

# Hybrid Human-AI Socio-Technical Environments in the Digital Age

Gerhard Fischer

University of Colorado, Boulder, CO USA

|          |                                                                                                  |          |
|----------|--------------------------------------------------------------------------------------------------|----------|
| <b>1</b> | <b>INTRODUCTION .....</b>                                                                        | <b>3</b> |
| <b>2</b> | <b>THE ARRIVAL OF AI-ASSISTED DEVELOPMENT: NEW POSSIBILITIES AND NEW CHALLENGES .....</b>        | <b>3</b> |
| <b>3</b> | <b>FROM PROBLEM SOLVING TO PROBLEM FRAMING.....</b>                                              | <b>3</b> |
| 3.1      | EXAMPLE FOR LIMITED PROBLEM FRAMING: FROM SEPTEMBER 11, 2001 TO THE GERMANWINGS FLIGHT 2015..... | 4        |
| 3.2      | THE GERMANWINGS DISASTER.....                                                                    | 5        |
| 3.3      | THE LIMITED AND INADEQUATE FRAMING OF “COCKPIT SAFETY” .....                                     | 5        |
| 3.4      | EXPLORING THE IMPACT OF AI ASSISTANCE TO IMPROVE PROBLEM FRAMING.....                            | 6        |
| <b>4</b> | <b>SUBSTRATE ACCOUNTABILITY: META-DESIGN AND EUD AT THE INTERACTION LEVEL ....</b>               | <b>6</b> |
| <b>5</b> | <b>DISCUSSION TOPICS (DTS) FOR THE WORKSHOP .....</b>                                            | <b>7</b> |
| <b>6</b> | <b>CONCLUSION .....</b>                                                                          | <b>8</b> |
| <b>7</b> | <b>REFERENCES .....</b>                                                                          | <b>9</b> |

## List of Figures:

|                                                                                    |   |
|------------------------------------------------------------------------------------|---|
| <i>Figure 1: The upstream/downstream distinction in software engineering</i> ..... | 4 |
| <i>Figure 2: Overview of the Case Study</i> .....                                  | 5 |

## Abstract

The rapid integration of large language models (LLMs) into our everyday lives intensifies a tension that the Meta-Design and End-User Development (EUD) communities have identified but never fully resolved: the gap between the systems people use and the systems people can meaningfully shape. My contribution argues that navigating this tension in the AI era requires two conceptual moves that existing frameworks do not yet make explicit.

The first is a shift in the fundamental orientation of socio-technical design — from *problem solving* to *problem framing*. Where problem solving optimizes within a given frame, problem framing treats the frame itself as the primary design object. I illustrate this distinction through a counterfactual analysis of the 2015 Germanwings disaster

resulting from the post-9/11 cockpit door hardening policy, which resulted in unforeseen consequences.

The second contribution is the introduction of a new design principle named “*Substrate Accountability (DP-SA)*”. LLMs are centrally trained, opaque, and largely non-modifiable artifacts, thereby contradicting the principles of design-after-design, evolvability, and informed participation that meta-design requires. Our prior design substrates examined in the meta-design tradition (e.g.: domain-oriented design environments [Fischer, 1994], and the Envisionment and Discovery Collaboratory (EDC) [Arias et al., 2016]) were open to inspection, local modification, and community-driven evolution. LLMs break this assumption at the architectural level, concentrating design authority in a small number of major companies.

Together, these two contributions reframe what it means to design *hybrid human-AI socio-technical environments* in the digital age: not merely integrating AI into existing participatory frameworks, but rethinking the conditions under which human agency, problem ownership, and cultures of participation can be sustained when the most powerful design substrate available is, by design, beyond the reach of its users.

