, 2021). If more researchers tried new and more innovative approaches to combining humans and machines, it would not only advance HC and bring it closer to solving some of today’s biggest problems but also make it more “interesting” for participants, he explained in one of our conversations (Jan. 14, 2021). Michelucci described this as the “duality” in HC. An HC project must be both interesting and “fulfilling” for participants (Michelucci, Jan. 14, 2021).

With his vision of HC, Michelucci might not necessarily represent the perspective of all HC researchers. However, as I learned during my ethnographic fieldwork at the Human Computation Institute, it is also shared with great enthusiasm by his colleagues. The CS coordinator Paul explained in our interview in October 2020:

I think human computation could do everything [laughs]. [...] I think the [...] only limit is sort of human imagination of what problem you can take and tackle with human computation and, of course, still the [...] existing resistance to nonconventional ways of doing that. Like, you still need to, [it is] still a huge job to convince scientists and founders—like what happened with Stall Catchers, the scientists didn't believe that this could be done but they did later and no funders would believe that this could be done, but then there was one who decided that he will take a risk and fund this nonconventional project so I think that's [...] the only limitations for now to tackle any problem with [...] human computation. (Oct. 14, 2020)

The spirit of working on something new and “revolutionary” drives the work at the Human Computation Institute, such as that of developer Kate, who explained that working on this project opened their eyes to the relations between humans and machines:

I'm really excited to work on a project [...] like that. I really gained a lot, learned a lot as well, I'm glad I took that chance actually. It made me think of things in a different perspective. I used to think of humans as separate from computers. I used to think people can do, cannot do what a computer can do or the opposite, but the more I've worked on [different projects at the institute], I started to [gain] a different perspective on things like [...] the line between humans and machine is getting more blurred and maybe in the future, Stall Catchers or human computation would have been the first step in a neural network or beaming people's brain into computers, I don't know about that [laughs], but maybe that's part of the revolution, that's part of the first step of how we do these things and I'm really happy to have been part of this project. (Nov. 19, 2020)

Sharing this excitement, developer Samuel half-jokingly explained to me in another conversation that combining humans and machines in a way that humans oversee machines could indeed lead to unprecedented capabilities: “[B]y bringing these two together, it's like we're having a god, we're building a god. [Laughs] We are creating a new era god” (Sept. 2, 2021). But in contrast to an AGI that will outperform and rise up against humanity, as one of the fears stated which was identified by Cave and Dihal as prevalent AI narrative prophecies (2019), HC, according to Samuel, presented

a big opportunity. It [...] can prove that AI will [...] help humans instead of replacing humans in the future. And HCI [the Human Computation Institute] is proving that AI can help humans instead of replacing them. because AI still has some disadvantages, still has some issues. It's not like we're all in [...] a science fiction movie, where AI is destroying everyone. [...] HCI has an opportunity to become, [...] the best, to become a role model for people I mean in the AI section, AI development. (Samuel, Sept. 2, 2021)

Human-in-the-loop computing in HC is imagined as leading to *good* sociotechnical systems, in which human control is ensured. Due the institute's dedication to the “betterment of society” (Human Computation Institute, n.d.), Samuel explains, it could become a role model for other AI researchers and developers (see above). The mission of the in-

stitute illustrates the ethical motivation behind the projects and work at the institute to develop HC systems to develop *better* futures.

While the Human Computation Institute's vision might be very specific and personal, the counter-imaginary of HC and HI as ethical projects can also be found in one of the first HI papers of the current research trend by human-centered AI researcher Ece Kamar (2016a; 2016b). Referring and connecting to HC advances, Kamar argues for "hybrid intelligence"³ systems as AI systems that are guided and supervised by humans to "prevent the mistakes and failures that would be caused by an AI system working alone" (Kamar 2016a, 4070). Rather than "designing AI systems that function alone, we should focus on hybrid systems that can benefit from partnership with humans" (Kamar 2016a, 4070), which can lead to a "virtuous improvement cycle for the system to continuously learn from" (Kamar 2016a, 4070). These, according to Kamar, can go beyond the current limitations of mere machine AI (2016b, 25). Kamar's HI definition clearly has great similarities with common definitions of HC. However, Kamar specifically refers to HC as crowdsourcing and emphasizes the fruitful environment of such crowdsourcing or HC platforms since they "function as testbeds for data collection and experimentation" (2016a, 4070) for accessing "human intelligence."⁴

Dellermann et al. (Dellermann, Calma, et al. 2019; Dellermann, Ebel, et al. 2019) further define HI as "socio-technological ensembles" in which both the human and AI "continuously improve by learning from each other." (Dellermann, Ebel, et al. 2019, 640) The paper "Mapping Citizen Science through the Lens of Human-Centered AI" (Rafner et al. 2022), which was initially published as a preprint with the title "Revisiting Citizen Science Through the Lens of Hybrid Intelligence" (Rafner et al. 2021), has been written by multiple authors and visionaries in the field of HC. It can, thus, be understood as an attempt to come to a shared understanding of the fields despite these "incompletely aligned views" (Hilgartner 2015, 35), for example, regarding the definition of HI, and what differentiates HC from HI. Divergent understandings are not uncommon for new or young research fields that have not stabilized into one prevalent narrative that most advocates can agree

3 While Kamar's article definitely describes an early take on HI, it is not clear when the term was first coined.

4 For Kamar, crowdsourcing platforms are "testbeds" for HI (2016a, 4070). Including the crowd in online CS projects also creates the starting point for exploring HI at the Human Computation Institute and the Center for Hybrid Intelligence at Aarhus University (Elmann 2022). In an article in the *Human Computation Journal*, the CS researcher and ecological informatics specialist Greg Newman elaborates on the opportunities of "Citizen Cyberscience" for HC: "Some citizen science projects introduce new human computation techniques or engagement modalities, thus directly contributing to a growing body of human computation methods. In this way we can see citizen science as applied human computation, a platform for human computation research, and a body of work that may innovate in the human computation space. As an example, citizen science often generates platforms for citizen engagement that can be used in human computation research" (2014, 108). Platforms such as Stall Catchers, Foldit, and ARTigo can, hence, be understood as laboratories for HC research. The game environment and framing of the project also plays an important role in the intraversions of human-technology relations due to affordances of the platforms as games, for example, which invite participants to play around and find ways to engage with the platform and game that go beyond the developer's intended designs.

on. Nevertheless, HC and HI visions seem to be able to agree on several shared understandings, the most prominent of which are the imaginations of alternative futures of AI to widespread AGI and strong AI sociotechnical imaginaries that place humans at the center of their approaches. While acknowledging that discussing these narratives collectively can result in oversimplification and generalization, examining them as an interconnected yet diverse bundle can provide a deeper understanding of the nuances and specificities of both HC and HI. As Forsythe has aptly put it regarding her research on informatics and AI:

While the beliefs held by actors within a given scientific setting or community are neither identical nor altogether stable over time, neither are they completely reinvented from one event to the next. Scientists construct local meanings and struggle to position themselves within a cultural context against a backdrop of more enduring beliefs, at least some of which competing actors hold in common. After all, in order to negotiate at all or successfully to frame actions or objects in a given light, there must be overlap in the interpretive possibilities available to the parties concerned. ([1992] 2001b, 2–3)

The overlap in HC lies specifically in the attempts to imagine AI futures with humans at the center and in control. While this utopian perspective is reminiscent of traditional techno-optimist imaginaries, according to which automation and AI will free humans from undesirable work, the HC imaginary as a counter-imaginary to AGI endeavors focuses on the “co-evolution of humans and technology,” so that the former retain control over the latter. “In-the-loop” refers not only to *being* in the loop but also to being *in control* of the system, as opposed to other uses of this term, such as in crowdworking applications, where humans are in the loop in the sense of directly training an ML model by annotating data.

Shared Paradigms

Even as advocates position HC as a counter-imaginary to other AI pursuits, they remain open to the possibility that AGI may become possible at some point in the future. This is evident in Michelucci’s words quoted above, in which he states that the only way to achieve “superhuman intelligence [...], at least today, is with humans in the loop” (Michelucci 2019a, 11:40–11:47). The insertion “at least today” indicates that Michelucci does not completely distance himself from the idea that AGI might someday be possible, similar to Dellermann, Ebel, and colleagues, who speak of the “near future” (2019, 637).

Despite the boundary work of HC to distance itself from AGI visions, the visions share much in common. Following the cultural studies scholar Toby Miller, “[e]very cultural and communications technology has specificities of production, text, distribution, and reception. But the utopias and dystopias of successive innovations share much in common. As private excitements and public moral panics swirl, they also repeat” (2006, 6). Although HC refers not only to previous innovations but also concurrent developments in AGI endeavors, Miller’s observation about the commonalities of utopias and dystopias of successive innovations can also be observed for these fields. This is not surprising given that many HC researchers and advocates commonly have a background in

computer and/or cognitive science. To understand the intentions behind the design of HC systems, it is, nevertheless, important to discuss these parallels.

