

Centering Knowledge Along the Responsible LLM Supply Chain: An Empirical Study & Multi-Stakeholder Taxonomy

AGATHE BALAYN*, Microsoft Research, USA

FANNY RANCOURT, ServiceNow, Canada

FABIO CASATI, ServiceNow, Switzerland

UJWAL GADIRAJU, Delft University of Technology, The Netherlands

Existing CSCW and organization management literature suggests that **knowledge** is a central construct when developing, deploying, and using technological systems that are embedded in multi-stakeholder supply chains. Yet, it has rarely been the focus of comprehensive empirical investigations across LLM and broader AI supply chains. In this work, we draw on semi-structured interviews with 71 LLM practitioners to examine the knowledge required to produce responsible LLM systems, and how practitioners acquire it within LLM supply chains. Our findings not only reveal knowledge blindspots, but also a knowledge access and translation gap, where practitioners recognize knowledge needs but cannot fulfill them. We explain these gaps by showing how knowledge exchanges occur (proactively or serendipitously), which organizational arrangements and tools facilitate them (e.g., knowledge intermediaries), and which barriers undermine them (including limited visibility, lack of mutual understanding, and unclear responsibility for sharing and maintaining knowledge). We further synthesize knowledge needs into a multi-dimensional taxonomy that characterizes knowledge based on its abstraction, theme, and lens. We then discuss how this taxonomy can serve both as a practical resource for organizations of the responsible LLM supply chain to improve actionability and accountability (for example, by supporting precise communication, prompting self-reflection, and facilitating mutual blindspot identification), and as a methodological tool to reframe, revise, and disambiguate research and policy works on responsible AI notions related to knowledge. We hope to inspire future work in the CSCW community by outlining research opportunities to support practitioners in accessing, exploiting, and sharing relevant LLM knowledge across the supply chain.

CCS Concepts: • **Human-centered computing** → **Empirical studies in HCI**; *Empirical studies in collaborative and social computing*; *HCI theory, concepts and models*; • **Computing methodologies** → **Artificial intelligence**; • **Social and professional topics** → *User characteristics*; **Management of computing and information systems**; *Socio-technical systems*.

Additional Key Words and Phrases: knowledge, information, transparency, literacy, explainability, AI supply chain, AI systems, empirical study, qualitative study, semi-structured interviews

ACM Reference Format:

Agathe Balayn, Fanny Rancourt, Fabio Casati, and Ujwal Gadiraju. 2026. Centering Knowledge Along the Responsible LLM Supply Chain: An Empirical Study & Multi-Stakeholder Taxonomy. *Proc. ACM Hum.-Comput. Interact.* 10, 6, Article CSCW061 (October 2026), 31 pages. <https://doi.org/10.1145/3816909>

*Research partially conducted while the first author was at ServiceNow and Trento University.

Authors' Contact Information: Agathe Balayn, Microsoft Research, USA, balaynagathe@microsoft.com; Fanny Rancourt, fanny.rancourt@servicenow.com, ServiceNow, Canada; Fabio Casati, ServiceNow, Switzerland, fabio.casati@servicenow.com; Ujwal Gadiraju, Delft University of Technology, The Netherlands, u.k.gadiraju@tudelft.nl.



This work is licensed under a Creative Commons Attribution 4.0 International License.

© 2026 Copyright held by the owner/author(s).

ACM 2573-0142/2026/10-ARTCSCW061

<https://doi.org/10.1145/3816909>

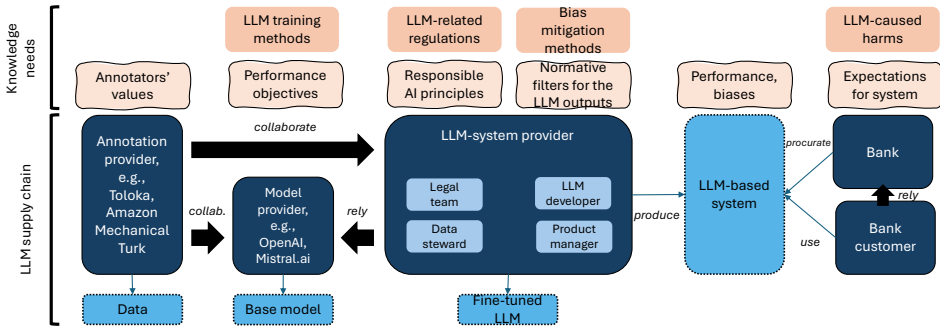


Fig. 1. Example of knowledge needs along an LLM supply chain that our interview participants were involved in. A bank considered deploying an LLM system to enable customer self-service, and hiring another organization to develop this system. The system's outputs may vary in accuracy and offensiveness across customer groups, potentially harming the customers and the bank's reputation. Addressing this requires multiple forms of LLM-related knowledge. The bank must understand potential LLM harms and request mitigation from the LLM provider organization. This organization, in turn, must grasp the bank's and customers' values to assess and reduce offensiveness. The developers in this organization need to accordingly select an appropriate base model, apply bias mitigation techniques, and coordinate with data annotation companies to align training data with the bank's values. Meanwhile, legal teams must ensure regulatory compliance of the system in light of applicable AI regulations.

1 Introduction

Knowledge is necessary to develop and use any technology [22, 61, 63], and is also one of the main assets of any organization [60, 119, 149]. Software systems powered by Large Language Models (LLMs) rely on a particularly complex technology. Technically, LLM-powered systems build on a foundation model containing a multitude of probabilistic components [72]. Socially, they can cause diverse harms if not developed responsibly [17, 130, 150]. Thus, knowledge in its various dimensions (as we unpack through this paper), is a necessary ingredient in creating actionable and accountable responsible AI practices. That is why we argue that studying the needs for knowledge of practitioners who produce LLM systems is the first step to foster responsible LLM systems. Normative notions of transparency, explainability, and literacy in AI testify to this: they are well-spread in both research and policy documents [7, 11, 19, 57, 71, 79, 96, 105, 111], they are being adapted to LLM systems [82, 91, 92, 99], and their common denominator is knowledge for responsible AI.¹

Organizationally, the development of LLM systems spans a complex and rapidly-evolving supply chain of practitioners and organizations [33, 135, 155]. In such a multi-stakeholder setting, the example in Figure 1 illustrates that different actors in the LLM supply chain need diverse knowledge. Furthermore, it shows that “having knowledge” is not simply an individual consideration: actors can acquire this knowledge via diverse types of knowledge exchanges (e.g., during collective decision-making activities, or when relying on others' artifacts). Responsible outcomes, thereby, depend on whether relevant knowledge can be recognized, accessed, shared, interpreted, and acted upon across myriad practitioners, artifacts, and organizational boundaries. Finally, the example also highlights how such knowledge can serve as a necessary vehicle for shaping actionable and accountable organizational structures: given the availability of knowledge and the appropriate facilitation of

¹We focus on knowledge requirements within LLM supply chains. Given that LLMs are a subset of broader AI systems, we can build on existing research in AI literacy, explainability, and transparency, assuming that the principles and challenges identified in these areas are also relevant to LLM-specific contexts.

exchanges, clear accountability boundaries can be charted out for stakeholders throughout the supply chain.

Despite this pivotal role of knowledge and the potential it can offer as a starting point that helps tackle barriers to responsible AI, relatively little is understood about the practices and needs of these actors, e.g., LLM developers, engineers, managers, policymakers, or designers, remain under-studied [37].² Despite some research on practitioners involved in designing AI algorithms and applications [28, 68, 88, 98, 127, 140], most of the focus has been cast on the users of the systems at the end of the supply chain or on stakeholders impacted outside the chain [14, 116]. Moreover, studies cover only a few knowledge types (e.g., excluding know-how) in comparison to information science and organization management literature [125, 147], and rely on the self-reported needs of individual practitioners [68, 89, 157, 159] without exploring their potential blindspots that the knowledge of others in the chain could make visible. Furthermore, these past studies do not comprehensively report on the benefits of, conditions for, and barriers to knowledge exchanges across the supply chain actors. Without such a comprehensive understanding of knowledge needs, and without taking stock of current knowledge strategies and challenges, we cannot orient future research efforts toward facilitating this necessary knowledge transmission and exploitation across the supply chain. Acknowledging this gap, we tackle the following research questions:

- **RQ1:** *How do knowledge exchanges occur among practitioners in the LLM supply chain, for what purposes, and under what conditions?*
- **RQ2:** *What kinds of knowledge might these practitioners have or need for setting up responsible LLM systems?*

To answer these questions, we identify knowledge exchange practices and recurring knowledge needs through 71 semi-structured interviews and in-situ observations with various practitioners working for organizations that develop, deploy, and use LLM systems. We then synthesize these into a comprehensive and unified taxonomy of LLM knowledge. Our study shows that the LLM supply chain does not only fragment knowledge across practitioners, but also constitutes an invaluable resource for learning about making LLM systems responsible. Yet, we need to tackle organizational problems (e.g., visibility, responsibility allocation) and leverage initial strategies for knowledge exchanges (e.g., shared spaces, asynchronous tools) to help practitioners fully take advantage of this opportunity. To do so, our knowledge taxonomy should be extended in the future to support different types of knowledge exchanges (e.g., serendipitous or proactive), while we enrich our understanding of knowledge needs and develop organizational incentives for knowledge exchanges.

Ultimately, this paper makes three contributions to CSCW research on knowledge, supply chains, and responsible AI:

- (1) An empirical account of knowledge gaps in supply chains and knowledge exchange practices including barriers and facilitators.
- (2) A multi-dimensional knowledge taxonomy that practitioners, researchers, and policymakers can rely on in their respective work.
- (3) A synthesis of implications which inform the future design of tools, processes, and organizational structures that facilitate knowledge exchanges.

²We also lack studies around the knowledge practices in supply chains focused on predictive AI systems (preceding LLMs systems). Yet, we scope our work around LLM supply chains due to their urgency and for consistency purposes.

2 Background & Related Work

2.1 Knowledge as Essential for the Development of AI and Other Technologies

While knowledge is recognized as essential by social science and human-computer interaction (HCI) researchers [5], its definition varies [84]. In this work, we adopt a broad definition that reflects diverse forms of knowledge identified in AI and HCI research [57, 93, 105, 107, 133, 141] – “*Knowledge is the combination of data and information, to which is added expert opinion, skills, and experience, to result in a valuable asset which can be used to aid decision making*” [125].³ Education researchers [22, 61, 63] have shed light on the types of knowledge about technologies, their functioning and usage, that is necessary to technology developers.