## Keywords

Problem Framing; September 11, 2001; Germanwings Flight 2015; Substrate Accountability; Large Language Models; Hybrid Human-AI Socio-Technical Environments

**Published in:** Proceedings of the 10th International Workshop on Cultures of Participation in the Digital Age (CoPDA 2026): Exploring the Relationship between EUD, AI-Assisted Development, and Meta-Design, June 2026, Venice, Italy. Copyright © 2026 for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

## 1 Introduction

The concepts of meta-design and EUD have been discussed, explored, and instantiated with the development of prototypical systems in numerous publications (e.g.: [Fischer & Giaccardi, 2006; Lieberman et al., 2006], [Fischer et al., 2017; Schön, 1983; von Hippel, 2005]) and previous CoPDA workshops (e.g.: CoPDA 2013 “Empowering End-User to Improve their Quality of Life”). More recent research (e.g.: [Fischer, 2021], [Barricelli et al., 2022]) has integrated and explored the additional possibilities that AI can contribute to support and enhance meta-design and EUD.

My brief arguments in this workshop contribution address the 10th CoPDA's focus on how EUD, Meta-Design, and Cultures of Participation should evolve in the AI era.

## 2 The arrival of AI-assisted development: New Possibilities and New Challenges

The integration of artificial intelligence into our framework was anticipated long before the current wave of generative models. We explored the consequences of two fundamentally different orientations [Fischer, 2024; Fischer & Nakakoji, 1992]:

- *AI replacing* humans (optimizing for automation, efficiency, and the substitution of human judgment) versus
- *AI empowering* humans (augmenting human agency, deepening distributed cognition, and supporting cultures of participation).

Advances in Artificial Intelligence provide new opportunities and challenges for meta-design to empower users more than in the past, giving them new perspectives, strategies, and tools to be exploited in the design-for-design process [Barricelli et al., 2024].

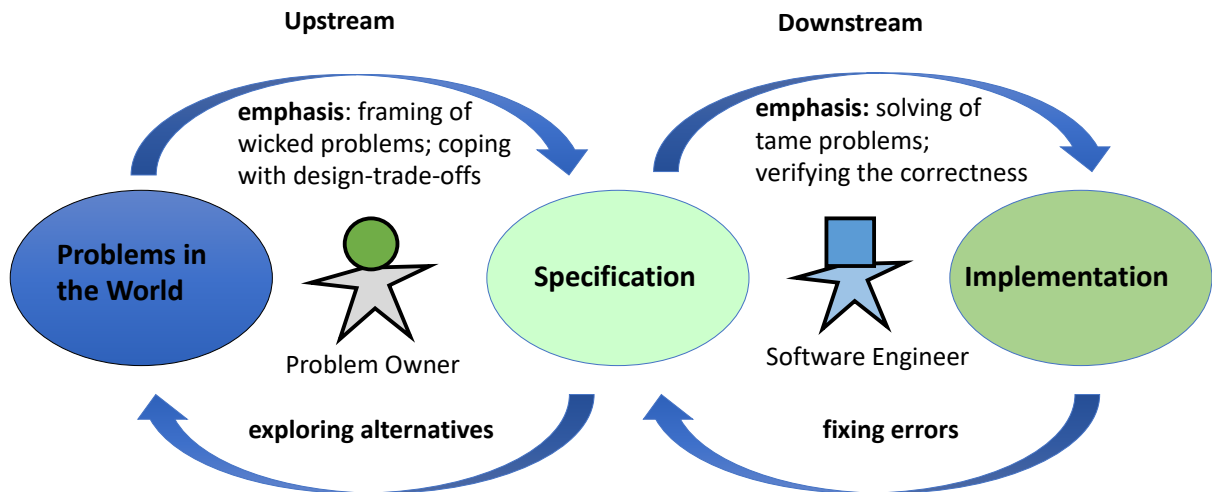
The central challenge for Hybrid Human-AI Socio-Technical Environments in the Digital Age is not choosing between these orientations abstractly, but designing socio-technical environments that sustain human agency with increasing AI capability.

## 3 From Problem Solving to Problem Framing

*"In real-world practice, problems do not present themselves to practitioners as givens — they must be constructed"*  
Schön, D. A. (1983). *The Reflective Practitioner*

This paper advances two conceptual moves. The first is a shift in the fundamental orientation of socio-technical design — from *problem solving* to *problem framing*. *Problem solving* [Simon, 1996] operates within a given problem space and optimizes toward a specified goal. *Problem framing* is the prior and more consequential activity: it

names what will be attended to, constructs the space within which solutions will be sought, and determines which consequences and whose interests will be visible. The upstream/downstream distinction in software engineering [Fischer, 1994] makes the practical stakes of this difference concrete: well-executed downstream work (formally correct implementations) is worthless, and can be harmful, when the upstream frame is wrong.