I here want to briefly mention two of the perspectives and paradigms that form parallels between HC and strong AI approaches. First, imaginaries of AGI (both utopian and dystopian) and HC have in common that they view these developments (AGI and HC) as having an impact on humanity as a whole. They describe phenomena that do not only solve specific problems but have a large-scale influence on humanity and its problems in general. In the case of HC, this can be seen in the example of Michelucci's vision and in the quotes that introduce the new *Human Computation Journal*. In their introduction, Michelucci and Simperl not only quote Cherne but also founder of evolutionist theory Charles Darwin (1859): "In the long history of humankind (and animal kind, too) those who learned to collaborate and improvise most effectively have prevailed" (cited in Michelucci and Simperl 2014, 1). By invoking Darwin and Cherne in this way, they frame HC as a revolutionary phenomenon with possibly existential implications. This way of framing parallels popular imaginations of the rise of AGI, for example, as physicist, cosmologist, and president of the Future of Life institute⁵ Max Tegmark writes in *Life 3.0. Being Human in the Age of Artificial Intelligence*: "In comparison with wars, terrorism, unemployment, poverty, migration and social justice issues, the rise of AI will have greater overall impact" (2017, 37–38).

Second, both HC and AI approaches generally build upon information processing as their fundamental paradigm. Information processing, which interprets cognitive processes as processes of information processing, is considered the key paradigm underlying AI research (Ahrweiler 1995, 15; cited in Koch 2005, 45; see also Simon 1996; Turkle 2005a; Becker 2023). It is also fundamental to HC, as Michelucci writes:

[T]he construal of computation as being equivalent to information processing seems to best fit the practical context of human computation. In HC, "computation" refers not just to numerical calculations or the implementation of an algorithm. It refers more generally to *information processing*. This definition intentionally embraces the broader spectrum of "computational" contributions that can be made by humans, including creativity, intuition, symbolic and logical reasoning, abstraction, pattern recognition, and other forms of cognitive processing. As computers themselves have become more capable over the years due to advances in AI and machine learning techniques, we have broadened the definition of computation to accommodate those capabilities. Now, as we extend the notion of computing systems to include human agents, we similarly extend the notion of computation to include a broader and more complex set of capabilities. (2013d, 84, emphasis i.o.)

Though Michelucci argues for a broad understanding of cognitive processing (see below), the examples provided still refer primarily to abstract, mental processes.

5 The Future of Life institute is a nonprofit organization with the mission of "steering transformative technology towards benefitting life and away from extreme large-scale risks" (Future of Life Institute, n.d.).

This understanding of information processing can be contrasted with Hutchins' understanding of computation and distributed cognition.⁶ Even though Hutchins' *Cognition in the Wild* was published in 1995 and his critique of the cognitive science approach might not be applicable to all current understandings, his elaboration on the problems of approaches that focus on information processing seems relevant here:

The model of human intelligence as abstract symbol manipulation and the substitution of a mechanized formal symbol-manipulation system for the brain result in the widespread notion in contemporary cognitive science that symbols are inside the head. [...] And while I believe that people do process symbols (even ones that have internal representations), I believe that it was a mistake to put symbols inside in this particular way. The mistake was to take a virtual machine enacted in the interaction of real persons with a material world and make that the architecture of cognition. (1995a, 365)

The problem, following Hutchins, is that an approach to cognition or intelligence as information processing makes it difficult to combine it with action, which is fundamental to an understanding of distributed cognition as process. "History and context and culture will always be seen as add-ons to the system [if cognition is reduced to information processing], rather than as integral parts of the cognitive process, because they are by definition outside the boundaries of the cognitive system" (Hutchins 1995a, 368). Such an understanding also has consequences for the imagination of the human in the loop, as I will discuss below. Although only briefly discussed so as not to go beyond the scope of this research, the two examples show how HC is, to a certain degree, situated within the broader field of AI research and, thus, rooted in the same foundational paradigms and perspectives with endeavors from which it distances itself.

The Imagined *Human* in the Loop

Who is the human in the loop? While the use of the term *human* in the sociotechnical imaginary of HC implies a very broad understanding, a closer look at the human-in-the-loop imagination reveals that the *human* in the loop in HC narratives and literature does not refer to developers or designers of HC-based CS systems but specifically to users and participants who are invited to contribute by the system designers. In the examples of Stall Catchers, ARTigo, and Foldit studied, however, the participants who are *in* the loop are not necessarily those who are *in control*, even if they play an active role in forming the systems, as will be discussed in the next chapters. Instead, the projects and sociotechnical systems have been initiated and created by researchers and developers, who, in the end, have decision-making power over the overall project.

Human-in-the-loop generally refers to a paradigm or approach within computer science research fields such as ML (Monarch 2021; Mosqueira-Rey et al. 2022). The US De-

6 For further discussion of the information processing paradigm, see, for example, Ahrweiler (1995), Koch (2005), Turkle (2005a), Becker (2023). Unlike Hutchins or literary theorist N. Katherine Hayles in her book *Unthought: The Power of the Cognitive Nonconscious* (2017), my research does not seek to define "cognition" but to examine how it is conceptualized in the field I investigate.

partment of Defense defined it in 1998 as a “model that requires human interaction” in their *Modeling and Simulation Glossary* (Under Secretary of Defense for Acquisition Technology 1998). More recently, it has been defined as “an interaction between humans and an artificial intelligence (AI or machine), with the goal of improving the machine’s AI” (Rueckert and Riedl 2022). Even though usually not explicitly stated, human-in-the-loop is understood as a way to model these interactions, which are defined by the designers and developers of such systems or interactions. With *human-in-the-loop*, authors using the term are most often not referring to themselves. They are referring to the imagined users. It should be noted that humans in the loop can also be, for instance, medical doctors using a system in the example of medical AI systems (Kieseberg et al. 2015). During my second research stay in Ithaca and when discussing this observation of the human-in-the-loop approach, Michelucci agreed: “I’m a human in the loop but not in a way we often talk about it” (fieldnote, Nov. 9, 2022). Rather, the participants in today’s HC systems are those in the loop who might later be replaced by automated processes as computational capabilities advance. Designers, developers, and researchers set the terms and decide what process should be automated and where participants should be invited to solve a specific task. In this sense, human control over AI and hybrid systems currently does not typically take place “in the loop,” but as control “of the loop.” Therefore, to better understand sociotechnical assemblages, it is crucial to consider the role of the designer and developer in the development and evolution of HC-based CS projects. Only then is it possible to contextualize the imagined *human* in human–technology relations of the future and analyze how the imagined humans actually relate to these. That is, how they actively engage and bring in their own ideas of how they want to be human in human–technology relations in HC-based CS.⁷ This observation of the human in the loop also implies not only a relational but also a processual perspective, focusing on both the becoming of the assemblage and the continuous processes of territorialization, reterritorialization, and deterritorialization, which are “mutually enmeshed,” following Deleuze and Guattari: “It may be all but impossible to distinguish deterritorialization from reterritorialization, since they are mutually enmeshed, or like opposite faces of one and the same process” (Deleuze and Guattari 2000, 258). This chapter, therefore, discusses both the imaginaries of HC advocates and the work of designers and developers on materializing them, while, at the same time, changing the imaginaries.

The imagination of *who* the human in the loop is already alludes to *how* humans are imagined by HC researchers and the projects that they create. The humans in the loop are not imagined as the researchers or experts in HC. But, in collating their nonexpertise responses, the information they provide can be computed into “expert-like” information or output. At the same time, “human intelligence” (Law and Von Ahn 2011), “cognition” (Michelucci et al. 2015, 2), “human intellect” (Hartman and Horvitz 2013, xi), and the “human processing power” (Von Ahn 2005) have remained central to the image of the human since

7 In fact, following a relational understanding of technology or Louise Amoore’s approach to “cloud ethics” (2020), “the human in the loop is an impossible subject who cannot come before an indeterminate and multiple *we*” (2020, 66, emphasis i.o.). This implies that we must focus on the relations and entanglements of the various human and nonhuman actors, which I attempt to do in this work.

the very beginning of HC, as I have shown in the previous sections. The idea of human processing power is linked to the broader idea of information processing, which does not only refer to computers, but applies to computation in general, which is performed by a wide range of organisms, including humans, animals, and even bacteria (Michelucci 2017a). According to Michelucci's introduction to the *Handbook of Human Computation*, computation, as understood in HC, includes computer algorithms and numerical calculations but, at the same time, "embraces the broader spectrum of 'computational' contribution that can be made by humans, including creativity, intuition, symbolic and logical reasoning (though we humans suffer so poorly in that regard), abstraction, pattern recognition, and other forms of cognitive processing" (2013b, xxxviii). Human computation then refers to such information processing processes "that derive from the computational involvement of humans in simple or complex systems" (Michelucci 2013b, xxxviii). Michelucci elaborated on his understanding in one of our interviews in early 2021 as follows:

[W]hen I think about human computation [...], I'm really thinking about information processing in the large. We have a lot of different ways to do information processing. We have human humans. We have other animals. [...] [W]e have [...] other kingdoms that can do information processing. And then we have inventions. We have machines we've created that help us do information processing. So, I'm always interested to look at this entire collection of things and [...] this is our tool kit, our tinker toys that we can build with. What can we build with these things to create new capabilities that didn't exist before and that can help us survive. (Jan. 14, 2021)

Michelucci has referred to this holistic approach to thinking as "organismic computing" (2013c) in earlier work.⁸ By considering not only humans in contrast to technology but also nonhuman entities, this approach is reminiscent of naturecultures and more-than-human perspectives in the humanities and social sciences research.⁹

The idea of HC here does not detract from the question of how to automate human tasks with machines or AI, for example, but describes them—including other nonhuman life—as different "information processing" tools that can be combined into something that exceeds the individual capabilities. The combination of these different elements can, in the institute's understanding, make possible the survival of humans, which depends on the survival of the entire ecosystem and world (fieldnote, Nov. 9, 2022). Humans in the loop are, thus, computational elements in HC systems that can solve a specific problem presented, for instance, an image recognition problem in the example of CAPTCHA, Stall Catchers, and ARTigo, that cannot currently be solved by *other* information processors, such as computers. Moreover, humans are imagined as creative problem solvers whose creativity can be freed and enhanced by HC. Finally, humans in the loop in crowd-based HC systems are imagined as aggregate individuals who can be networked together in the

8 In Michelucci's understanding, HC and "organismic computing" are not two separate concepts but "organismic computing is an instantiation of human computation" and "multiple human computation methods can be implemented in" different parts of organismic computing systems (fieldnote Aug. 19, 2022).