Given this complexity, such research needs to be extended to AI technologies and, in particular, generative LLM systems [150]. In the past, AI research has implicitly made a case for adopting the lens of *knowledge*, hinting at AI practitioners’ lack of knowledge and low AI literacy as a cause of flawed practices and harmful AI systems [13, 13, 46, 140]. Studies have uncovered the knowledge of AI that some AI stakeholders have, primarily end-users, decision-subjects, or developers (e.g., [52, 68, 89, 98]). Subsequently, researchers have developed explainability tools [10, 53, 75, 94], transparency tools [18, 24, 52, 57, 62, 89, 90, 94], and literacy frameworks [15, 109, 158] to specify the knowledge these actors require for their work, and support them in acquiring it.

⇒ We extend these efforts by empirically studying the knowledge practices of the AI stakeholders that are typically not covered yet play an important role in the development of AI systems (e.g., product managers, and quality assessors). We also complement known knowledge needs that were previously identified using self-reports with additional relevant knowledge nuggets that might otherwise remain in one’s blindspots.

2.2 Knowledge as a Vehicle that Shapes Supply Chains

Outside of AI, plenty of research has argued that individual knowledge alone is insufficient in a multi-stakeholder context. A single stakeholder cannot be aware of all relevant knowledge within a large multi-stakeholder initiative. Hence, getting acquainted with the perspectives of the other stakeholders fosters reflections on one’s own knowledge blind-spots [9, 59, 131, 145], while developing a common understanding of a problem enables dialogues and fosters effective and inclusive multi-stakeholder initiatives [59, 65, 147, 148]. Direct and offline knowledge exchanges within and across organizations particularly facilitate collaborations or reliance and ultimately ensure the production of value [29, 55, 60, 114, 119, 147, 149]. Methodologies such as participatory design [107], contextual inquiry [74], value-sensitive design [56], and empathic co-design [131] rely on these ideas of collecting and sharing knowledge from diverse stakeholders.

Producing AI systems is inherently a multi-stakeholder endeavor [103], that shapes how responsible the systems are [45, 47, 103, 121]. No single individual can independently manage the range of required expertise [6], data [156], interdependent software components [155], and computational infrastructure [146, 152]. The concept of “AI supply chain” [33, 38, 51, 155] alludes to this interdependence of actors, and their collaborations and collective deliberations, e.g., between ML developers and domain experts [97, 100, 113, 117], business owners [115], UX practitioners [108, 140, 151], or potential end-users [21, 43, 161]; and researchers call for even more collaborative practices [73, 76, 117, 120, 160]. LLM systems are especially concerned with such multi-stakeholder efforts as LLM technologies require even more resources than prior types of AI systems [12, 152].

³“Knowledge” has been explicitly mentioned in the context of AI to refer to a) information used to develop rule-based AI systems [104] or fed into the training of machine-learning algorithms [44], or b) when AI is used as a tool to improve knowledge management [27]. Such contexts and the knowledge involved -domain knowledge- is different from our work.

⇒ Aligned with the non-AI-focused literature, prior studies argue that practitioners in AI supply chains could similarly benefit from direct knowledge exchanges [98, 117, 140] as they may lack awareness of each other’s roles and values [33, 143, 155], and regulatory efforts enforce such exchanges between AI provider and deployer organizations [99, 102]. In this work, we embrace this multi-stakeholder perspective and study *how* knowledge exchanges transpire across the supply chain, and aim to explore and identify key challenges and limitations in current practices.

2.3 Knowledge Taxonomies as a Facilitator of Knowledge Exchanges

In CSCW and organizational scholarship, coordination across diverse stakeholders often depends on shared artifacts that allow people to produce a shared mental model of relevant knowledge [140] without requiring full agreement. Boundary objects [137], materialized as “repositories” [3, 137], common knowledge spaces [3], or ontologies [9], serve as communication intermediaries by stabilizing notions and vocabularies while remaining flexible enough to accommodate local interpretations. In particular, *knowledge taxonomies* [84] arrange and log knowledge comprehensively using a classification scheme. They provide an easily interpretable knowledge map for uncovering persistent gaps, and facilitate synergies and reveal stakeholders’ knowledge blind-spots [74, 107, 131]. They are seen as facilitators of multi-stakeholder initiatives and knowledge exchanges in organization management [65, 147, 148], HCI and CSCW research [73, 107, 131]. They also constitute the basis for most knowledge logging and retrieval tools from prior organization management, HCI and CSCW works [3, 34, 69, 77, 80, 101, 112, 128].

In the context of AI, the few studies that explore multiple AI stakeholders [47, 103, 121, 122, 145] do not comprehensively describe their knowledge needs, but show some challenges for sharing knowledge—for instance what to document remains uncertain for upstream actors of the LLM supply chain who tend to restrain the documentation to technical considerations only [142]—and call for more common knowledge [100, 113, 140]. More broadly, AI researchers and policymakers propose structuring principles [64, 106, 154], frameworks and guidelines [15, 109], toolkits [16, 85], checklists [93, 123], and other documentation tools [8, 57, 105], to support responsible AI development within and across organizations. These artifacts all implicitly rely on an underlying knowledge taxonomy, and constitute accessible vectors of knowledge around AI fairness, explainability [49, 133, 141], and transparency [7, 8, 52, 57].

⇒ Recent meta-research in HCI has shown that taxonomies and other descriptive frameworks are essential for creating shared vocabularies that enable coordination across heterogeneous stakeholders, particularly in complex, multi-actor settings [54]. Building on this, our empirically grounded taxonomy provides the structured conceptual scaffolding needed to support consistent knowledge exchange across the LLM supply chain.

3 Knowledge Practices in LLM Supply Chains: Qualitative Insights (RQ1)

3.1 Method

3.1.1 Semi-Structured Interviews. To study knowledge exchange practices across the LLM supply chain, and identify the knowledge nuggets used or needed in relation to LLM systems and responsible AI topics, we conducted semi-structured interviews.

Participants. We interviewed 71 participants across 12 private organizations. Participants were recruited via our professional network with snowball and convenience sampling. We ensured that all participants are involved in developing or deploying LLM systems, and span a wide variety of roles within organizations. We also ensured to recruit multiple employees from the same organization, and from multiple organizations working on the same LLM supply chain to reflect on potential synergies in-between participants or organizations. The recruited participants cover a plurality of

domains in the LLM supply chain, including engineering (LLM developers, LLM researchers), user experience (user research, content design), business (product managers, business analysts, customer service), and governance (legal and risk teams, government relations); see Table 1. Our participants' organizations participating in the production of LLM systems span various specializations across the LLM supply chain, e.g., data annotation or LLM fine-tuning. The organizations deploying and consuming the LLM systems span various sectors, e.g., banking or insurance.

Characteristics	Values and counts
Location in the chain	LLM provider (48), software application consumer (19), data and annotation provider (2), end-user (2)
Primary role	Legal and governance (10), business-oriented (10), product manager (9), UX research (9), AI research (8), AI engineer (7), software engineer (5), ML governance (5), UX design (4), customer service (2)
Gender	Man (43), woman (28)
Work location	US (34), Canada (16), India (8), Netherlands (7), Switzerland (4), Italy (1), Belgium (1)
Educational background	Computer science and AI (20), computer science (15), UI, UX and psychology (13), legal and governance (12), business (9), no particular study mentioned (2)
Past experience with AI	- Developing AI algorithms (solely): 2- years (2), 2-3 years (6), 3-5 years (5), 5+ years (7) - Designing around AI algorithms (solely): 2- years (7), 2-5 years (4), 5+ years (2) - Selling or adopting AI systems (solely): 2- years (4), 2-5 years (6), 5+ years (2) - Managing AI development teams: 2- years (1), 3-5 years (4), 5-7 years (3), 7+ years (1) — 7 of these participants also had between 1 and 13 years of experience developing AI algorithms - Governing AI systems: 2- years (2), 3-4 years (11), 5+ years (2) — 4 of these participants in the 3-4 years category also had respectively 10, 6, and 3 years of AI development or 5 years of designing with AI

Table 1. Statistics about our interview participants. We do not specify further the amount of experience with LLMs, as all participants have been working with LLMs for the last three years maximum.

Procedure. Stage 1: We asked our participants to describe their day-to-day activities related to LLM system design, development, deployment, or use, their understanding of the organizations and stakeholders related to their LLM supply chain, and how they might work or communicate with, or simply be impacted by these stakeholders (including through which means, e.g., documentation). From these discussions, we started extracting some of the limitations in their knowledge, their knowledge acquisition and exchange practices, and the challenges they face. Stage 2: We fostered explicit reflections about knowledge by directly asking for their broader thoughts related to data, information, expertise, and knowledge, in terms of usages, wishes, and needs. We also prompted them to discuss their primary concerns, challenges, and needs concerning responsible LLM supply chains (questions that we broadly framed around the notion of “multi-stakeholder efforts”), as such discussion could reveal additional knowledge needs and gaps. For instance, we asked them to talk about the potentially harmful impact of LLM systems and their wishes and contributions around that, which triggered some participants to describe their desire to act on AI harm and their lack of knowledge to do so. Stage 3: As the following responsible AI notions are highly concomitant with knowledge, we probed their understanding of explainability and transparency: what these notions evoke, how important they are, and what the entailing challenges and wishes might be. We also enquired about their AI literacy, avoiding leading questions about their prior experiences with AI and their type and level of understanding of AI, to identify the types of knowledge nuggets they refer to when describing their own literacy.

3.1.2 Interview Data Analysis. We conducted an inductive thematic analysis [144]. After creating anonymized transcripts for each interview, the authors familiarized themselves with a few of them. Then, the first author coded each transcript for insights pertaining to RQ1. They especially paid attention to quotes related to the pros and cons of knowledge exchanges, the modes and format

that exchanges take, the tools and processes that structure and might facilitate these exchanges, the explicit or implicit limitations of these exchanges, and the challenges faced and wishes expressed by the interview participants in this context. In discussion with the authors and with a particular inclination towards the opportunities and barriers created by the LLM supply chain, the first author refined these codes and clustered them in order to create themes that summarize the insights. We now describe our five main themes.

3.2 Knowledge Opportunities in the Supply Chain

Across the interviews, we identified two key limitations in the ways the actors of the supply chain formulated and fulfilled their knowledge needs, that the supply chain can help address. In the supplementary material, we discuss in detail cases where practitioners miss the opportunities provided by the supply chain, and the implications of the resulting knowledge gaps (e.g., misuse of LLMs, duplicated efforts).