**Figure 1:** The upstream/downstream distinction in software engineering

This paper adopts a different starting point. It argues that the central challenge of AI is not what these systems can do for users, but how they transform the socio-technical conditions under which users can participate, learn, and shape systems over time. While AI systems reduce many forms of implementation complexity, they simultaneously introduce a more consequential risk: they stabilize provisional problem framings into infrastructure, compress evolutionary processes into continuous optimization, and obscure the very breakdowns that historically triggered reflection and redesign. The result is a shift from *implementation failures* to *design failures* [Curtis et al., 1988]—systems that solve the wrong problems correctly while making this misalignment increasingly difficult to detect.

### 3.1 Example for Limited Problem Framing: From September 11, 2001 to the Germanwings Flight 2015

The September 11, 2001 attacks ([https://en.wikipedia.org/wiki/September\\_11\\_attacks](https://en.wikipedia.org/wiki/September_11_attacks)) were a coordinated series of terrorist suicide attacks perpetrated by al-Qaeda against the United States. Nineteen terrorists hijacked four airliners, then flew one into each of the Twin Towers at the World Trade Center in New York City. The third plane crashed into the Pentagon. The fourth plane crashed in a rural Pennsylvania field during a passenger revolt. The 9/11 hijackers had exploited a simple weakness: cockpit doors were flimsy, unlocked, and easily breached.

To avoid that such a terrorist attack could be repeated, designers, regulators, and airlines worked on new design guidelines on cockpit security. They were operating within a

*problem frame* defined by the available failure data: hijackings by external actors. The major design requirements (mandated to airlines) resulting from their deliberations were:

- Reinforce cockpit doors to be bulletproof and blast-resistant;
- Install locking mechanisms controllable only from inside the cockpit;
- Mandate that doors remain locked during flight.

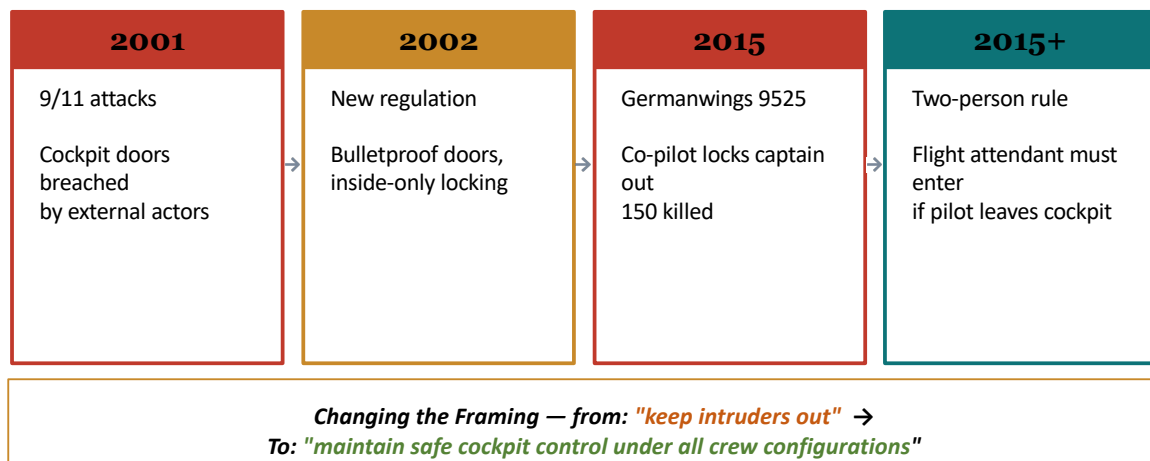
### 3.2 The Germanwings disaster

On a flight on March 24, 2015, the co-pilot deliberately locked the captain out of the cockpit during a toilet break and flew Airbus A320 flight 9525 into a mountain in the French Alps, killing all 150 people on board ([https://en.wikipedia.org/wiki/Germanwings\\_Flight\\_9525](https://en.wikipedia.org/wiki/Germanwings_Flight_9525)). It was later found out that the co-pilot suffered from mental problems and had previously been treated for suicidal tendencies and declared unfit to work by his doctor.