9 On naturecultures and more-than-human approaches, see Chapter 1, footnote 3.

loop. Before delving into these aspects, it is important to consider a human or user HC imagination specific to HC systems designed as GWAPs, as in the case of the examples studied. This imagination connects the image of the user with the image of the loop, since in GWAPs, according to Dabbish and Von Ahn's definition, the user is imagined as a contributor who must be entertained and participates due to their interest in playing games, not solving AI problems (2008). These user images play a fundamental role in the design and actual implementation of HC systems in the field of CS, which is why I here want to briefly discuss how and why the examples studied were designed as GWAPs.

Imagining Humans as Players and the Loop as a Game

The examples of Stall Catchers, Foldit, and ARTigo, while addressing very different scientific problems, have in common that they are designed as computer games, and, as such, share several game features. Most prominently, all of the games include a scoring feature so that participants can earn points and compete against each other on the leaderboards. The projects are not accidentally designed as games. In fact, many online CS projects are not game-based but rather designed as more classic scientific studies, where participants fill in surveys or collect data that they upload to a dedicated platform (e.g., the Great Influenza Survey [Land-Zandstra et al. 2016] or iNaturalist [n.d.]). In the field of HC, however, the idea of developing human-in-the-loop systems as games has been part of the HC idea from the very beginning.

Von Ahn introduced “human algorithm games” or the more often referenced term “games with a purpose” in his aforementioned doctoral thesis. Instead of following “traditional approaches” concentrating on the improvement of software to solve problems that computers are not yet able to solve, Von Ahn aimed at “constructively channel[ing] human brainpower using computer games” (2005, 3). Games present a fruitful environment for this goal because, according to Von Ahn, they offer the incentives required to make humans contribute to computational systems (2005, 11). “In every case, people play the games not because they are interested in solving an instance of a computational problem, but because they wish to be entertained” (Von Ahn 2005, 12). Von Ahn thereby builds on previous research, arguing for the success of gamification, or making “work fun,” and the need to design enjoyable UIs but transfers this to using game-design to solve AI problems. These “human algorithm games” could be thought of as algorithms in which humans perform the computational steps. “Instead of using a silicon processor, these ‘algorithms’ run on a processor consisting of ordinary humans interacting with computers over the Internet” (Von Ahn 2005, 11). In his thesis, Von Ahn points out the importance of asking the following question when assessing such games, which is also one of the normative principles guiding the development of HC systems at the Human Computation Institute:

“[C]an a computer play this game successfully?” If the answer is yes, then the game is of little or no utility because the computer could simply play the game to solve the computational problem. If the answer is no, then the game has truly captured a human ability and the time invested by the players may provide a useful benefit. (Von Ahn 2005, 81)

In contrast to the normative principle expressed by Michelucci that it is unethical to have humans perform a task that computers can solve, the question here is more about the “utility” of the human algorithm game for a given AI problem. The understanding that people contribute to such games because they “enjoy the game [...], in turn producing more useful output” (2005, 76) is at the core of Von Ahn’s idea. The ARTigo project, developed just a few years after Von Ahn’s doctoral thesis, built directly on his understanding of GWAPs. This approach seemed particularly attractive for collecting large amounts of image tags, since, as the team had calculated, paying student assistants to do the work would not have been affordable (Bry, Mar. 2, 2020). As a result, ARTigo was designed as a casual game that can be played occasionally, either for a very short time between subway stops or for a longer period of time. In a similar way, Stall Catchers was also designed as a casual game that can be played without much time commitment (Vepřek 2023b). The idea behind Stall Catchers’ game design builds on the user image the Human Computation Institute had in mind when developing the system:

I thought it was going to be 30 something people. I thought it was going to be the casual game that people in the workforce, [...] 20 something, 30 somethings, they’re going to be standing in line at the bank and they’re going to be catching stalls because it gives them a sense of purpose it’s like a kind of fun, casual game. (Michelucci, Jan. 14, 2021)

The team also saw the advantage of designing it as a casual game, rather than a crowdworking task, in that annotating the data in a game format allows one to build up a “self-sustaining community” (Michelucci, Jan. 21, 2021). They viewed the game approach as advantageous due the possibility of educating people: “[I]t feels like we’re making a bigger difference not just in analyzing the data but in educating people, in building a sense of community and purpose around finding a treatment for Alzheimer’s disease” (Michelucci, Jan. 21, 2021).

In the example of Foldit, crowdsourcing the human participants’ contributions in a crowdworking context, also did not seem feasible due to the complexity of Foldit, which would require training crowdworkers to perform the task. Developer Daniel explained in our interview:

[T]here is a huge learning curve. And if you try crowdsourcing this, there is a lot of overhead costs and up-front costs to teaching the crowdsourced workers how to play and [...] how to solve the task. And, in doing so, they are not very motivated to learn this huge thing. Imagine if you were being paid some small amount to learn a different language so that you could translate some paragraphs. That is not very interesting, not very motivating, it seems like a [...] huge cost up-front just to do some task. And so, by framing it as a game, there is a reason, there is an intrinsic motivation to learn how to play, there is a reason to stick around and continue doing it, [...] there is a sense of long-term investment. (Jan. 24, 2020)

A game-design would make the complex matter of protein design—which exceeds the complexity of crowdworking microtasks—more accessible as players would be more motivated to invest time and effort in learning to play Foldit. In addition, as Gidon, a team

member and one of the co-creators of Foldit, explained, the reason for designing it as a game emerged from the possibility of tapping an existing large user base:

[O]riginally the motivation for the game was that people spend a lot of time playing games already and games kind of engage people in solving problems and so maybe we could make a game that takes all of the time and effort that people spend solving problems in games and put that towards a real-world problem. (Jan. 31, 2020; cf. Miller et al. 2019)

This understanding aligns with the GWAP idea formulated by Von Ahn and Dabbish (2008) and reflects the user image the Foldit team had in mind when developing the project. The aim was to address “gamers,” people already spending a lot of time playing games. As Foldit researcher Brian Koepnick explained in an interview with the *German Video Game Awards (Deutsche Computerspielpreis)*, interest in and enjoyment of games can also lead to an interest in science: “Gaming is a great way to get people excited about science. The ‘average citizen’ has a lot to contribute to scientific research, but most people don’t want to read a textbook or even get a PhD. If you can solve a scientific problem through play, it becomes accessible to nonscientists.” (Koepnick 2020)

In a paper in which the authors analyze how learning and motivation frameworks can be utilized to improve player experience in Foldit, game scholar Josh Aaron Miller and colleagues note, based on their work, that Foldit, in fact, attracts participants from both the gaming and CS community (2019, 8).¹⁰

However, as I will show in Chapter 5, for many of the Stall Catchers and Foldit participants I interviewed, it was not the game design that brought them to the projects but the role Alzheimer’s disease played in their everyday lives. Some of the participants with a personal connection to Alzheimer’s disease even rejected the project being called a game. Specifically in the example of Stall Catchers, the team was well aware of this fact and acknowledged it. Nevertheless, and coming back to the imagination of the *human* in the loop and the imagination of the *loop*, it can, thus, be summarized that the human in the loop in HC design in the field of CS is often first envisioned as a gamer or someone interested in playing games who can be incentivized to contribute to the computational systems through entertainment. The *loop* in GWAPs, therefore, is imagined and designed

10 Additionally, GWAPs, as Von Ahn had already noted in his dissertation, are not suitable for every AI or scientific problem. In the example of Foldit, game design seemed to fit naturally with the way protein folding had been addressed in science for a long time. Foldit team member José explained: “[T]he way that we think about proteins and protein folding, it actually lends itself to competition pretty easily. So, the way we think about proteins already is a very, it’s a kind of competitive paradigm. [...] The way that we think about the problem is very much game-like already” (Jan. 22, 2020). One example of the game-like and competitive paradigm is the biennial *Critical Assessment of protein Structure Prediction* (CASP) experiment which has been taking place since 1994 (University of California, Davis, n.d.). Even though it is framed as an experiment, it represents science as a playful approach to the world (Dippel 2020). It is organized as a protein structure prediction challenge and resembles a sports tournament in which participating research groups computationally or experimentally try to come up with structures for certain amino acid sequences (University of California, Davis, n.d.).

as a game. In the field of HC in general, the imagination of the loop further refers to how humans and technology or AI should collaborate in future HI systems.