3.2.1 Incomplete formulations of knowledge needs. The actors of the supply chain are not always aware of their own knowledge needs. By this, we do not mean that they are unable to obtain information they know they need, but rather that they often frame their information requests too narrowly, overlooking additional knowledge that would be necessary to interpret that information responsibly. For instance, most participants who decide whether to deploy an LLM system (e.g., product managers) discussed requesting high-level information about its performance such as accuracy scores or toxicity levels. Only a few of these participants further looked at the conditions under which these metrics were produced, such as the benchmark used to evaluate performance and reflected on its limitations, e.g., a participant reported that LLM systems might be trained and evaluated using the same benchmark samples over several months, which can lead to overfitting, and so they typically inquired about how the benchmark is used concretely. Such incomplete formulations of knowledge needs became particularly visible when participants discussed potential blindspots of other actors in the supply chain. For example, P25 (product manager) pointed that many customers deciding which LLM system to adapt only look at accuracy scores computed by the LLM providers, and consequently struggle to interpret accuracy for their decision-making tasks:

“Accuracy is a challenging concept for people because they don’t have any AI expertise and go through their gut sense of it. Look at [LLM] summarization, and how there can be different types of inaccuracies, and how much it’s subjective to the person reviewing it. We still hear from customers ‘we won’t do it if it’s not 88% accurate.’ But how are you even defining it? ‘Well, if the end-user thinks it was right and solves their problem.’ But that is over simplified: the LLM could miss information, it could have hallucinated and made up something completely, it could have misinterpreted something.”

On the positive side, participants noted that interactions with other actors in the chain often helped them identify incomplete knowledge needs. For instance, P25 also explained how they support user researchers by helping them realize limitations in their study designs:

“A lot of UX researchers who haven’t worked in the AI space do not know that AI is probabilistic.⁴ If we aren’t helping them do their studies or come up with questions for when they’re assessing a new experience that has generative AI in it, they don’t even know to ask if users think the system is going to be correct. We help them know the limitations of AI to do informed research.”

⁴Note that in this quote and future ones, the participant referred to “AI” and not specifically to “LLMs”. Yet, all participants are involved in LLM supply chains. This mismatch in vocabulary employed is sometimes attributed to the non-technical participants lacking a good understanding of the specificity of LLMs, and sometimes to the participants having prior experience with other types of AI systems and reflecting about their practices for such systems too.

3.2.2 Unmet knowledge needs. Because knowledge is fragmented across the supply chain, practitioners may recognize their own needs but still be unable to meet them. This is evident in an excerpt from P5 (UX researcher) who expressed not having enough knowledge about the technical aspects of the LLM systems being developed:

“I don’t know what we do on the back-end to put safeguards in place. I don’t feel like I know a lot about how we’re developing the LLM, what we’re training it on, what we put in place to mitigate unintended consequences. It would be good to understand if we want to report back on certain things within the experience designed for the end-user.”

Fortunately, the interviews revealed that the supply chain creates opportunities for practitioners to draw on each other’s expertise. This dynamic was evident in the many knowledge-sharing practices participants described—both those already embedded in their daily work and those they wished were more common (notably, every participant mentioned either receiving or sharing knowledge with others). For instance, P56 (quality assessor) described having developed expertise in testing LLM toxicity but lacking the domain knowledge needed for more comprehensive evaluations:

“I know there is hate against Muslim folks. And I input some random text about that. And I would like the model to not respond when the user is saying something that they shouldn’t be speaking of. But I have only limited knowledge. With more people coming in, they will have more ways of testing how our model should be responsible enough for this problem, and I can develop the evaluations further.”

3.2.3 Subjectivity of knowledge needs. Note that participants did not always agree on what constitutes a genuine knowledge need. In some cases, individuals in similar roles articulated essentially the same need, even if they described it using different examples or terminology. For instance, across interviews with UX practitioners, we observed a shared need for technical knowledge—though they referred to it variously as knowledge of the “backend”, “algorithm”, “mechanisms”, “training”, “dataset”). In other cases, participants described different knowledge requirements for achieving the same goal, without any requirement being clearly wrong. For instance, one developer preparing a fine-tuned model emphasized the importance of understanding existing base models and their performance characteristics. Another developer, working toward the same goal, focused instead on knowing which organizations had developed particular LLMs, preferring to select a base model based on the practices and offerings of its developer rather than on model performance alone.

3.3 Modes of Knowledge Exchanges in the Supply Chain

From our interviews, we identified different triggers, drivers, and formats of knowledge exchanges.

3.3.1 Proactive and serendipitous exchanges. Participants were often reflexive about their own knowledge gaps. For example, P5 acknowledged:

“At the back-end like fine-tuning algorithms, I don’t understand much. I’d be fully transparent here. I have still more learning to do.”

In such cases, participants proactively sought to bridge these gaps. For instance, P9 (UX researcher) described how they relied on the specialized responsible AI knowledge of their colleagues to progress in their work:

“Even though I didn’t have specific background in AI, I’ve been really interested. And I’m learning from the people on this team who have deep knowledge on AI. My manager has a lot of knowledge, she’s been working in AI for a few years and GenAI recently. So getting to learn from them has been cool.”

In other cases where knowledge needs are incomplete, being exposed to some communication channels enabled the participants to realize that additional knowledge could be useful. For instance, P56 (quality assessor) explained that informal discussions with other teams enabled them to serendipitously uncover ways to evaluate the LLMs more comprehensively:

“I got some ideas from discussions with others or the groups I’m participating in. In one meeting, the research team talked about our models not detecting the difference between Canadian French and the regular French, how it offends people... This gave me new valuable information. So I try to listen to any sorts of discussion anywhere that I can get my foot. That’s how I got into testing this profanity.”

3.3.2 Exchanges initiated by knowledge seekers or holders. The preceding examples illustrate exchanges initiated by actors who self-identify a knowledge gap. Our interviews also revealed exchanges prompted by actors who possess additional expertise. On the one hand, many of these practitioners within the supply chain proactively anticipate potential blind spots among other actors and share relevant knowledge to prevent misinformed decisions. For instance, P12 (member of a responsible AI team) explained proactively meeting with engineering teams to teach them about relevant evaluation metrics that otherwise might be overlooked:

“If I know there’s an effort that’s starting in engineering, I bring my knowledge of evaluations and metrics that we need to do responsible AI. You have to be specialized enough to do that, have read lots of literature and know about the current models and current datasets. And then reflect about the LLM task to tell ‘it should be nice to use this metric or dataset to test if the summarization task is good or toxic.’ That won’t come naturally to business units or engineering, they will evaluate the accuracy (whatever that means) but not the other axes of responsible AI.”

On the other hand, some participants explained serendipitously identifying the knowledge needs of others. For instance, a developer (P24) regularly has random chats with other practitioners and identifies their limited practices and knowledge:

“Recently someone told me that they showed to their manager some machine learning metric that they thought was a good metric for their problem. And the manager said: ‘I don’t understand this metric, just pick accuracy.’ I told them that doesn’t make sense, it’s not because the manager didn’t understand the metric that it’s the wrong metric. Pick metrics that make sense for the problem you’re solving.”

3.3.3 Diverse formats of the knowledge exchanges. We observed that a single knowledge need can be addressed through diverse exchange formats. These include oral or written conversations, email exchanges, as well as one-way communications (e.g., presentations, company newsletters). Exchanges can occur directly between actors or be mediated through multiple chained interlocutors. For instance, P3 (user-experience designer) discussed how developers should explicitly communicate to user researchers the limitations in the LLM system, that the UX teams should then communicate to the end-users:

“Anything involving model drift and bias needs to be overseen by anyone who’s in the back-end. That’s harder for someone like me or a designer to have influence on. Those limitations within the model have to be reported to UX design and UX research. We [UX teams] can then communicate that to the end-user in the user-interface. For example, if a model has a very significant error rate for a specific task, we need to make sure that the end-users are aware of that, and they’re not over-relying on it. It’s a pipeline that needs to be managed across different teams.”

3.4 Facilitators of Knowledge Exchanges in the Supply Chain

We identified three approaches to foster knowledge exchanges in the supply chain.

3.4.1 Shared spaces. Shared spaces enable a reciprocal exchange of knowledge, and take the shape of workshops, conferences, and regular presentations. For instance, P10 (program manager) described various existing and potential avenues for supporting knowledge exchanges within the LLM provider they work for:

“One day, all parts of the organization gave their roadmap. All employees were invited and that put everyone on the same page. If we could do something like that for generative AI, from basics of ‘this is how we incorporate our products’ to establishing ‘if you want to do this, this is how you do it’, or if we can start pointing people in the right direction for what’s happening now, it’ll get them in the mindset of knowing there are things taking place. That would help establish a central forum for people not to repeat work. Also, every week we have presentations where customers explain how they’re using our LLM products. When they ask ‘will you develop responsible AI?’, our employees could then have a message to give.”

Several participants reported on efforts to create interdisciplinary reflection groups or boards that force communications across departments, and progressively enable practitioners to make more informed decisions. For instance, P52 similarly described:

“We don’t have a single person whose job is responsible AI. A small stakeholder committee is the best we can do, with a representative from each of our main areas (research, engineering, user research), to cross-check each other and fill in each other’s gaps and knowledge. That’s what’s gonna create a much stronger governance.”

3.4.2 Intermediary roles. Certain practitioners play the role of intermediaries to connect practitioners whose insights could benefit one another. For instance, several participants pointed to product managers as having the relevant connections within the organization and the relevant understanding of AI to serve as middle-men to create conversations or facilitate them: (P34)

“Product managers are at the crossroads of technical requirements, visual design, product design, UX research, content design, etc. They’re connected to everyone whereas, for example content designer might not really interface with developer. So they should spearhead efforts.”

In some cases, these roles are becoming more structured, such as P17 who explained their own work:

“We have people like me, we’re called strategists. We act as communication layers between different orgs. We have technical and business strategists that facilitate between the customer and applied research. And then me that communicates between the product design and the product creation side and the applied researchers.”

3.4.3 Tools. Participants mentioned several kinds of tools that supported uni-directional knowledge exchanges. On the one hand, they highlighted various educational programs—both those currently in use and those they wish were available. Such programs could consist in reading lists; onboarding sessions for specific responsible AI-related tools (data governance ones); or learning sessions when joining the company and other trainings depending on one’s role, e.g., P18:

“You need to educate people so that they adopt some level of responsible AI in the process. Like the product manager who is writing the requirements and including the acceptance criteria — ‘dear developer, make sure that you’re following this policy or style for implementation’. And then for people who are deploying the LLM, it’s like ‘don’t forget to add

additional software processing to filter for PII, and policy compliance like pre- and post-processing’.”