### 3.3 The Limited and Inadequate Framing of “Cockpit Safety”

The design requirement "create safe, locked cockpit doors" had embedded a hidden and catastrophic assumption: that the threat to passenger safety would always come *from outside* the cockpit. It had not considered the cockpit crew themselves as a potential threat vector (see Figure 2). This is a textbook *wicked problem* consequence [Rittel & Webber, 1984]— the solution was framed around one definition of the problem, and this framing excluded the possibility that the solution itself could become the instrument of harm.

## The Post “9/11/2001” Framing Failure Downstream Success and Upstream Failure



**Figure 2:** Overview of the Case Study

The section argues this is not a failure of engineering, but of framing: the problem was defined as "keep intruders out," not "maintain safe cockpit control under all crew

configurations." The Germanwings case illustrates convincingly that a technically sound solution to a wrongly framed problem is not an engineering success — it is a design disaster.

The response to the Germanwings incident was to mandate a *two-person cockpit rule* — requiring that a flight attendant enter the cockpit whenever one pilot leaves, so the cockpit is never occupied by a single person. The future will show to what extent this framing is sufficient to avoid further terrorist attacks.

### **3.4 Exploring the impact of AI assistance to improve problem framing**

In exploring the impact of AI assistance to improve problem framing, a counterfactual analysis could be employed by comparing observed outcomes with a hypothetical scenario. Research should explore the question:

*could AI-assisted design with LLMs available in 2026 (not available in 2002) have helped to move beyond local learning, which refines a solution within a fixed problem frame (exterior danger to the cockpit) to global learning, which questions the frame and may have helped to avoid the Germanwings disaster by considering crew psychology, two-person cockpit protocols, and mental health screening?*

Some LLM possibilities (e.g.: well-crafted prompts such as "prevent unauthorized entry to the cockpit") could potentially have generated ideas to avoid the narrow framing that led to the locked door design guidelines. In 2002, this kind of structured assumption audit was done by expert committees, whose composition (security specialists, aeronautical engineers, regulators) systematically underrepresented aviation psychology and crew welfare. An LLM synthesizing across aviation incident databases, human factors engineering, psychiatric literature on occupational mental health, and pilot suicide case histories from other countries could have potentially avoided such disciplinary blind spots.

While LLMs could, in principle, broaden framing by integrating heterogeneous knowledge sources, empirical evidence [Shin et al., 2025] suggests that they do not reliably improve problem framing in practice. This apparent contradiction highlights a core meta-design challenge: LLMs possess latent framing capabilities, but lack mechanisms to surface assumptions, interrogate problem boundaries, and support reflective reframing. Without such scaffolding, they reinforce existing frames rather than transforming them.

## **4 Substrate Accountability: Meta-Design and EUD at the interaction level**

Traditional EUD research was structured around two interdependent layers. At the *artefact layer*, users create, adapt, and evolve computational artefacts to fit their situated goals and practices. At the *environment layer*, meta-designed socio-technical platforms provide the open architectures, collaborative scaffolding, and participatory infrastructures

within which that artefact-level activity becomes possible and sustainable [Fischer et al., 2017]. The emergence of LLMs introduces a third, qualitatively distinct layer: the *substrate layer* consisting of *foundation models*.

LLMs represent a qualitatively new kind of design substrate: centrally trained, largely opaque, and non-modifiable by the communities that use them [Bender et al., 2021]. To support the objectives of meta-design and EUD, AI substrates must be subject to community governance rather than purely vendor control. Addressing this challenge requires extending meta-design principles beyond the artifact and environment levels to include the AI substrate itself.

This is the basis for the following design requirement:

*DP-SA (Substrate Accountability)*: “Socio-technical environments that embed LLMs must provide model transparency, version awareness, substrate plurality, and community oversight, so that AI infrastructures remain accountable to the communities they serve and compatible with EUD-aligned design.”

This requirement ensures that AI remains a *supertool* supporting human intelligence [Engelbart, 1995; Shneiderman, 2022] rather than an opaque authority shaping outcomes.