Human–Technology Conversations of the Future

The role of humans in HC systems is described in the current HC literature as computational contributors (Michelucci 2013d, 84) or assistants for AI systems (Kamar 2016a, 4071).¹¹ Authors such as Quinn and Bederson draw on terms from crowdworking applications to define computer, worker, and requester as the roles that exist in any HC systems (2011, 1410). These roles, although the order in which the roles are performed may differ from application to application, are considered consistent (Quinn and Bederson 2011, 1410). Taking the case of reCAPTCHA, for instance, and focusing on the workers and the computers in Quinn and Bederson's example, the authors describe the relation between computers and workers as follows: “[A] computer first makes an attempt to recognize the text in a scanned book or newspaper using OCR. Then, words which could not be confidently recognized are presented to web users (workers) for help” (Quinn and Bederson 2011, 1410). Here, the worker assists the computer by solving the text that the computer cannot recognize. They are solving the same task. In a different scenario, workers provide image labels that are subsequently aggregated computationally to delete irrelevant labels (Quinn and Bederson 2011, 1410). In this case, the tasks performed by the workers or the computer are different, while they are working together toward the same goal (even though the workers might not be aware of this goal). The point is that these different role assignments and relations are assumed to remain the same throughout the life of a project or system.¹² However, as I will show in later chapters, the roles of humans and the computer, or more precisely software or AI, in HC-based CS are not fixed and consistent but are, instead, constantly in flux and subject to reshaping and rearrangement over time.

Similar to Quinn and Bederson's second scenario described above, Michelucci envisioned humans and computers in the loop as working together to solve different tasks. Human–technology relations—or *the loop*—were imagined as a “conversation” in which the roles and tasks were distributed depending on the respective abilities of the computer and human. Imagining current and future relations between humans and technology, Michelucci explained in one of our meetings, using the example of a new, more complicated version of chess than the common one, that “it’s our scrappiness[,] our ability to [...] very quickly figure out, [...] which part of the search space is not worth exploring” (Jan.

11 It should be noted that developers sometimes also view AI as serving as an assistant to humans, depending on the approach taken, for instance, when a computational problem requires human assistance, as in the examples described here, or when humans receive help from AI technology in solving a problem, such as in the example of AlphaFold in Foldit (discussed in Chapter 6). In either view, developers and designers consider these role distributions to be rather fixed at the system level; I argue, however, that they are typically dynamic and intraverting.

12 Notably, some of the literature on HI disagrees with such a static understanding and advocates for a dynamic understanding of roles, tasks, and relations in hybrid human–AI systems (e.g., Akata et al. 2020; see also Chapter 6).

21, 2021). The loop between humans and computers, he continued to explain, could then be imagined as follows:

[L]et's say I'm a human player of chess and I have one hundred and fifty possible moves ahead of me, I can immediately exclude 90 percent of those because I know they would be pointless or most likely. Intuitively I just know and from all my experience with chess, there's no point. So, I've narrowed it down to maybe five moves. So now if a human does that, the human can go to the machine and say, I think these are the best possible five moves to be considered here. And then the machine can now search out that space, which has been substantially narrowed. It can use its brute force computation to search out possibilities. And the machine can put a bunch of possibilities in front of a human and say, OK, what do you think about these? And the machine says: Here are one hundred possible moves of the human says, no, no, no, no, these seven. Now try these seven and so on, and they can have this sort of conversation where they help each other. (Michelucci, Jan. 21, 2021)

Here, human–technology relations are envisioned as a “conversation” or “partnership” (Michelucci 2017a) in which the human first delimitates the search space for a problem for which the AI then makes suggestions that can again be narrowed down or corrected by the human in a dialogue. In the description above, the human retains decision-making power within their sociotechnical relation with the AI. For “making this all work,” Michelucci argued, it is important that “our automated systems [...] adapt to humans as opposed to getting humans to adapt to our automated systems” (Jan. 21, 2021). This idea of getting the AI to adapt to humans and not vice versa goes beyond the dialogue described above. The humans in the loop should decide how they want to engage in such systems, for example, by dancing or writing poetry, and to get there, “we just need to figure out the mappings, I think, somehow,” Michelucci explained (Jan. 21, 2021). Such an approach to hybrid human–AI systems, according to Michelucci, differs from an approach that focuses on “information processing efficacy,” in that it follows a human-centered and “ethical” approach:

[W]e often [...] look from the standpoint of information processing efficacy, so we say, OK, well, what are the computers best at? What are the humans best at? And then how do we combine those complementary strengths in the most effective ways? But what's most effective isn't always what's most enjoyable for a person. There are organic aspects to this, and you don't have to motivate a machine-based computer, but you do have to keep a human interested somehow if you're going to be ethical about it and [...] I think those things are very intertwined. [...] I mean, if the fundamental question is about how do we get humans and machines to work best together, I think that's evolving because the machines are getting more human-like in the things they can do. So that's a moving target. I mean, we're always going to be reengineering that as we discover new techniques for AI, for example, and as we discover new techniques for getting humans to work together and humans and AI to work together as well. (Michelucci, Jan. 14, 2021)

The “ideal” human–AI or human–technology relations in HC systems are not always the most efficient or effective solutions. Allowing humans to decide and set their own terms for how they want to be involved in such sociotechnical systems is, in this sense, described as an ethical choice. I, therefore, consider the development of HC systems to be an ethical project in which “beneficial” takes precedence over engineering/technical optimization. Moreover, as Michelucci explained, what is considered the best combination of humans and machines is constantly evolving. Acknowledging that there is a gap between the future imaginations of a self-chosen human–computer conversation and current instantiations of HC and its human–technology relations, Michelucci concluded with: “So, it sounds kind of crazy and how could that ever work? How could writing a poem solve a problem? But I actually think it’s not so far-fetched” (Jan. 14, 2021).

Considering HC-based CS systems and how humans are currently involved in projects like Stall Catchers, there seems to remain a discrepancy between the idea of placing humans—and their primarily human attributed ability, creativity—at the center of these hybrid systems and the need to abstract problems and the tasks of humans in HC systems in a machine-interpretable way. This abstraction, however, almost inevitably leads to the alienation of people and to the replaceability of the individual. This reduction of the human seems to form a dilemma of crowd-based HC-based CS projects that the Human Computation Institute team is aware of. In an attempt to solve this dilemma, if only partially, they created the category “The humans of Stall Catchers” on the institute’s blog where they present interviews and stories of participants and team members (Human Computation Institute, n.d.b). While there are only ten portraits of participants at the point of writing, these introduce, for instance, the participant’s background and motivation for contributing to Stall Catchers.¹³

Linked to this imaginary of a conversation between humans and AI at the Human Computation Institute is the self-imposed “oath” or normative principle that guides development and working practices (Michelucci and Egle [Seplute] 2020). The oath is adapted from the Hippocratic Oath, which was named after the ancient Greek physician Hippocrates, although its root and authors remain uncertain (Shmerling 2015). “It represents a time-honored guideline for physicians and other healthcare professionals as they begin or end their training. By swearing to follow the principles spelled out in the oath, healthcare professionals promise to behave honestly and ethically” (Shmerling 2015). At the institute, this oath is customized to “First use no humans” (Michelucci and Egle [Seplute] 2020). In a blogpost, the institute discussed an ML challenge it organized together with the data science crowdsourcing platform Driven Data (DrivenData Labs, n.d.) and the company and MATLAB¹⁴ creator MathWorks (The MathWorks, Inc., n.d.) to see whether ML engineers could develop ML models to solve the analysis task currently performed by humans in Stall Catchers. In this context, the post described the adapted Hippocratic Oath as follows:

13 By contrast, in the example of the ARTigo project, participants do not even have to register to contribute and are, thus, sometimes not identifiable beyond their IP address.

14 MATLAB is a programming language and platform for programming and numeric computation.

[I]f there is a way to solve a problem using machines, we should not ask a human to do it. [...] This might seem counter intuitive given that our mission is “*dedicated to the betterment of society through novel methods leveraging the complementary strengths of networked humans and machines*”. But we believe that sometimes action with the best of intentions can have costs that outweigh the benefits. If there is a job that machines can do, we think it would be unethical to waste volunteer human cognitive labor on that job, when there are other, more pressing societal needs that require the unique mental faculties of the magnificent human mind. (Michelucci and Egle [Seplute] 2020, emphasis i.o.)

The aim is not to develop *any* kind of conversation between humans and AI. Instead, humans should only be part of HC systems when the problem cannot be solved by computational systems. This oath reveals the contradiction inherent in HC, that is, the commitment to an ethic that values human labor on the one hand, and the simultaneous treatment of humans as a functional element on the other. Nevertheless, the normative principle implies a constant reevaluation of the possibility of automating the task performed by humans in Stall Catchers, as in the ML challenge mentioned above. As I show in this work, it also implies that Stall Catchers and HC systems in general must be built as systems that remain open to new human–technology relations with the “moving target” (see Michelucci’s quote above) and to the development of computational solutions to tasks previously performed by humans. HC systems, therefore, can never present complete or finished products but must remain at the edge of AI developments. This is where the imaginaries play a driving role for the intraversions of human–technology relations.