On the other hand, many participants discussed the need for written communication about their organization or its LLM systems. Beyond learning from regular company or department-wide communication emails, they emphasized the need for documentation, such as model cards. For instance, P2 in the policy team explained that while each actor of the supply chain plays a different role in building the LLM system, one last actor should document the work of all these actors:

“The problem with responsible AI is that it’s too broad for one single actor to know all about it, and to ensure that it gets infused in every aspect of it. So ideally responsible AI would be broken down in different components that different actors would have responsibility for. And engineers would document for transparency ‘Here’s what it was tested for, the data, the foundational model, the whole process through which we went to create this model before it was shipped.’”

3.5 Barriers to Knowledge Exchange within the LLM Supply Chain

Our interviews also revealed that knowledge exchange is concomitant with several challenges.

3.5.1 Lack of organizational visibility. We found that a general lack of organizational visibility leads to knowledge remaining untapped. On the one hand, interview participants in need of knowledge mentioned not necessarily knowing where to look for it. For instance, P5 explained:

“I didn’t know till talking to you today that we had a governance team. Working at other companies, I knew that there was a team dedicated to these questions around responsible AI, I would know who to go to for specific questions, and they would regularly communicate out things they thought were important for employees to know about.”

Similarly, participants who do possess knowledge pointed to in-need participants not knowing about them. For instance, multiple UX researchers lamented over product teams not being even aware of the user-research department, and repeating studies (with lower quality) that they had already conducted to fulfill their knowledge needs. On the other hand, several participants who possess potentially insightful knowledge reported that they did not always know who might need this knowledge although they would be open to share it. P29 (program manager) explained the need for organizations to create more structured spaces to remedy to this gap:

“Sharing information across the enterprise about what you’re doing is really important. The 5 W’s: who, what, where, when, and why? Especially who’s doing what, who are the groups that are working with the AI, you gotta have that inventory.”

A few participants proposed that creating structured documentation template could help them understand what they should log in for others to consume (which was also subject to further interrogations, e.g., when to update the documentation, how feasible it is to keep this documentation up-to-date continuously). The taxonomy we propose in Sec. 4 can serve as such a template.

3.5.2 Lack of mutual understanding. Our interview participants discussed various ways in which they would need to better understand the actors of the supply chain to be able to better communicate knowledge with them. They described basic communication issues across teams: (P24)

“I was working on a project and heard about this team that does the same thing. I scheduled a call to understand what they’re doing. And I understood only 50% of what they’re doing.”

Such misunderstandings are reinforced by many instances of conceptual confusions and terminological ambiguities. For instance, P3 who studies customer needs explained facing difficulties in

choosing the right words for appropriately communicating with non-LLM experts (the taxonomy we propose in Sec. 4 could serve as a reference point to clarify terminologies across stakeholders):

“It’s difficult to land on a shared vocabulary. That makes it very challenging to communicate and conduct research. When I did a research project on customer-service agents and their perceptions and risk tolerance to generative AI creating inaccurate results, it was really difficult to figure out: do I call it inaccuracies, or hallucinations, or signals? Which language to use when the agents are not even really familiar with these concepts?”

Our participants discussed that better knowing the other actors would help in developing appropriate communications. For instance, they emphasized the difficulty for them to gauge the extent of the knowledge one actor already has, and hence to what extent they should work on communicating further knowledge. Some of them reflected on the organizational structure, explaining that being in the same team further helps to create opportunities to understand the need of the others better.

3.5.3 Unclear allocation of responsibility. Participants highlighted the lack of clear responsibility allocation pertaining to knowledge sharing. It was typically unclear who is in charge of documentation, or of creating documentation tools and other knowledge exchange programs. This was partially explained because it remains unclear who is best suited to do so (e.g., would they need to have the knowledge to create spaces for knowledge exchange?). P18 (program manager) for instance pointed out that nobody is responsible for enforcing appropriate logging of knowledge in their organization:

“Unless best practices for documentation are part of the software development lifecycle, it’s not prioritized. And trying to do catch up after the fact doesn’t really work, because by that time, there’s a good chance that the people who have got the actual knowledge have either left the company, left the team, or forgotten.”

P24 (developer) further explained that due to this lack of responsibility allocation, some actors of the supply chain simply consider that they do not have the time to share knowledge with others. Unclear allocation of responsibility can also create misalignments across actors. Some participants discussed that other actors assume they are responsible for providing them with knowledge, yet these assumptions are not always obvious to others. P64 who works for a data collection and annotation company shared:

“The client [LLM developer] usually doesn’t know themselves what they need. And they rely on us to provide them with good data. And one part of this puzzle is writing the instructions / the guidelines for the data annotators.”

4 Taxonomy of Knowledge Needs Across the LLM Supply Chain (RQ2)

Our interviews revealed various knowledge nuggets⁵ that are essential for actors across the LLM supply chain to perform their work. We synthesize them into a taxonomy that organizes them into key categories. Unlike a checklist or documentation template that presupposes which exact nuggets matter, a taxonomy supports the exploration of “unknown unknowns” by making categories of possible knowledge visible. And unlike a repository that accumulates content without structure, a taxonomy enables purposeful search, comparison, and maintenance by stabilizing categories. Hence, a taxonomy is particularly relevant for the LLM supply chain because it can help practitioners address several challenges identified in our qualitative analysis (cf. Sections 3 and 2.3). For instance, it can act as a tool to help them log useful information systematically, facilitate *purposeful* searches

⁵Generalizing the term “information nugget,” which typically refers to factual knowledge, we adopt the term *knowledge nugget*—already employed in education and in knowledge management [48]—to refer to any valuable piece of knowledge that an LLM practitioner may require.

of relevant knowledge and mitigate organizational opacity, expand opportunities for *serendipitous* discovery of knowledge nuggets, and establish a shared vocabulary to reduce *miscommunication*.

4.1 Method

To achieve these goals, we define requirements for the taxonomy along two dimensions (R1, R2):

- **(R1) Format:** The taxonomy should be organized so as to:
 - (a) **Unify the knowledge needs of diverse practitioners** unlike existing resources that target specific roles (e.g., toolkits aimed at AI developers solely [4, 20]) or narrow objectives (e.g., gaining understanding about the AI model [141]).
 - (b) **Provide a unified vocabulary** that clarifies the interpretation of knowledge nuggets.
 - (c) **Offer a simple structure** that enables easy discovery and logging.
 - (d) **Be extendable**, allowing the addition of new knowledge nuggets as needed.
- **(R2) Content:** The knowledge nuggets included should:
 - (a) **Cover a broad range of actors** in the LLM supply chain, including those with limited AI literacy, those in organizational or policy roles [94, 122], in contrast to prior work focused on end-users [28, 52, 81, 87, 116], developers [68, 140], or UX designers [88, 157, 159].
 - (b) **Encompass diverse knowledge types**, avoiding the narrow scoping seen in prior research (e.g., transparency or explainability [57, 105, 133, 141] or AI expertise [93]).
 - (c) **Pertain specifically to LLM systems**, defined as systems powered by LLMs to deliver services to customers, ensuring rigor and focus.
 - (d) **Relate to responsible AI**, limiting scope to knowledge relevant for developing safe, fair, transparent, and accountable LLM systems that comply with regulations. [31, 70, 99, 130]. For example, this allows us to exclude knowledge about the underlying hardware infrastructure, despite its relevance to LLM development.

To meet these requirements, we developed a set of dimensions that distinguish practitioners' knowledge nuggets. These dimensions apply uniformly across all actors (R1a), enable simple organization of knowledge nuggets (R1c) while clarifying their meaning (R1b), and ensure flexibility (R1d) by allowing easy addition of dimensions or nuggets.

To establish these dimensions and populate them with nuggets, we inductively coded the interview transcripts for explicit or implied knowledge nuggets and the contexts in which they were needed (i.e., stakeholder roles and goals). We grouped similar nuggets across varying levels of detail (especially when appearing only a few times). We then conducted open axial coding to categorize these nuggets, identifying distinctions such as technical vs. socio-technical knowledge, or general vs. artifact-specific information. We allowed for multiple categorizations of the same nugget to serve different analytic needs, and acknowledged the categorizations made by our interview participants. To enrich our taxonomy, we reviewed recurring and diverging knowledge categorizations from psychology and education research [41, 63], and proceeded to a deductive coding round to check for overlooked knowledge categories. We finalized the taxonomy by reconciling inductive and deductive insights, ensuring our taxonomy dimensions are mutually exclusive and actionable (e.g., in relation to patterns between knowledge categories and their uses).

4.2 Overview of the Dimensions in the Taxonomy

We characterize the knowledge nuggets we identified using three dimensions summarized in Figure 2. We now explain these dimensions and illustrate them using the *example* in Figure 1.

4.2.1 Knowledge abstraction. This dimension (rows in Figure 2) goes from **specific** knowledge about one LLM supply chain that revolves around one engineered LLM system (e.g., *the components of the LLM system deployed at the bank, such as the dataset or the interface designed for user interactions*

		Theme			
		Technology	Production	Use	Context
Abstraction	Generalized Knowledge	□ LLM functioning ✓ Responsible AI concepts ➤ Human-AI collaboration frameworks ❖ Governance frameworks	□ ✓ ➤ ❖ artifacts ↳ Skills relevant to processes (doing, reflecting, instrumental)	□ LLM (dis)advantages ✓ Harms ➤ LLM use-cases ❖ Industrial sectors of application	□ AI research conventions ✓ AI policies & regulations ➤ Perceptions of AI ❖ Organizational players
	Specific Knowledge	□ ✓ ➤ ❖ artifacts ↳ Description ↳ Qualities (e.g., results of evaluation)	□ ✓ ➤ ❖ artifacts ↳ Process (design, development, deployment, testing) ↳ Actors involved	□ ✓ ➤ ❖ artifacts ↳ Usage (user, amount of use, purpose of use) ↳ Impact (degree of satisfaction, externalities)	□ System requirements, intended use ✓ Applied policies and regulations ➤ Customer needs, end-user characteristics ❖ Constraints of deployer & provider

Lens			
□ Technical lens (artifacts, e.g., base-model, data flow)	✓ Socio-technical lens (artifacts, e.g., safeguards)	➤ Interactional lens (artifacts, e.g., user-interfaces)	❖ Organizational lens (artifacts, e.g., appropriate use policies)

Fig. 2. Our knowledge taxonomy with its three dimensions—abstraction, theme, and lens—in blue italicized boxes) and their sub-dimensions (in gray bold boxes). The types of knowledge nuggets (plain text in the white cells) needed by the actors of the LLM supply chain are organized by the three dimensions. These types can be mapped one-to-one to a lens, or they can be mappable to different LLM artifacts corresponding to different lenses—in which case we indicated “artifacts” in the cells and used arrows to list the types of nuggets associated to these artifacts.

with the LLM application); to **generalized** knowledge about any LLM technology (e.g., *regulations around LLM that the bank might have to comply with*). In psychology, specific knowledge is typically termed “situational” [41], “descriptive” [63], or “declarative” knowledge [58], as well as “know-that” or “information” [126] or “task information” [26], or “know-what” [95].