It is an interesting question to ask: will future foundation models (owned and governed by OpenAI, Meta, Google, Anthropic,..) change from their current form incorporate EUD control mechanism at the substrate level?

Claims to be discussed at the workshops could be:

- “Future foundation models are unlikely to evolve toward fully user-modifiable substrates.” The economic costs, technical complexity, and safety requirements of large-scale AI systems make centralized governance by model developers likely to persist.
- Emerging mechanisms (e.g.: Retrieval-Augmented Generation (RAG) [AWS, accessed March 2026] such as fine-tuning, adapter layers, configurable safety policies, and retrieval integration are introducing controlled forms of substrate customization. These mechanisms do not provide full end-user development at the substrate level but create configurable extension points that allow limited adaptation while preserving centralized control.

## 5 Discussion Topics (DTs) for the Workshop

This section enumerates a set of discussion topics that could be further explored during the workshop:

- **DT-1:** Is *problem framing* teachable in the abstract or does it require domain immersion and background knowledge? Can the workshop participants identify from

their own research problems of successful reframing activities with the help of human and/or computational environments?

- **DT-2:** Is the Germanwings case a good example of a framing failure that EUD-aligned AI could have helped prevent? Is it easy with counterfactual analysis (what *would have happened* if something had been different) to construct compelling narratives about how things could have gone differently, relying on hindsight but providing no foresight?
- **DT-3:** Is DP-SA an addition to our meta-design framework (or is it an extension to existing DPs such as evolvability, informed participation, and human problem-domain interaction)?
- **DT-4:** Is the opacity of current LLMs a new kind of substrate problem for EUD that requires genuinely new design responses? Is LLM opacity an unchangeable architectural fact or a design challenge that might be overcome?
- **DT-5:** Is Meta-Design and EUD desirable for all problems or settings — or is it a (unnecessary) burden for personally irrelevant problems in which users are better served with *acceptable defaults*? This issue is broadly discussed with the concept of “*Libertarian Paternalism*” in [Thaler & Sunstein, 2009]: not everyone who uses a tool passively is making a mistake — they may be making a rational allocation of attention.
- **DT-6:** Are the communities most harmed by wrongly framed systems (e.g.: marginalized groups whose needs are systematically excluded from design framing) precisely the communities least able to participate in the upstream framing processes? The paper by [Shin et al., 2025] provides evidence that LLMs increase the competence gap between experienced and inexperienced designers: the former were able to extract more novel frames from LLM interactions, while the latter did not benefit.
- **DT-7:** Is *Agentic AI* (minimize human involvement, maximize automation such as VIBE coding, applicable to downstream activities) and *EUD-oriented AI* (humans as primary agents in environments such as the Envisionment and Discovery Collaboratory (EDC), applicable to upstream activities) in fundamental opposition or can they be combined in meaningful ways?

## 6 Conclusion

My contribution is focused on two related arguments:

- *Problem Framing as Primary Design Object:* LLMs are, at present, predominantly downstream tools. They are capable at operating within given frames and can successfully be used to support new levels of distributed cognition.
- *Substrate Accountability (DP-SA):* The design principle DP-SA (specifically in the context of wicked problems) is necessary for creating future LLMs that provide new levels of support for meta-design and EUD.

The arrival of LLMs in socio-technical environments does not simply add a new tool for broadening the possibilities for distributed cognition. It introduces a new kind of *substrate* (one that is centrally trained, largely opaque, and non-modifiable by the

communities that depend on it) in opposition to our meta-design tradition, whose deepest commitments have always been to evolvability, co-adaptivity, and informed participation. The tension between these two approaches is the defining design challenge for *hybrid human-AI socio-technical environments* in the digital AI era with a focus on the specific objectives:

- shifting focus from downstream activities to upstream framing processes;
- extending participatory design principles to the AI substrate level;
- designing interaction mechanisms that make assumptions visible and contestable;
- addressing emerging inequalities in framing competence amplified by LLM use.