From the imagined human–computer conversation described above that is ethical and enjoyable for humans, and the imaginaries’ consequences for HC development at the Human Computation Institute, I return to the question of who is in the loop in the next subchapter, but this time, focusing on how crowds—rather than individuals—are imagined to be in the loop.

Crowds in the Loop

Although HC systems do not always or necessarily involve a collective or crowd of humans, the imagination of the human in the loop as a “crowd in the loop” is a fundamental idea behind the HC-based CS examples studied. Here, the understanding is that the power to build HI lies in the combination of machines with (digitally) interconnected humans (Human Computation Institute, n.d.). In this way, HC is linked to both the idea of the “wisdom of crowds” (Surowiecki 2005) and the concept and research field “collective intelligence.”¹⁵ The term “collective intelligence” gained popularity in the 2000s, in part

15 In the *Handbook of Collective Intelligence*, in which the editors and organizational theorist and management scholar Thomas W. Malone and computer scientist Michael S. Bernstein aim to form the new interdisciplinary field “collective intelligence,” they define “collective intelligence” following Malone, Laubacher, and Dellarocas broadly as “groups of individuals doing things collectively that seem intelligent” (Malone, Laubacher, and Dellarocas 2009; cited in Malone and Bernstein 2015, 1). They do not specify further what they mean by “intelligence” to keep it adaptable to different understandings of “intelligence” and point to the perspective with “seem” (Malone, Laubacher, and

with the journalist James Surowiecki's publication *The Wisdom of Crowds* (2005; cf. Malone and Bernstein 2015, 6). The wisdom-of-the-crowds approach is a fundamental aspect of all crowd-based HC systems. Building on examples ranging from computer algorithms to stock prices and votes, Surowiecki defines four characteristics for a group or crowd to be "wise:"

diversity of opinion (each person should have some private information, even if it's just an eccentric interpretation of the known facts), independence (people's opinions are not determined by the opinions of those around them), decentralization (people are able to specialize and draw on local knowledge), and aggregation (some mechanism exists for turning private judgments into a collective decision). (2005, 10)

With these characteristics, Surowiecki argues, the likelihood of an accurate crowd answer is high as each individual person's errors will cancel each other out (2005, 10). The characteristic of diversity of opinion implies that individuals do not have to be experts on the problem in question. In an interview with Neil Savage for the *Communications of the ACM* (Association for Computing Machinery) magazine Adrien Treuille, computer scientists and one of the creators of Foldit, reflected on the astounding insight of Foldit's functioning: "[l]arge groups of non-experts could be better than computers at these really complex problems," he argued (Treuille in Savage 2012). "I think no one quite believed that humans would be so much better" (Treuille in Savage 2012).

The calculation of crowd answers in Stall Catchers is also built on the assumption of nonexperts contributing to the game: "we don't assume that every catcher is as accurate as a trained laboratory technician. Instead we use 'wisdom of crowd' methods that can derive one expert-like answer from many people" (Michelucci 2017b). Humans in the loop in HC are, thus, not imagined as individual experts, capable of solving the problem at hand but as crowds of nonexperts capable of producing "expert-like" answers that are better than computational results.

However, the "wisdom-of-the-crowd" approach, while emphasizing the collective effort of combining human efforts, still relies on the isolated individual human in the crowd, as Surowiecki's second characteristic ("independence") states. Furthermore, if we follow Surowiecki's definition, these individuals are then aggregated in a way that their "private judgements [are turned] into a collective decisions" (Surowiecki 2005, 10, see above). The "crowd" describes an imagination of aggregated individuals not interacting, communicating, or even aware of each other. This imagination of the crowd, however,

Dellarocas 2009; cited in Malone and Bernstein 2015, 1). Moreover, with "acting," they constrain that "intelligence to be manifested in some kind of behavior" (Malone and Bernstein 2015, 3) that emerges from individuals or groups that act together. Michelucci defines "collective intelligence" in the *Handbook of Human Computation* in a similar way but with specifying "intelligence" as "problem-solving." "A group's ability to solve problems and the process by which it occurs" (2013d, 84). "Collective intelligence" can, hence, be understood as a subdiscipline of HC or HI, since the latter focus on the combination of crowds, groups of people, with computers and AI, while "collective intelligence" could also refer to groups of people or biological systems without any computational elements.

does not completely align with those of Foldit's and Stall Catchers' teams, who, instead, invest in building "communities."¹⁶

Even though participants in crowd-based HC-based CS at the Human Computation Institute are considered the same as classification providers, the scientific purpose of the projects requires careful evaluation of the individual contributions. The "wisdom-of-the-crowd" approach generally does not account for individuals who do not contribute with the best of intentions. In order to ensure that no harmful input occurs, algorithmic control functions are implemented in Stall Catchers. At the same time, these control functions take into account that user input may vary; it is assumed that humans do not always perform equally well. The control mechanisms will be described in detail in Chapter 6. In a nutshell, while humans are currently doing the annotating of research videos because there are currently no algorithms that can solve this task, it is still an algorithm that evaluates how good the human input was in each case by measuring the participants' "sensitivity." At the same time, individual participants' responses are weighted according to their sensitivity level to the research data. Thus, although the actual task cannot be solved computationally, the ability of humans is, nevertheless, continuously reviewed by algorithms and the skill level dynamically adjusted if necessary. Humans are integrated into these systems in a way that is primarily technical, reflecting the previously discussed understanding of humans in the loop as information processors. Stall Catchers relies on users processing information differently (regarding other humans) and accounts for fluctuations in individual performance in near real time (fieldnote, Nov. 16, 2022).

In this sense and summing up the three imaginations of the human in the loop, they build upon crowds consisting of individual input providers or information processors whose annotations and skill levels are evaluated, rated, and weighted by computer algorithms. This contrasts with social scientific understandings¹⁷ of individuals as "embodied, socially situated, and 'cultured' human beings" (Beck 2012, 136). While computational models cannot yet take over the analysis, the analysis results are only created in human– or crowd–technology relations by the combination of different participants through computer algorithms based on statistical methods.

From these human-in-the-loop visions, I now turn to the HC imaginary of the *Thinking Economy*, to which HC is understood to lead and in which the idea of the human in the loop is understood to fully unfold.

16 For example, for both Foldit and Stall Catchers, this includes creating ways for participants to connect, such as through forums or in-game chats.

17 Similar observations have been made by Forsythe with the example of the understanding of knowledge when studying knowledge engineers who "treat knowledge as a purely cognitive phenomenon. Knowledge, to them, is located solely in the individual mind; expertise, then, is a way of thinking. This contrasts with the social scientific view of knowledge as also being encoded in the cultural, social, and organizational order. Given the latter view, contextual factors are seen to play a role in expertise, and knowledge appears to be a social and cultural phenomenon as well as a cognitive one" (Forsythe [1993] 2001h, 52).

Humans in the Loop in a Future *Thinking Economy*

The potential of HC is considered to be particularly significant in the field of CS, which has played an important role in HC research since its beginning.

In our interview, Bry, who was one of the principal investigators (PI) of the ARTigo project, editorial board member of the *Human Computation Journal*, and contributor to the *Handbook of Human Computation*, saw the potential of HC and CS as “tremendous” (Mar. 2, 2020):

So, I'm completely convinced that in many, many [scientific] areas, such as ethnology for example, human computation, crowdsourcing, citizen science have a *huge* future, but even more than that. I believe that this is the approach of the future in science, [...] in the cooperation of people, also at the workplace in general. The future of human work is either caring for people or creativity *and* traditional professions, which will remain, but will be so optimized by software that compared to today only a fraction of people will do it. [...] There is nothing better than human computation to tap into creativity. Science lives from people without expertise coming up with ideas that experts don't have. [...] So, I believe the future of work will be in the direction of human computation, citizen science. (Bry, Mar. 2, 2020)

According to this vision, HC enables creativity and particularly supports “out of the box thinking.” In the example of science, this is achieved by including perspectives of people without specific training in a particular field. According to Bry, the focus on care and creativity will change work in general, also in the scientific field at universities (Mar. 2, 2020). He further explained that “[y]ou don't need this tight frame, these working hours, a permanent room and so on. People work much more creatively when they are freer. And one will get this freedom. [...] And the complement to that will be more creativity. This will make this freedom affordable *and* justify it. (Bry, Mar. 2, 2020) With HC, Bry argued, most routine tasks will become more efficient and largely automated, freeing people in all sorts of jobs to think creatively, such as the working conditions now enjoyed by a few professors. The potential of HC is imagined to be endless; HC will unleash humans from undesired work so they can focus on their creativity. When I asked him if there were also limits to what HC could do, Bry declared that while there would be some limits, he did not currently know what they would be, “we will see” (Mar. 2, 2020).