This dimension is especially useful to bring clarity in certain terms used by our participants. For example, it was striking to observe that some participants understood the notion of *transparency* in terms of provider organizations (a) sharing *abstract knowledge* about LLMs to the deployer organizations (e.g., how do LLMs generally work), or instead *specific knowledge* about a particular LLM system, be it about (b) disclosing that an LLM end-user is interacting with an LLM system, or (c) communicating information about the different artifacts that make up the LLM system at hand.

4.2.2 Knowledge theme. This dimension (columns in Figure 2) describes the kinds of questions a knowledge nugget answers. a) About the underlying **technology**: *what* kind of engineering, scientific, or procedural *components* or *principles* are involved with a given LLM system or general LLM technologies? b) About the technology **production** approach: *how* is or can an LLM system be *produced* from design to implementation? c) About the **usage** of the technology: *how* can an LLM system be *used* or is typically used, and what are the impacts thereof? d) About the technological **context**: *where* does or can the production and use of LLM systems take place (e.g., with which motivation, under what kinds of constraints)? *For instance, the knowledge about the LLM components (e.g., the data, annotations, and base model), the knowledge developers make use of for mitigating LLM unfairness, the harm caused by the LLM system, and the normative expectations of the bank respectively correspond to technology, production, usage, and context-related knowledge.*

Again, this dimension is useful to further clarify the responsible AI notions that participants used. For instance, when participants referred to transparency as specific knowledge, some imagined sharing the dataset specifications or the model specifications (underlying technology) while others thought of sharing how the dataset and model had been built (production knowledge). As for

participants' notions of explainability, some understood it as the provider providing the deployer with a simple and basic understanding of the functioning of LLM models (underlying technology), while others saw it as 'algorithmic transparency' displaying the online behavior of a specific LLM system (usage knowledge), i.e., what kinds of inputs are entered by users, what kinds of outputs the system returns and what mechanisms it uses to return these outputs (e.g., related to works on local explanations [86, 132]). Additionally, while it was not necessarily termed "explainability" or "transparency", several LLM deployers we interviewed expected the providers to inform them about the AI regulations they followed to develop the systems (contextual knowledge).

4.2.3 Knowledge lens. This dimension (bullet points in Figure 2) specifies the types of system artifacts a knowledge nugget refers to or characterizes the flavor of a nugget. This goes from the **technical** underpinnings of LLM technologies (e.g., *the technical LLM components such as the particular base LLM fine-tuned by the LLM developer, or the basic functioning of an LLM*), to the **socio-technical** embedding of this technology in society (e.g., *the bias mitigation methods applied to the LLM, the harms that LLMs might cause*), to the **interactions** with the technology (e.g., *the user-interfaces designed to warn end-users of potential biases, the perceptions of LLMs that end-users typically have*), to the **organizations** and organizational processes around it (e.g., *the data stewardship processes and AI principles of the provider, AI regulations*).

This dimension reflects distinctions that actors of the LLM supply chain often operate (which increases familiarity with our taxonomy and facilitates its use) when discussing specific roles in the chain. For instance, UX researchers and data scientists were often said to be in charge of respectively the "interactional" knowledge around the end-users, and the technical artifacts of the LLM systems. Policy stakeholders were typically envisioned to have a better understanding (and be responsible for) the governance processes in an organization, and dedicated responsible AI teams were tasked to infuse socio-technical responsible AI notions across the supply chain.

4.3 Detailed Description of the Knowledge Categories

We now describe in more detail the knowledge categories and associated knowledge nuggets. Additional examples for each of the categories can be found in the supplementary material.

4.3.1 Specific technology. This category describes the artifacts of the four lenses that are implemented within the LLM supply chain. For instance, P33 discussed specific needs of the system deployers:

"A great degree of transparency in how features are implemented [is needed], at the level that we can say: if you do this thing in the user interface [interactional], here is what the data flow looks like [technical]. This would matter in our adoption processes."

The provider's governance team (P10) explained asking clarifications about the ways the LLM outputs are treated by the deployer, and especially wanted to know whether the deployer had set up any output filtering process [socio-technical] and any employee training to foster responsible LLM usage [organizational]. These descriptions can include information about the performance and limitations of these artifacts: (P7)

"Algorithmic transparency is telling customers the algorithms used and their limitations."

4.3.2 Specific production. This category shows how the LLM system was produced. Its nuggets pertain to each LLM artifact and the phases they go through (e.g., design, development, deployment, updating, testing). They include descriptions of the choices, justifications thereof, and information about the actors involved in each phase (to bridge the lack of organizational visibility). For instance,

a customer service representative (P60) discussed needing deployment-related knowledge from LLM product managers:

“I would like to see these managers tell me they’ve tried the AI out with some customers first to see how helpful it actually was, before deploying it broadly.”

4.3.3 Specific use. This category focuses on the use of the artifacts. It includes knowledge about the users of each artifact, their goals, and their amount of use; and knowledge about the impact of the artifacts, in terms of user satisfaction or potential negative effects they might have. For instance, P8 (user researcher) emphasized the importance of knowing about user satisfaction to develop products:

“What to build, ideally, it should be impacted by customer need and the quantitative data on the usefulness.”

P30 (customer satisfaction) hinted that how end-users use the LLM system might differ from its intended use (which might be relevant information for policymakers):

“We provide them [customers] with a product intended to work in a specific way. If we educate them appropriately, but then they decide to bypass some of the guardrails, then it should be their responsibility.”

4.3.4 Specific context. This category specifies the influences for the system that is developed, such as information about the end-user (e.g., their needs and expectations), the constraints of the provider and deployer organizations that impact the system production, or the regulations that the system is built to follow. For instance, P6 (UX researcher) investigates user needs to design the systems:

“We do discovery to understand what customers are looking for and their needs. And then concept testing to see their thoughts.”

On the technical lens, this category specifies the intended use of the system and its requirements. For instance, P18 (program manager) explained sharing with engineering teams how the system should function:

“In your requirements management system, you write: as a user, I want to do X, Y and Z. The best practice is to provide some acceptance criteria and context, so that the person going to develop the requirement has not only the use-case and user story.”

4.3.5 Generalized technology. This category refers to the conceptual ideas that rein the functioning of LLM systems. The technical nuggets are related to the theories behind the functioning of LLMs (sometimes referred to as the “laws of nature”, “mathematical reasoning”, “know-why”, or “conceptual” knowledge) [41, 42, 95, 124]: the basic LLM artifacts, in-depth knowledge (e.g., about the different LLM algorithms that exist, the bounds and guarantees that LLM algorithms have), or the main types of existing AI technologies, as several participants had to explain to prospective deployers the differences between predictive and generative AI to talk about the advantages and risks of the LLM systems they offer. The socio-technical lens calls for nuggets about the responsible AI concepts developed to cater to issues in the outputs of LLMs (e.g., unfairness metrics, adversarial attacks). For instance, a technical developer (P32) wondered:

“Is it better to first have a very accurate model, and then bits by bits, make it aligned and helpful? or is it better to first put some guardrails and train within that limited bound?”

The interactional lens invites for knowledge about potential issues in human-LLM collaborations and best practices to tackle these issues, and the organizational lens refers to existing governance

frameworks that have been established to oversee LLM systems (e.g., documentation tools, multi-stakeholder decision processes).

4.3.6 Generalized production. This category (often termed “know-how” or “procedural knowledge” or “prescriptive knowledge”) is related to the idea of experience-based learning, and the skills [41, 63, 95, 124, 126] necessary for the production of LLM systems. We identified three types of such knowledge. a) Instrumental knowledge refers to the awareness of and ability to use instruments that support work around LLMs. These instruments can be technical, such as evaluation benchmarks; socio-technical such as explainability toolkits; as well as organizational such as documentation tools for LLM systems. b) Knowledge for doing refers to best practices and rules of thumb for each phase of an LLM artifact. This is often tacit knowledge, that corresponds to “tricks of the trade” [63]. This can be technical skills, such as fine-tuning a model. For instance, an LLM researcher (P33) said:

“Selecting the appropriate loss function is an experience: you do that a few times, and then hopefully you’re good at it. And you guess right every single time.”

This can also be social skills, such as knowing how to best work with annotators to develop datasets by crafting appropriate data annotation instructions. c) Knowledge for reflecting refers to the ability to reflect on the preconditions and consequences of LLM work to ensure the morality of a supply chain activity [124, 126]. For instance, an engineer (P53) highlighted a typical gap in evaluation processes:

“One of the big questions is: how do we evaluate potential risk? Some of the ways we evaluate are algorithmic, but also there are some societal aspects. You have to think of the threats from a social perspective.”

4.3.7 Generalized use. This category refers to general indications about the typical uses of LLMs (e.g., their regular use-cases, who their predominant users are), about the potential (negative) impacts LLM systems might have, and their recognized advantages and disadvantages. For instance, a product manager (P13) in an LLM provider organization explained the need to inform LLM deployers:

“We have to empower customers for them to make informed decisions. So we have to make them aware of trade-offs, risks, and opportunities of AI.”

4.3.8 Generalized context. This category describes the context in which LLMs are being developed. This might be “domain information” [26], such as existing AI regulations, as well as knowledge of the main organizations playing in this field (e.g., leading LLM developers, annotation companies). This also refers to information about individuals, such as the typical opinions that society has on LLMs. Finally, it also encompasses knowledge of broader LLM and AI conventions, e.g., for scientific researchers to publish papers using the correct benchmarks, for stakeholders of different backgrounds to be able to communicate by adapting the vocabulary used, etc. For instance, a product manager (P20) explained the knowledge requests from deployers:

“[LLM provider] is considered the de facto large language model provider. So for everything that we need to develop, everyone –including the customers– are interested in knowing how well our solution performs against them.”