## 7 References

- Arias, E. G., Eden, H., & Fischer, G. (2016) *The Envisionment and Discovery Collaboratory (EDC): Explorations in Human-Centered Informatics*, Morgan & Claypool Publishers, San Rafael, California.
- AWS (accessed March 2026) *What Is Rag (Retrieval-Augmented Generation)?*, <https://aws.amazon.com/what-is/retrieval-augmented-generation/>.
- Barricelli, B. R., Fischer, G., Fogli, D., Mørch, A., Piccinno, A., & Valtolina, S. (2022) *Copda 2022: "AI for Humans or Humans for AI?" Proceedings of the Sixth International Workshop on Cultures of Participation in the Digital Age*, <https://ceur-ws.org/Vol-3136/>.
- Barricelli, B. R., Fischer, G., Fogli, D., Mørch, A., Piccinno, A., & Valtolina, S. (2024) "Advancing the Integration of Artificial Intelligence in Meta-Design," Proceedings of the Conference on Advanced Visual Interfaces (AVI 2024), Genoa, Italy (June); published in ACM Digital Library.
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021) "On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?" in *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, pp. 610-623.
- Curtis, B., Krasner, H., & Iscoe, N. (1988) "A Field Study of the Software Design Process for Large Systems," *Communications of the ACM*, 31(11) pp. 1268-1287.
- Engelbart, D. C. (1995) "Toward Augmenting the Human Intellect and Boosting Our Collective IQ," *Communications of the ACM*, 38(8) pp. 30-33.
- Fischer, G. (1994) "Domain-Oriented Design Environments," *Automated Software Engineering*, 1(2) pp. 177-203.
- Fischer, G. (2021) "End-User Development: Empowering Stakeholders with Artificial Intelligence, Meta-Design, and Cultures of Participation" in D. Fogli, Tetteroo, D., Barricelli, B.R., Borsci, S., Markopoulos, P., Papadopoulos, G.A. (Ed.), *Is-EUD 2021 Proceedings*, Springer, LNCS 12724, pp. 3–16.

- Fischer, G. (2024) "A Research Framework Focused on AI and Humans Instead of AI Versus Humans," *Interaction Design and Architecture(s) - IxD&A Journal* Winter 2023-2024 pp. 17-36, DOI: <https://doi.org/10.55612/s-5002-059>.
- Fischer, G., Fogli, D., & Piccinno, A. (2017) "Revisiting and Broadening the Meta-Design Framework for End-User Development" in F. Paterno, & V. Wulf (Eds.), *New Perspectives in End User Development*, Kluwer Publishers, Dordrecht, The Netherlands, pp. 61-97.
- Fischer, G. & Giaccardi, E. (2006) "Meta-Design: A Framework for the Future of End User Development" in H. Lieberman, F. Paternò, & V. Wulf (Eds.), *End User Development*, Kluwer Academic Publishers, Dordrecht, The Netherlands, pp. 427-457.
- Fischer, G. & Nakakoji, K. (1992) "Beyond the Macho Approach of Artificial Intelligence: Empower Human Designers - Do Not Replace Them," *Knowledge-Based Systems Journal, Special Issue on AI in Design*, 5(1) pp. 15-30.
- Lieberman, H., Paterno, F., & Wulf, V. (Eds.) (2006) *End User Development* Kluwer Publishers, Dordrecht, The Netherlands.
- Rittel, H. & Webber, M. M. (1984) "Planning Problems Are Wicked Problems" in N. Cross (Ed.), *Developments in Design Methodology*, John Wiley & Sons, New York, pp. 135-144.
- Schön, D. A. (1983) *The Reflective Practitioner: How Professionals Think in Action*, Basic Books, New York.
- Shin, J., Polyanskaya, A., Lucero, A., & Oulasvirta, A. (2025) "No Evidence for LLMs Being Useful in Problem Reframing" in *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, pp. 1–25.
- Shneiderman, B. (2022) *Human-Centered AI*, Oxford University Press.
- Simon, H. A. (1996) *The Sciences of the Artificial*, third ed., The MIT Press, Cambridge, MA.
- Thaler, R. H. & Sunstein, C. R. (2009) *Nudge — Improving Decisions About Health, Wealth, and Happiness*, Penguin Books, London.
- von Hippel, E. (2005) *Democratizing Innovation*, MIT Press, Cambridge, MA.