The understanding of the future of work enabled and shaped by HC is shared by Michelucci, who uses the term *Thinking Economy* to describe the future he envisions. While Bry focuses on the creativity of humans which cannot be replaced by computers, Michelucci centers “problem-solving” instead:

I think the bottom-line is, I think we're scratching the surface of how people can be involved in science. And I think science, what is science? To me it's problem-solving. So, what we're basically saying is: Science is a methodology that helps us do problem-solving effectively and we can get lots of people involved in that process in many different ways and citizen science is helping us begin that process of involving people. And frankly, I think, that's what people are going to be doing. I mean, 50 years from now, factories will be automated. So, we don't need people working in

factories, I mean, basically we're just going to have, like okay, we've created all these technologies, we've created a bunch of problems [...], we have problems we didn't create but we didn't fix, and we need to solve them. So, we just need all the people to be spending their time on the problem solving. So, this is a way to do that. (Jan. 21, 2020)

Michelucci envisions a future where, science, i.e., problem-solving, will be a collaborative effort involving not only professional scientists but potentially anyone whose work may be automated in that future. In addition to (scientific and societal) problems in general, problem-solving will need to address the problems created by humans while developing technologies.

In contrast to marketing and information management researchers Roland Rust and Ming-Hui Huang, who argue that we currently live in the *Thinking Economy* (Rust and Huang 2021) and are in the midst of the development of a "Feeling Economy" (Huang and Rust 2019), Bry and Michelucci locate the *Thinking Economy* in the future. In this development toward a *Thinking Economy*, the COVID-19 pandemic was "sort of a catalyst for the *Thinking Economy*" (Michelucci, Jan. 14, 2021). It has "accelerated the shift [...] that I envision toward a *Thinking Economy*. And so now that it's sort of greased those wheels, there's even more opportunity" (Michelucci, Jan. 21, 2021). In 2020, with the pandemic causing lockdowns across many countries, more and more people started to connect online and engage in CS. Online CS projects in general but especially those focused on coronavirus research, such as Foldit, Folding@home (The Folding@Home Consortium (FAHC), n.d.), and Rosetta@home (University of Washington, n.d.), which are all HC-based projects, experienced a huge upswing (Vepřek 2020). "[B]ecause of this, it'll be easier to get people who are already engaged online to start [...] to participate in human computation systems. And I think people do it already without knowing it." (Michelucci, Jan. 14, 2020) The potential of HC is understood to be infinite by both Michelucci and Bry. In this imaginary, HC will enable people in the future *Thinking Economy* to spend their time with problem-solving in creative and enjoyable ways. The chosen modes of engagement should then be integrated into the sociotechnical system, in which the computational parts are supposed to be built around these individual engagements.

The HC imaginary of the *Thinking Economy* provides insights into the thinking in which HC-based CS are embedded and how the human-in-the-loop visions are understood to ultimately materialize and unfold. As I have discussed earlier, there is a discrepancy between the human-in-the-loop ideas and current implementations. In the *Thinking Economy*, this gap is understood as being closed, and it is, thus, both a utopian aim for HC development and a contrast foil to current implementations. These efforts flow into and partly define HC-based CS's intraverting human–technology relations.

Weaving Together the Imaginaries

Taken together, the imaginaries of HC presented, including the imaginations of the human in the loop (and the loop itself), though not exhaustive, describe the narratives, understandings, and visions driving HC-based CS development as open systems at the edge

of AI and scientific research and, therefore, at least to some extent, the intraversions of human–technology relations. According to these imaginaries, HC systems are to be developed in such a way that they allow people to participate in a self-determined way and leave room for human creativity. Humans and AI in HC imaginaries are partners in a dialogue in which the AI is an assistant to the humans. But, as I will show in the following chapters, the roles of the AI or computer and humans in HC-based CS games are not fixed. Instead, they change continuously and intravert over time due to properties that may be inherent in HC assemblages and external forces acting upon them.

Moreover, there is a mismatch between the HC imaginaries and current implementations. I argue that it is precisely this mismatch that plays an important role in the intraversions of human–technology relations because the very gap to full HI that this mismatch represents is what drives developers and scientists to push the limits of the system further toward their envisioned goal. Human computation advocates and developers play important roles in forming these systems and their human–technology relations. However, their imaginaries also face resistance and conflict. Participants and researchers with their own aspirations and motivations, and computational models not meeting scientific requirements, for example, interfere with these visions in everyday life.

Human Computation visions of creating “hybrid thinking systems” (Bowser et al. 2017) with humans in the loop remain largely abstract. The question of how humans and/or human intelligence could and should be involved remains at the core of current research efforts in the HC field in general and is often negotiated on a case-by-case basis. The practices of designing and developing HC systems that are shaped and created by the prevalent visions and counter-imaginaries (while, at the same time, forming and creating them) do not generally follow established procedures or frameworks for combining humans and AI.¹⁸ Instead, these “future practices,” to borrow sociologist and cultural theorist Andreas Reckwitz’s (2016) term, consist of continuous experimenting, sketching, prototyping, and infrastructuring. By introducing path dependencies, these practices not only materialize imaginaries but also constrain them in a particular direction. I now turn to these practices, which, along the way, negotiate questions of practical human involvement and the distribution of roles between AI and humans to realize enhanced performance and desirable futures.

18 The *Handbook of Human Computation* can be considered a first attempt to address this gap (Michelelucci 2013a). However, on the level of software architecture or development, for example, developer Kate of the Human Computation Institute explained that she could not rely on existing code to build HC-based CS as she would when working on other software projects, such as e-commerce platforms. For the example of Stall Catchers, Kate argued that “Stall Catchers is something unique in a sense that you don’t really have a lot of people that have worked on similar types of projects so there are not a lot of open-source tools that you can use if you want to do something. So, most of the stuff I had to code from scratch” (Nov. 19, 2020).

Imagining as Practice: Infrastructuring and Experimentation

Imagination, not only but especially when visions have not risen to a hegemonial or sociotechnical imaginary in the understanding of Jasanoff and Kim (2009; 2015), does not happen in intangible space. Instead, it relies on “hands-on” (Hilgartner 2015, 45) practices and exploration. In the field of HC, the question of *how* superhuman capabilities can be achieved by combining humans and machines or AI in new ways, or how imaginaries of AI overcoming humans can be fought against to find the “human place” in the future (Carlsbergfondet 2019, 00:58-1:00), is negotiated alongside the development through infrastructuring and experimenting with sociotechnical systems.¹⁹

I apply the term “infrastructuring” to focus on the material-semiotic practice and stress how infrastructures are constantly in the making (Niewöhner 2015, 5; cf. Bossen and Markussen 2010). Every sociotechnical system, no matter how established and common, requires continued infrastructuring. While infrastructuring practices are, thus, part of every computational project or software platform development, I here want to discuss two different forms of infrastructuring I observed at the Human Computation Institute. Just as “infrastructures operate on differing levels simultaneously” (Larkin 2013, 330), infrastructuring takes place on different levels. Because HC systems are at the edge of AI and scientific research, they often cannot build upon existing frameworks, and developing HC requires very “fundamental” infrastructuring. I first turn to an example of *fundamental* infrastructuring pursued at the Human Computation Institute before discussing instances of *everyday* infrastructuring in the Stall Catchers project. While I focus here on infrastructuring related to the Stall Catchers platform performed by team members of the Human Computation Institute, infrastructuring related to the development of the data pipeline in the laboratory that prepares the data to be sent to the platform is analyzed in Chapter 6.

Infrastructuring Toward “Sustainable Human Computation”

Infrastructuring can be understood as part of the bigger goal of materializing and establishing HC and the vision of hybrid human–AI systems. One such infrastructuring endeavor pursued by Michelucci and colleagues at the Human Computation Institute is to make the development of HC systems more “sustainable” (in the sense of self-sustainable as opposed to environmentally sustainable). This is one of the specific aims of their Civium initiative, “an integration platform and commerce engine for sustainable human computation” (Michelucci 2019b). Michelucci described the idea of this platform to be built in an article on *Medium* from 2019:

Civium is an operating system for a new class of supercomputers powered by computer hardware and cognitive “wetware” that will enable us to build, improve, and deploy transformative human/AI systems in support of open science and innovation.

19 In this sense, HC practitioners share with AI practitioners studied by Forsythe that their goals are very broad and the “meaning and appropriate scope of ‘artificial intelligence’ [or in this case HC,] are subject to ongoing negotiation” (Forsythe [1988] 2001a, 76).