4.4 Validation of the Taxonomy

Method. To validate our taxonomy, we organized feedback sessions with practitioners of the LLM supply chain. We conducted four individual and three group feedback sessions (with around 15 people each). The group feedback sessions were advertised within one of the organizations as

informal in-progress-research talks about AI topics, especially AI research and LLM governance, focusing especially on information and knowledge needs and practices. The call for attendants emphasized the space that would be given to them to discuss with the presenters and provide comments. Employees interested in the topic could join the sessions freely. We presented the study and its findings in around 15 minutes and asked the attendees to provide any reflections (around 30 minutes), e.g., whether they recognize their or others' needs, whether any need is missing, whether any additional knowledge nugget might be missing, whether the categories resonate with them. Across sessions, feedback consistently indicated that the taxonomy was not merely a descriptive artifact, but usable as a shared structure for communication and governance ideation. Below we detail four recurring forms of demonstrated utility.

4.4.1 Specifying communications. The taxonomy proved useful in clarifying communications during group discussions: it helped make vague knowledge requests more precise, by letting discussants point to dimensions and categories. For instance, when a UX researcher expressed a lack of “technical knowledge” about LLMs, discussions initially risked remaining at the level of a general deficit (“I don’t know the backend”). However, using the taxonomy, a product manager encouraged them to unpack what “technical” referred to, and specify the type of knowledge nuggets they were referring to, using the **abstraction** dimension –e.g., whether the request concerned specific system artifacts (e.g., a particular data flow) or generalized understanding of how LLMs work–, and the **theme** one –e.g., whether it concerned technology facts versus production processes.

Similarly, the taxonomy helped participants resolve ambiguities in terminology related to responsible AI concepts such as *transparency*, by enabling participants to reference specific nuggets of interest and point to particular **abstractions** and **themes**. Examples provided in the previous section are illustrative of the misunderstandings that could have arisen.

4.4.2 Triggering self-reflections about blindspots. The taxonomy also prompted participants to recognize knowledge categories they had not previously prioritized or had fully ignored, effectively functioning as a structured tool for ideation. For instance, one LLM developer (P24) described how their model selection practice focused on training and test data, architecture, and evaluations (i.e., technology-focused knowledge). After we introduced the taxonomy to them, seeing the taxonomy’s distinctions between **themes** prompted them to consider whether model selection should also explicitly incorporate knowledge about downstream use cases and risks (use and context), reframing model choice as not only a technical optimization but also a responsible deployment decision:

“I’m thinking I could think more broadly about the use-case and the risks to make my choice.”

Similarly, a researcher (P32) whose work focuses on developing LLM assessments reflected that the taxonomy prompted new questions:

“It’s nice, it [the taxonomy] prompted me to reflect: was this model built for a specific purpose [context] that is different from what the benchmark is intended to measure? [production].”

Here, the taxonomy’s **theme** dimension, and especially the separation of production-related knowledge from technology facts, helped articulate a potential mismatch that would be difficult to express if “knowledge” was treated as an undifferentiated construct. While such limitations of performance scores might be well-known to researchers, they were not to our practitioners and the taxonomy helped them reflect.

4.4.3 Triggering reflections about others’ blindspots. The dimensions of the taxonomy were particularly useful to practitioners to guide the identification of the knowledge blindspots of others in

the supply chain. Several participants picked nuggets that they were aware of, and discussed how other practitioners could benefit from them. For instance, the generalized production knowledge triggered a researcher (P8) to lament the lack of responsible AI knowledge of product managers who evaluate LLMs and make deployment decisions. They used the **lens** dimension to point out that these managers typically only request information about the accuracy-related quality of the LLM outputs (*technical*) as they are not aware of other *socio-technical* metrics. In this context, they appreciated the comprehensiveness of the taxonomy and the **abstraction** dimension, that enabled to interrogate simultaneously the factual information that should be shared with others and the generalized knowledge needed to appropriately use this information.

The flexibility provided by the taxonomy design in terms of the granularity of the categories it creates (and the different sub-categories it implies) also revealed useful in this context. For instance, delving deeper into the technical knowledge, a participant with a governance background explicitly questioned whether to share with customers solely what's legally required by the AI Act or also what customers might request themselves. They liked the taxonomy's distinction between the different technical artifacts (examples of technical artifacts –i.e., sub-categories– were provided in the description of the **technical lens** –i.e., category), discussing whether the customers would need specific information about the base model, the fine-tuned model, or only the system built around the LLM.

4.4.4 Developing governance for knowledge exchanges. Several participants were triggered by our taxonomy to ideate about existing governance strategies that facilitate knowledge exchanges. For instance, one participant with a governance role was particularly excited by taxonomy as they recognized the ability the taxonomy would provide them to systematically identify the knowledge gaps of the employees in their organization (what they called “negative knowledge spaces”) as well as the knowledge owners (when the specific production information is populated as it calls for identifying production actors). They discussed that they could then develop effective and efficient remedies to the needs, such as common trainings for all practitioners manifesting a same need, or other targeted knowledge exchange spaces.

As for the development and revision of governance tools, a model owner (P13) who is involved in company-wide efforts for promoting transparency discussed how the taxonomy reveals potential extensions to existing documentation tools. The different categories it provides created a structured way for P13 to think through the large amount of potentially useful knowledge nuggets, going from category to category to reflect on its utility, and delving deeper into the examples of knowledge nuggets for the potentially relevant categories. For instance, when looking at specific knowledge (**abstraction**) related to technical artifacts (**lens**), they discussed:

*“Where datasets come from, how they were collected, whether there is any copyrighted material in them [examples from the production **theme**]. This is perhaps not important for the engineering team, but that’s probably important for other parts of the value chain and governance process. That makes me think we might want to include it in our model cards. [Instead they typically include technology-related knowledge].”*

5 Discussion

5.1 (RQ1) Fragmentation of Knowledge in the LLM Supply Chain: a Resource and a Challenge

5.1.1 The two facets of the LLM supply chain. We showed that the supply chain functions as a double-edged sword: it fragments knowledge but it also constitutes a resource for bridging knowledge gaps. From a CSCW perspective, this reflects a classic tension discussed in distributed cognition research,

where knowledge is necessarily spread across people, artifacts, and organizational boundaries, but coordination failures can turn this distribution into a liability rather than an asset [67, 129]. This finding also aligns with the few prior AI-focused studies that have shown that practitioners within the AI supply chain acquire knowledge to develop responsible AI systems via inter-personal learning and resource foraging [98, 134]. Beyond reaffirming the importance of knowledge work for responsible AI, we extend prior findings by demonstrating that these avenues are used to acquire both generalized knowledge (including competencies highlighted in earlier research) and specific knowledge, and by providing a concrete account of how interpersonal exchanges occur.

5.1.2 Modes of knowledge exchanges. We revealed how AI practitioners (could) learn from other practitioners, via proactive knowledge inquiries or serendipitous ones, triggered by self-reflections or the attention to blindspots that interlocutors displayed, and happening (a)synchronously. These findings resonate with CSCW work showing that knowledge often circulates through informal communities of practice rather than through formal repositories alone, as practitioners rely on peer networks to make sense of complex systems [25, 110]. We also identified the strategies that facilitate knowledge exchanges, such as the creation of shared spaces across or within an organization, the establishment of intermediary roles, and the development of tools for logging and spreading knowledge. While most of the responsible AI research is currently focused on documentation-centered tools for enabling knowledge exchanges [8, 57], these results broaden our horizon by highlighting other ways in which practitioners do or wish to do knowledge exchanges with the aim of producing more responsible LLM systems.

5.1.3 Barriers to knowledge exchanges. We found that limitations in the LLM supply chain challenge knowledge exchanges. These barriers resonate with different types of “knowledge boundaries” in CSCW [30], be it the lack of information to identify relevant knowledge (syntactic boundary), the divergent interpretations of a knowledge nugget (semantic boundary), and the incentives hindering the exchanges (pragmatic boundary). Some of these findings might seem unsurprising considering the challenges generally associated to supply chains that are considered to be catalyzers for the creation of untrustworthy AI systems, e.g., ineffective communications [134], the “accountability horizon” [33], and the lack of clear responsibilities [155]. Yet, we are the first to illustrate how these concepts play a role in knowledge exchanges. Interestingly, as many practitioners were cognizant of these barriers, they slowly started to engage in efforts to build knowledge exchange capacity as theories of infrastructuring in information technology outline [118]. Our work further showed not only that the knowledge taxonomy we offer can increase awareness of these barriers and reflexivity around mutual knowledge blindspots, but that it also constitutes a concrete tool that helps establishing governance strategies to tackle them.

5.2 (RQ2) Diversity of Knowledge Needs within the LLM Supply Chain

Contrary to prior research in CSCW and responsible AI or AI policies that often discuss the need for practitioners to have more knowledge about AI in an undifferentiated manner [2, 13, 99], our study provides an extensive list of knowledge nuggets (via the knowledge taxonomy) that actors across the LLM supply chain possess and might need.

Method. To explore in more details how the knowledge needs we identified compare to prior work, we collected 79 research publications that touch upon questions of AI knowledge needs and related responsible AI notions (i.e., AI literacy, explainability, and transparency), 7 policy and regulatory papers (e.g., AI Act [99]) that were published by governmental organizations across various countries and that are regularly discussed in prior research publications, and 7 AI principles and guidelines published by private organizations (e.g., Microsoft Responsible AI Standard [1]). We

mapped the content of these publications to the knowledge nuggets covered in our taxonomy, and reflected on the insights.

5.2.1 Coverage of knowledge nuggets. This exercise revealed that the findings in our study comprehensively describe the knowledge nuggets discussed in prior publications: we did not find any nugget that is not contained in our taxonomy. Besides, it highlighted the types of nuggets that have been neglected in prior works. For instance, despite all these categories being essential to some of our participants, we found that the literature primarily focuses on knowledge around *technical artifacts* (58% of publications), while only 4% of the publications deal with *organizational knowledge*. Furthermore, in contrary to *specific technology* knowledge that is mentioned in 83% of publications, the *specific production knowledge* (15%) and the *generalized contextual and use knowledge* (11%) had been discussed only scarcely in the past. Note that more recent publications start to recognize the importance of such knowledge [158]. These findings mirror those in our qualitative study and especially the knowledge gaps we identified among supply chain practitioners (Sec. 3), which is likely explained because most of prior work has not explicitly acknowledged the potential blindspots of practitioners when establishing lists of knowledge nuggets (they rely on self-reports of needs solely). Similarly, existing publications are centered around a few LLM artifacts (the ML algorithms and datasets), yet our participants also required information about other artifacts such as the data flows which are only reported in 15% of publications, or organizational processes that we did not find reported in prior works. Our more comprehensive account of knowledge needs is a first step to bridge the gap between the theoretical need for knowledge exchanges and its practical fulfilling.