It is also a bazaar, for sharing, trading, and finding the widgets and services we need to create and sustain the capabilities we seek, breathing new life into unsupported projects and platforms. Ultimately, we believe Civium has the potential to seed a new thinking economy that rewards the uniquely human cognitive abilities needed to tackle our most pressing societal issues. (Michelucci 2019b)

Civium's goal of building a foundation for future HC systems and currently unsupported projects originated in part from the "sustainability problem in citizen science" (Michelucci, Jan. 21, 2021). Most HC-based CS projects are developed by nonprofit research or academic institutions and are, therefore, highly dependent on funding. If funding ceases, most projects cannot be sustained (cf. Miller et al. 2023). Addressing this problem is important and valuable to Michelucci, as he explains, since these projects

are serving [...] a valuable purpose in society. So, then it is a question of well, who finds it valuable and why aren't the people who find it valuable paying for it?! If there's a need for this, if there's no need for it, there's no need for it! But it exists because there was a need in the first place, and somebody is benefitting from that need. So, who is benefitting and why aren't they getting sustained?! (Jan. 21, 2021)

From this starting point, Michelucci, together with collaborators, began to identify where the problem was coming from:

I realized that part of the problem is that citizen science, like many research activities, has become isolated in our broader economy in that there are these sort of cultural divides between academic research and industrial research and that there's an opportunity there. [...] So even though there may be distrust ... there may be exploitation, there might be, in other words, reasons for distrust. (Michelucci, Jan. 21, 2021)

Civium was conceived with the aim not only of enabling projects to overcome the sustainability problem but also of enabling individual users—from CS participants to developers, researchers, and other actors in the future *Thinking Economy*—to decide for themselves the conditions under which they wish to participate in human/AI systems (Michelucci, Jan. 21, 2021). "And we use our actual human computation methods to enable these things" (Michelucci, Jan. 21, 2021). The idea also aimed to reflect on and address the power hierarchies that exist in "top-down" online CS projects. Here, designers, developers, and researchers currently set the terms under which participants can choose to contribute to scientific research. This does not mean that participants do not play an active role in shaping the projects and cannot resist inscribed meanings or intended usage (see Chapters 5 and 6). However, even when participants are invited to co-shape the projects and their involvement by providing feedback to the developers, in the case of disagreement or differing priorities about which features to implement or which bugs to fix (cf. Miller et al. 2023), it is the project teams that decide, not the participants. All they can do then is decide whether to leave or to stay. By contrast, within Civium, Michelucci envisions that participants

can decide whether they're comfortable with that role and if they're not, they don't have to adopt that role, or, maybe they feel like they should be getting something back that they're not and they can advocate for themselves more easily to do that [...] in so far as ... human cognition becomes a more critical resource for executing scientific research. I think that gives the individual contributors leverage which can provide pressure against hierarchy to say, hey, you need us! (Jan. 21, 2021)

Additionally, Michelucci explained, Civium is intended to solve the problem of "duplication" in CS, referring to the fact that CS projects are often built from scratch, rather than building on existing mechanisms and approaches, because the infrastructures required are missing (Jan. 21, 2021). Even if "one-size doesn't fit all, [...] there are ways to do community outreach, there are ways to build communities. And there are people who are very good at it. So, if these communities exist already, do you need to spend two years building one from scratch?" (Michelucci, Jan. 21, 2021).²⁰ Michelucci's vision was for Civium to become an "eco-system" that combines all these functionalities and becomes "a marketplace and an integration platform and [...] a policy engine so that people can set terms" (Jan. 14, 2021) to enable sustainable HC. As such, Civium is a "matter that enable[s] the movement of other matter" (Larkin 2013, 329) and is itself a sociotechnical system building on HC. As Larkin formulates for infrastructures, "[t]heir peculiar ontology lies in the facts that they are things and also the relation between things. As things they are present to the senses, yet they are also displaced in the focus on the matter they move around" (2013, 329).

At the time of writing, Civium is still under construction. The Human Computation Institute and external collaborators have started working on implementing the structures of and for the platform, such as an "experimentation toolkit" that allows the creation of sandbox versions of online CS projects to run experiments without affecting the actual "live" project and the scientific research behind it (Vepřek, Seymour, and Michelucci 2020). To this date, however, Civium, which is supposed to provide the infrastructure for HC, is "a neat idea and it looks like we're doing a few interesting things related to that but it's not a thing yet" (Michelucci, Jan. 21, 2021). Here, the "duality" of infrastructures described by Larkin (2013, 329) becomes apparent, since infrastructuring itself is dependent on resources to become technology supporting systems.

This brief excursus shows how infrastructuring not only materializes but also shapes the not yet stabilized visions of HC through the introduction of "path dependencies" (Klausner et al. 2015). By promising to facilitate HC development, these path dependencies also constrain the possibilities of what is possible to imagine. Just as the development of Civium's modules or experiment toolkit is shaped by Stall Catchers as the institute's

20 In addition, building on experiences of applying for ethical review for HC-based CS projects in the US, which to date is not specifically tailored to emerging fields, such as online CS or AI research, Civium should include a new approach to ethical review to ensure that research and CS are conducted ethically and that no participants are harmed. During my fieldwork and collaboration with the institute, I contributed to the discussions and analysis of the current Institutional Review Board or ethical review approach in the US and how it could be improved to better fit new emergent research and application fields (Vepřek, Seymour, and Michelucci 2020; Vepřek 2022b).

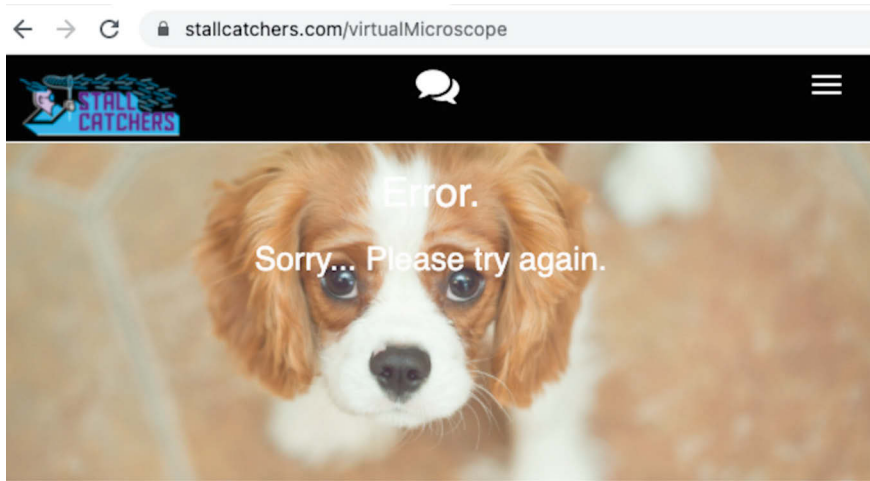
main project, Stall Catchers itself was inspired by and built upon the existing online CS project Stardust@Home (Westphal et al. 2005; Stardust@home, n.d.).

Stardust@Home is an online CS project launched in 2006 after NASA's Stardust mission returned with a collection of interstellar dust embedded in an aerogel collector. It invites participants to search for micron-sized interstellar dust particles in images, respectively, short videos consisting of image stacks, which are based on the data collected. The data is presented in a customized virtual microscope, similar to the Stall Catchers platform. In fact, Stall Catchers' UI was built from the Stardust@Home's UI, including the virtual microscope, and even its code served as a starting point for Stall Catchers. "Who knew stardust and blood vessels could be so similar?," asks Stall Catchers' website (Human Computation Institute, n.d.c). By drawing on the existing CS platform in the field of astronomy, the way in which stalled blood vessels in Alzheimer's disease research are to be identified was pathed. And since Stall Catchers then served as the base model for future HC-based CS projects, such as the Human Computation Institute's Dream Catchers project on sudden infant death syndrome (Ramanauskaite 2020) and other future projects to be built on Civium, specific ways of analyzing research data and engaging participants guide the development of future HC-based CS further. As in the general example of evolutionary histories of machines described by computer scientist Iyad Rahwan and colleagues (2019, 481), parts are reused in different contexts which constrains future performance but also enables new innovations. Or, as assemblages, which are always to be thought of as multiples, extend, they also change (Deleuze and Guattari 2013, 7). To understand how path dependencies emerge and guide developments, it is important to trace both the emerging (counter-)imaginaries and their situatedness as well as the infrastructuring that materializes and shapes these narratives and visions. Civium presents such an infrastructuring project that, if completely realized, will probably path the way and form of future HC-based CS development.

Puppies in Stall Catchers: Everyday Infrastructuring

Shortly before the official start of the Catchathon on April 29, 2021, at 7:00 pm CET, several participants, including the new AI bot GAIA and myself, had problems accessing Stall Catchers. Instead of its UI, a screen with a picture of a puppy looking at the user expectantly and an error message saying "Error. Sorry...Please try again." appeared (see Figure 2). At 6:25 pm CET Michelucci sent a "mayday" to Stall Catchers' lead developer. Something had gone wrong with the activation of the dataset. While two team members were busy setting up the live streaming for the kickoff event and entertaining those participants who had already joined the Zoom meeting for the kickoff, Stall Catchers' lead developer, the bot creator, Michelucci, and myself were trying to figure out where the puppies came from (fieldnote Apr. 28, 2021).

Figure 2: *Puppies in Stall Catchers. Error page*



Source: Screenshot taken by LHV on May 17, 2021 (<https://stallcatchers.com/virtualMicroscope>)

While a moment like this was exceptional in the sense that the Stall Catchers team was constantly working to prevent such situations, such breakdowns are not uncommon in software and technology design in general. Things and processes fail; breakdown is “a condition of technological existence” (Larkin 2008, 234). Even if a process works well most of the time, it can fail for various reasons. In the example of the Stall Catchers puppies, these reasons included data format issues in the dataset to be activated. After creating workarounds and debugging the problem, the puppies went away. But shortly thereafter, a new problem arose that required the team’s full attention and troubleshooting once more. Finally, after two turbulent hours, the Catchathon was underway, AI bot GAIA was catching stalls, and participants were able to access the platform and event pages. On the same day and one of the institute’s main slack channel, Michelucci summarized this incident as follows: “Well that was exciting [face with tears of joy emoji] We’ve learned to expect the unexpected. Within minutes of the event starting the new dataset appeared to crash the server, but then we got that working in time, and then GAIA Bot errored out, and thanks to quick team support we got that working again.” (Apr. 29, 2021)

This example, while not an everyday experience, illustrates how “put[ting] out fires” (fieldnote Apr. 29, 2021) and dealing with the unexpected is part of HC-based CS design, experimentation, and maintenance. Not only did the introduction of a new element (GAIA) cause problems but also routinized processes, such as activating a dataset, failed. Infrastructuring is a crucial part of the everyday maintenance of HC-based CS projects.

With limited time and financial resources, the Human Computation Institute’s team not only had to put out fires at special moments, such as a Catchathon, but also often had to interrupt the day-to-day development of new features, work on Civium, or other new projects. The institute’s developers (at most times there was one, sometimes two developer(s) working for the institute and the developers changed over the course of my research), for example, were not only full stack developers but also responsible for the

entire troubleshooting, maintenance, and operational processes, as developer Kate described in our interview:

I basically oversee everything about the platform, making sure that everything is running smoothly. So, development of course, getting new features and debugging any issues with the platform but also making sure that the infrastructure is running smoothly. So, anything related to the servers, the database, things like that. (Nov. 19, 2020)

It was not unusual for the institute's developers (as well as other members of the institute) to abandon their carefully prepared work plan to intervene: “[S]ometimes I have some intervention I need to deal with [if] Stall Catchers is having some issues. I have to ... wake up in the night and try to debug it—sometimes, not all the time,” explained developer Samuel (Sept. 2, 2021). In these moments, acute measures and actions were required that relegated the imagination of how HC-based CS should be developed to a secondary concern until the problem was resolved. Since these interventions were not infrequent but rather common and recurring, they not only pushed back set timelines but also influenced the imaginations of future systems and human–AI collaborations (an example of the latter will be provided in Chapter 6).

Keeping the system up and running was critical but not the only practice of everyday infrastructuring. I would like to focus here on two brief examples of such practices which also demonstrate how the imaginations of HC-based CS as ethical projects are continuously materialized and performed in the everyday. I see these practices as infrastructural work because of their essential role in creating a participant base and, thus, a foundation for the projects. The examples are 1) creating a seamless and effortless experience for participants, and 2) being transparent and responsive to participants. They show that for the HC system to *work* in the CS domain, from the perspective of the Human Computation Institute team, it is not only necessary to continuously work on the infrastructures as described but also to present the project to participants in a meaningful and enjoyable way.

The first example of the creation of a seamless and effortless experience relates to game design:²¹ HC-based CS projects rely on game design to motivate participants to

21 The example also relates to practices of developing HC, which are also shaped by contingencies and uncertainties that are specific to it due to its human-in-the-loop approach. While one could say that, in a simplified way, software developers generally write code for it to be executed by computers, in HC, humans perform part of the computation. As Gray and Suri have described using the example of ghost work: “Normally, when a programmer wants to compute something, they interact with a CPU [Central Processing Unit] through an API defined by an operating system. But when a programmer uses ghost work to complete a task, they interact with a person working with them through the on-demand labor platform’s API. The programmer issues a task to a human and relies on the person’s creative capacity—and availability—to answer the call. Unlike CPUs, humans have agency: they make their own decisions. While CPUs just execute whatever instruction they are given, humans make spontaneous, creative decisions and bring their own interpretations to the mix. And they have needs, motivations, and biases beyond the moment of engagement with the API. Given the same input, a CPU will always output the same thing. On the other hand, if you send a hungry human into a grocery store, he or she will walk out with

complete what are often monotonous tasks and to present complex processes in a simple way. Therefore, the role of the game interface is central. Human Computation Institute team member Paul explained in our conversation that the experience of the participants had to be “as seamless, effortless as [...] possible” so that participants were not “burdened with [...] the whole interaction” (Oct. 14, 2020) taking place behind the game interface. By “whole interaction,” he was referring to the complex participant–software interplay required to perform data analysis valuable to the scientific research in the laboratory. Creating such a “seamless” experience involved practices of not only ensuring the platform’s availability but also introducing new game-specific features or organizing special events and playful challenges that motivated participants to continue contributing. With these efforts, the team aimed to move the human–software interplay that generates crowd answers into the background of the participant experience. In my conversations with Stall Catchers participants, it was, thus, not surprising that HC itself was only mentioned in very few conversations.

Being transparent and responsive to participants was also important for the Stall Catchers team. It aimed to be as transparent as possible about the decisions made at the institute, the science in the laboratory, and the scientific processes and progress behind the projects (fieldnote Oct. 14, 2022). The institute saw itself as a mediator between scientists and participants, for example, by explaining and translating the scientific research and sharing it through blog posts (fieldnote Oct. 19, 2022). They also sought to be transparent about design decisions, problems, and mistakes they made in relation to the platform or project: “[W]e showed our human side right away and [...] we made ourselves very humble in front of the community and we let the community know that that’s how we were coming to the relationship with them, with this deep sense of humility,” explained Michelucci (Jan. 14, 2021). Besides communicating with participants through blog posts and public events such as “hangouts,” where participants could engage in conversations with the developers and scientists, this principle also includes being “responsive” (Paul, Oct. 14, 2020) by monitoring the different communication channels and answering questions from participants. The same principle was also shared by the Foldit team whose developers, for example, shared a “good faith agreement” about transparency:

I think that as long as we as the developers really have a good faith agreement that we are going to try and always make it fun, and we are always going to try to be as transparent as we can for players about the science that we’re doing. And then it’s their choice entirely as to whether they want to participate. We just hope that they want to continue participating. (Hugo, Jan. 28, 2020)

a dramatically different bag of groceries than if they were not hungry. In exchange for this impetuosity and spontaneity, humans bring something to work that CPUs lack: creativity and innovation.” (2019, xiv) Thus, this defining characteristic of HC itself introduces new challenges in developing the human–software system due to the unpredictability of human engagements (see Chapter 5). These uncertainties and the resulting contingencies not only interfere with grand HI imaginaries and project development plans but require continuous attention, immediate response, and practices beyond software development.

Nevertheless, for CS games, there remains an unresolvable tension between the principle of being transparent to participants and not revealing too much about the game mechanics to ensure the scientific data quality. I will address this “tricky balance” (fieldnote Oct. 14, 2022) that must be maintained in Chapter 5.

Between Counter-Imaginary and Infrastructuring

In this chapter, I explored the visions and imaginaries of future human–technology relations or combinations building “hybrid thinking systems” that underlie HC systems and drive their design and development. I showed how advocates present HC as an alternative to strong and general AI endeavors. They distinguish themselves from strong AI narratives by positioning HC as a counter-imaginary to such pursuits and emphasizing the importance of the human-in-the-loop approach, which they see as crucial to achieving capabilities beyond purely computational ones, while mitigating the dangers voiced in dystopian AGI imaginaries. Despite these definitory efforts, however, HC shares common paradigms with such AI efforts. At the level of visions and imaginaries, then, HC is also rooted and situated in the very reference points from which it seeks to distinguish itself.

Furthermore, I discussed different human-in-the-loop imaginations that describe *who* is envisioned to be in the loop and how the *human* is imagined, how the *loop* is envisioned as a game and as human–AI conversations of the future, and the imaginations of *crowds-in-the-loop*. Eventually, I explored the imagination of HC as leading to a future *Thinking Economy* unleashing human creativity. While these imaginations form and inform the everyday design and development of HC systems, as “social practices” (Jasanoff 2015b, 323) they simultaneously rely on infrastructuring and experimenting. How these imaginaries or vanguard visions are materialized and explored through examples of everyday practices of designers and developers was the focus of the second part of this chapter. Mackenzie described for code that

[c]ode understood as a collective imagining seems a long way from code as a program of instructions for a machine to execute. However, practices of imagining are not purely mental operations; in no way does imagining reduce to a detached, abstract fantasy. It constitutes collective relational realities. Software attaches different localities to each other because it diffuses relations between them. The composite texture of software is reliant on unfinished exchanges between code and coders. (2006, 138)

Similarly, HC systems emerge in practices of designing, developing, and maintaining concrete sociotechnical systems, and such practices not only fill the idea of combining humans and AI to achieve superhuman capabilities with hands-on examples and meaning but also renegotiate and transform these imaginaries.

Yet, imagining and infrastructuring HC does not only take place within AI discourses, but design and development of HC systems are just as much shaped by their own everyday becoming in the interplay of all human and nonhuman actors involved. Often, they change and adapt through serendipitous discoveries or other coincidences,

through object potentials and timely moments seized by actors or through the human–technology relations unfolding in (and forming) the sociotechnical assemblages. While this chapter focused on how HC is imagined and HC-based CS systems are designed to “stabilize practice,” the following chapters will also focus on how these systems or assemblages are “destabilized” through (creative) practice (Beck 1997, 296).