5.2.2 Knowledge heterogeneity. Our findings also confirm that different practitioners need different knowledge nuggets, and show that the nature of the knowledge needed varies along multiple dimensions (abstraction, theme, and lens). This heterogeneity helps explain why existing governance tools centered around knowledge (e.g., model cards, evaluation reports) often fall short when applied across the supply chain. Not only do they only cover partially the needs of the diverse practitioners they are supposed to support, but their formatting also makes them challenging to exploit by these different practitioners. They rely, on the one hand on lists with too many irrelevant nuggets that increase cognitive load, and on the other hand on hard-to-interpret notions for the diverse actors whose AI literacy is often oriented towards role-specific abstractions. This interpretation of our findings further motivate why an easily-navigable knowledge taxonomy could constitute an appropriate alternative to existing tools –what we discuss next.

5.3 Knowledge Taxonomy as a Boundary Object for Supply Chain Practitioners

Our taxonomy is composed of three dimensions that form 32 knowledge categories (2 abstractions $\times$ 4 themes $\times$ 4 lenses), themselves referring to diverse knowledge nuggets. It offers analytical categories to concretely breakdown the plural knowledge nuggets exchanged in the LLM supply chain. It is the first taxonomy that is explicitly built to support knowledge exchanges across practitioners of the LLM supply chain –and Section 4.4 hints at its effectiveness.

5.3.1 Structure of the knowledge taxonomy. Our taxonomy is multi-dimensional. Instead, prior ones contain uni-dimensional knowledge categories (e.g., [32, 93, 139] for literacy, [8, 23, 57, 105] for transparency, and [133] for explainability). Prior works that do refer to categories of knowledge rely on some of the dimensions we proposed, corroborating the value of our dimensions (we did not identify any reference to dimensions that are not reflected in our taxonomy). For instance, based on a qualitative study, Madaio et al. [98] point to the “computational” and “procedural” orientations of the knowledge resources AI practitioners have access to –what our taxonomy also reveals with the technical and organizational lenses. Besides, prior frameworks cover only (subsets of) the

knowledge nuggets contained in sub-dimensions of our dimensions, e.g., generalized knowledge for most literacy frameworks and specific knowledge for documentation tools. While we can explain this limitation by the potentially more narrow goals of these prior works, the high coverage of our taxonomy revealed insightful as a prerequisite to meaningfully craft future actionable interventions.

5.3.2 Utility of the knowledge taxonomy to supply chain practitioners. Without the three dimensions in our taxonomy, it would have been more difficult for us to identify and interpret the missing knowledge in existing publications (above), and it would be more challenging for AI practitioners to conduct and organize knowledge exchanges (cf. Section 4). Our findings show that the central challenge in responsible LLM development is not merely awareness of disparate knowledge pieces, but the persistent difficulties practitioners face in accessing, translating, and aligning knowledge across the supply chain. The taxonomy is useful not because it enumerates all possible knowledge practitioners might need, but precisely because it provides a shared *structure* for articulating, interpreting, uncovering, and coordinating knowledge needs in a setting where such needs are often partial, fragmented, and contested. This aligns with broader HCI literature showing that descriptive and generative frameworks are essential precursors to effective contextualization and action [54].

In practice, this easy-to-understand structure with multiple entry points (different practitioners started familiarizing with the taxonomy using different dimensions) supports different forms of sense-making across roles without requiring uniform expertise. It helps actors recognize when a knowledge request is underspecified, situate their own expertise relative to others, and reason about what kinds of knowledge are missing or misaligned. It also constitutes an ideation scaffold and a documentation backbone for individual use and for the further development of governance strategies. Besides, by stabilizing categories of knowledge without prescribing content, it fosters the necessary negotiation work that our interview participants highlighted with the subjectivity of knowledge needs. In this sense, the taxonomy functions as a boundary object for the LLM supply chain: it does not resolve disagreement about responsibility or decision-making nor does it solve knowledge gaps, but it renders those disagreements and blindspots legible and actionable. This helps explain why a taxonomy (rather than a checklist or role-specific framework) is well suited to supporting responsible AI work in supply chain settings.

6 Implications

6.1 Enriching our understanding of knowledge needs

Our study showed that practitioners are not always aware of their own knowledge needs, that these needs are subjective, and that miscommunications often arise. This has implications for future research on responsible AI and knowledge. Researchers should not rely solely on self-reported knowledge needs, but also on indirect and relational methods for eliciting knowledge needs (in light of existing knowledge elicitation methods [35, 36, 66, 78, 83]), potentially making use of our taxonomy as an elicitation tool but also as a terminological disambiguation means. Besides, they should treat knowledge needs as situated and negotiated, rather than as static checklists to be satisfied. With these ideas in mind, we encourage future research to investigate the following research directions.

6.1.1 Extending the set of knowledge nuggets and sites of the supply chain. Since our study constitutes one of the first attempts at comprehensively studying knowledge needs across the LLM supply chain, we had to limit its scope. We encourage researchers to broaden it along multiple dimensions. We have not explored in-depth all types of practitioners and organizations within the LLM supply chain (the LLM providers were all large organizations and all the feedback sessions were conducted

at the same LLM provider). We narrowed knowledge nuggets to those having to do with responsible AI, but each actor needs and possesses additional nuggets that are specific to their own work, cf. Xie et al. [158]. We developed the taxonomy around LLM supply chains, and additional research is necessary to corroborate its applicability to other types of AI systems.

6.1.2 Verifying which nuggets are indeed needed. Whether a knowledge nugget in the taxonomy is actually needed for an actor of the chain merits further exploration. Although this will prove challenging due to the subjective nature of needs, this appeared especially urgent during the feedback sessions where some participants tended to develop many new needs without being able to articulate the reason therefor. E.g., a model owner (P13) said “*it forces me to think. [...] I would be happy to see what was tried out [in terms of experiments on the LLM training parameters]. I don’t know why exactly but I would like to see it.*” Researchers might want to prioritize the still under-explored but seemingly-essential categories of knowledge our study revealed (Section 5).

6.2 Developing tools for knowledge exchanges

Tools revealed to be one of the promising avenues for knowledge acquisition and exchanges. However, they were not widespread across all the organizations we covered, and they remained the objects of some interrogations (e.g., how to ensure that the documentation is continuously up-to-date). Besides, most participants were not aware of research progress along these lines, such as simple or interactive AI documentation [39, 57, 105], interfaces [40, 153], or training programs [136, 138]. Hence, we encourage researchers to investigate to what extent and how existing research could be adapted to the context of LLM supply chains to facilitate knowledge exchanges. Particularly, our knowledge taxonomy is a direct answer to some of the needs we identified, and could be extended in multiple ways.

6.2.1 Ensuring appropriate use of the nuggets. Our findings revealed that possessing a knowledge nugget alone does not represent a necessary nor a sufficient condition to achieve responsible LLM systems. For instance, practitioners often need knowledge nuggets from multiple categories to achieve one goal. Particularly many examples in Section 3 point to combining generalized knowledge which might refer to owning competencies of the generalized production category, with factual information in the specific knowledge categories, and comprehending scientific and policy principles in the other generalized categories. Hence, future extensions of the taxonomy should investigate how to support LLM practitioners in leveraging knowledge appropriately. A tool could for instance suggest nuggets that are often combined together. For that, one could include metadata such as whether a need is self-reported, externally perceived and validated, and which practitioners it pertains to and what goals they have [141]. We caution however that the taxonomy should not become a prescriptive tool due to the subjectivity of needs. Instead the design should allow for personalization.

6.2.2 Improving usability. During the feedback sessions, the participants did not express being overwhelmed by the taxonomy, however it took them time to explore the lists of knowledge nuggets and identify relevant ones. Hence, the usability of the taxonomy could be refined. Aligned with prior findings around AI explainability [50], we recommend investigating how to build usable tools on top of the taxonomy that are tailored to their audience. Different types of practitioners might require accessing different knowledge categories at different levels of granularity (e.g., those with expertise in technical knowledge might want to be presented with more granular technical knowledge but high-level organizational one). Based on our findings, the intention of the knowledge seeker might also call for diverse designs, e.g., serendipitous knowledge discovery could be facilitated by a simple overview of the 32 knowledge categories that make up the taxonomy or by an easy

toggling between the different dimensions of the taxonomy, while targeted knowledge search could potentially benefit from an interactive search function among all the included knowledge nuggets. More broadly, whether the concrete knowledge exchange should happen through the tool, or whether it should simply be facilitated by the tool remains to be investigated. Interrogating which of the 32 knowledge categories are adapted to tool-based exchanges would be particularly insightful. Tacit knowledge contained in the generalized production category might for instance be challenging to formalize without deep conversations [78], and the tool could instead focus on connecting knowledge seekers and holders together (e.g., by having metadata specifying who is a knowledge holder in an organization for each nugget).

6.3 Improving knowledge exchange processes

An important implication of our study is that responsible AI practices do not only require more tools, but also call for institutional forms of knowledge exchanges, that merit further research.

6.3.1 Crafting exchange spaces. While organizational processes and shared spaces revealed essential, we do not know how to best operationalize and facilitate these modes of exchanges. How to ensure that the processes allow for both serendipitous and targeted knowledge exchanges, driven by either the needer or their interlocutor? Note especially that serendipitous and interlocutor-driven exchanges lack supporting artifacts until now. An idea could be to craft task-based exercises where multiple actors would give feedback to each other. These exercises might however be met with typical challenges associated with feedback giving (e.g., the vulnerability it requires to expose one's own reasoning to others for scrutiny and criticism [134]). Should we attempt at streamlining *chained* exchanges across multiple actors of the chain (e.g., if they happen simply due to a lack of organizational visibility), or is there a reason for them to be chained (e.g., if it is difficult for actors who are distant in the chain to communicate)? To address accountability needs, is it sufficient to extend knowledge intermediaries and create knowledge advocate roles?

6.3.2 Creating the conditions for knowledge exchanges. Unclear responsibility for producing and maintaining knowledge repeatedly prevented knowledge from being shared or kept up to date. This suggests that artifacts and workflows alone are insufficient: effective knowledge exchange also depends on organizational arrangements that assign ownership, legitimize time spent sharing knowledge, and align incentives across roles. Finally, incentivizing knowledge exchanges can also be facilitated via policies, especially across organizations (e.g., the AI Act [99] only starts specifying exchanges between LLM providers and deployers). In that regard, we invite policymakers to expand and revise their understanding of knowledge across the LLM supply chain, since we found that knowledge needs were not necessarily reflected in the obligations they formulated.

7 Conclusions

This paper proposed the notion of knowledge as a precondition for responsible LLM systems, and examined knowledge as a central coordination and exchange problem in responsible LLM supply chains. We empirically identified the primary knowledge gaps faced by actors within the LLM supply chain, namely a blindspot recognition gap and an access / translation gap, and demonstrated that the supply chain offers a unique opportunity to address these gaps. Furthermore, we examined how knowledge exchanges occur across the supply chain, which strategies facilitate them (e.g., knowledge intermediaries), and which barriers impede them (e.g., lack of organizational transparency and accountability). Finally, we developed a comprehensive multi-dimensional LLM knowledge taxonomy that summarizes the concrete knowledge needs of the actors. Feedback sessions indicate that the taxonomy can function as a practical boundary object: practitioners used it to disambiguate vague knowledge requests, reflect on previously unprioritized knowledge categories,

identify others' blind spots, and ideate about governance interventions. This can ultimately support responsible AI work in the supply chain, by enabling more effective and better-informed dialogues and reliance decisions. Combined with our qualitative findings, researchers and policymakers can leverage our taxonomy to better account for the role of LLM-related knowledge in shaping responsible practices.

References

- [1] 2022. Microsoft Responsible AI Standard, v2. <https://query.prod.cms.rt.microsoft.com/cms/api/am/binary/RE5cmFl?culture=en-us&country=us>
- [2] Mark S Ackerman. 2000. The intellectual challenge of CSCW: the gap between social requirements and technical feasibility. *Human-Computer Interaction* 15, 2-3 (2000), 179–203.
- [3] Mark S Ackerman, Juri Dachtera, Volkmar Pipek, and Volker Wulf. 2013. Sharing knowledge and expertise: The CSCW view of knowledge management. *Computer Supported Cooperative Work (CSCW)* 22 (2013), 531–573.
- [4] David Adkins, Bilal Alsallakh, Adeel Cheema, Narine Kokhlikyan, Emily McReynolds, Pushkar Mishra, Chavez Procope, Jeremy Sawruk, Erin Wang, and Polina Zvyagina. 2022. Prescriptive and descriptive approaches to machine-learning transparency. In *CHI conference on human factors in computing systems extended abstracts*. 1–9.
- [5] Leah Hope Ajmani, Jasmine C Foriest, Jordan Taylor, Kyle Pittman, Sarah Gilbert, and Michael Ann Devito. 2024. Whose Knowledge is Valued? Epistemic Injustice in CSCW Applications. *Proceedings of the ACM on Human-Computer Interaction* 8, CSCW2 (2024), 1–28.
- [6] Saleema Amershi, Andrew Begel, Christian Bird, Robert DeLine, Harald Gall, Ece Kamar, Nachiappan Nagappan, Besmira Nushi, and Thomas Zimmermann. 2019. Software engineering for machine learning: A case study. In *2019 IEEE/ACM 41st International Conference on Software Engineering: Software Engineering in Practice (ICSE-SEIP)*. IEEE, 291–300.
- [7] Gloria Andrada, Robert W Clowes, and Paul R Smart. 2023. Varieties of transparency: Exploring agency within AI systems. *AI & society* 38, 4 (2023), 1321–1331.
- [8] Ariful Islam Anik and Andrea Bunt. 2021. Data-centric explanations: explaining training data of machine learning systems to promote transparency. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*. 1–13.
- [9] Dimitris Apostolou, Gregoris Mentzas, Bertin Klein, Andreas Abecker, and Wolfgang Maass. 2008. Interorganizational knowledge exchanges. *IEEE Intelligent Systems* 23, 4 (2008), 65–74.
- [10] Agathe Balayn, Natasa Rikalo, Christoph Lofi, Jie Yang, and Alessandro Bozzon. 2022. How can explainability methods be used to support bug identification in computer vision models?. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*. 1–16.
- [11] Agathe Balayn, Panagiotis Soilis, Christoph Lofi, Jie Yang, and Alessandro Bozzon. 2021. What do you mean? Interpreting image classification with crowdsourced concept extraction and analysis. In *Proceedings of the Web Conference 2021*. 1937–1948.
- [12] Agathe Balayn, Mireia Yurrita, Fanny Rancourt, Fabio Casati, and Ujwal Gadiraju. 2025. Unpacking Trust Dynamics in the LLM Supply Chain: An Empirical Exploration to Foster Trustworthy LLM Production & Use. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. 1–20.
- [13] Agathe Balayn, Mireia Yurrita, Jie Yang, and Ujwal Gadiraju. 2023. “Fairness Toolkits, A Checkbox Culture?” On the Factors that Fragment Developer Practices in Handling Algorithmic Harms. In *Proceedings of the 2023 AAAI/ACM Conference on AI, Ethics, and Society*. 482–495.
- [14] Gagan Bansal, Besmira Nushi, Ece Kamar, Walter S Lasecki, Daniel S Weld, and Eric Horvitz. 2019. Beyond accuracy: The role of mental models in human-AI team performance. In *Proceedings of the AAAI conference on human computation and crowdsourcing*, Vol. 7. 2–11.
- [15] Vita Santa Barletta, Danilo Caivano, Domenico Gigante, and Azzurra Ragone. 2023. A rapid review of responsible ai frameworks: How to guide the development of ethical ai. In *Proceedings of the 27th International Conference on Evaluation and Assessment in Software Engineering*. 358–367.
- [16] Rachel KE Bellamy, Kuntal Dey, Michael Hind, Samuel C Hoffman, Stephanie Houde, Kalapriya Kannan, Pranay Lohia, Jacquelyn Martino, Sameep Mehta, Aleksandra Mojsilović, et al. 2019. AI Fairness 360: An extensible toolkit for detecting and mitigating algorithmic bias. *IBM Journal of Research and Development* 63, 4/5 (2019), 4–1.
- [17] Emily M Bender, Timnit Gebu, Angelina McMillan-Major, and Shmargaret Shmitchell. 2021. On the dangers of stochastic parrots: Can language models be too big?. In *Proceedings of the 2021 ACM conference on fairness, accountability, and transparency*. 610–623.
- [18] Avinash Bhat, Austin Coursey, Grace Hu, Sixian Li, Nadia Nahar, Shurui Zhou, Christian Kästner, and Jin LC Guo. 2023. Aspirations and practice of ml model documentation: Moving the needle with nudging and traceability. In

Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems. 1–17.

- [19] Maalvika Bhat and Duri Long. 2024. Designing Interactive Explainable AI Tools for Algorithmic Literacy and Transparency. In *Proceedings of the 2024 ACM Designing Interactive Systems Conference*. 939–957.
- [20] Sarah Bird, Miro Dudik, Richard Edgar, Brandon Horn, Roman Lutz, Vanessa Milan, Mehrnoosh Sameki, Hanna Wallach, and Kathleen Walker. 2020. Fairlearn: A toolkit for assessing and improving fairness in AI. *Microsoft, Tech. Rep. MSR-TR-2020-32* (2020).
- [21] Abeba Birhane, William Isaac, Vinodkumar Prabhakaran, Mark Diaz, Madeleine Clare Elish, Iason Gabriel, and Shakir Mohamed. 2022. Power to the people? Opportunities and challenges for participatory AI. In *Proceedings of the 2nd ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization*. 1–8.
- [22] Roger E Bohn. 2009. Measuring and managing technological knowledge. In *The Economic impact of knowledge*. Routledge, 295–314.
- [23] Rishi Bommasani, Kevin Klyman, Shayne Longpre, Sayash Kapoor, Nestor Maslej, Betty Xiong, Daniel Zhang, and Percy Liang. 2023. The foundation model transparency index. *arXiv preprint arXiv:2310.12941* (2023).
- [24] Karen L Boyd. 2021. Datasheets for datasets help ML engineers notice and understand ethical issues in training data. *Proceedings of the ACM on Human-Computer Interaction* 5, CSCW2 (2021), 1–27.
- [25] John Seely Brown and Paul Duguid. 1991. Organizational learning and communities-of-practice: Toward a unified view of working, learning, and innovation. *Organization science* 2, 1 (1991), 40–57.
- [26] Katriina Byström. 1999. *Task complexity, information types and information sources: examination of relationships*. Tampere University Press.
- [27] Federico Cabitza, Andrea Campagner, and Carla Simone. 2021. The need to move away from agential-AI: Empirical investigations, useful concepts and open issues. *International Journal of Human-Computer Studies* 155 (2021), 102696.
- [28] Carrie J Cai, Samantha Winter, David Steiner, Lauren Wilcox, and Michael Terry. 2019. "Hello AI": uncovering the onboarding needs of medical practitioners for human-AI collaborative decision-making. *Proceedings of the ACM on Human-computer Interaction* 3, CSCW (2019), 1–24.
- [29] Alexandre Canny, Elodie Bouzekri, Célia Martinie, and Philippe Palanque. 2019. On the Importance of Supporting Multiple Stakeholders Points of View for the Testing of Interactive Systems. In *1st Workshop on Research and Practice Challenges for Engineering Interactive Systems while Integrating Multiple Stakeholders Viewpoints (EISM 2019)*, Vol. 2503. CEUR-WS: Workshop proceedings, 113–121.
- [30] Paul R Carlile. 2002. A pragmatic view of knowledge and boundaries: Boundary objects in new product development. *Organization science* 13, 4 (2002), 442–455.
- [31] Lu Cheng, Kush R Varshney, and Huan Liu. 2021. Socially responsible ai algorithms: Issues, purposes, and challenges. *Journal of Artificial Intelligence Research* 71 (2021), 1137–1181.
- [32] Thomas KF Chiu. 2021. A holistic approach to the design of artificial intelligence (AI) education for K-12 schools. *TechTrends* 65, 5 (2021), 796–807.
- [33] Jennifer Cobbe, Michael Veale, and Jatinder Singh. 2023. Understanding accountability in algorithmic supply chains. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*. 1186–1197.
- [34] Jeff Conklin and Michael L Begeman. 1988. gIBIS: A hypertext tool for exploratory policy discussion. *ACM Transactions on Information Systems (TOIS)* 6, 4 (1988), 303–331.
- [35] Nancy J Cooke. 1994. Varieties of knowledge elicitation techniques. *International journal of human-computer studies* 41, 6 (1994), 801–849.
- [36] Nancy J Cooke. 1999. Knowledge elicitation. *Handbook of applied cognition* (1999), 479–509.
- [37] Eric Corbett and Emily Denton. 2023. Interrogating the T in FAccT. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*. 1624–1634.
- [38] Kate Crawford and Vladan Joler. 2018. Anatomy of an AI System. *Anatomy of an AI System* (2018).
- [39] Anamaria Crisan, Margaret Drouhard, Jesse Vig, and Nazneen Rajani. 2022. Interactive model cards: A human-centered approach to model documentation. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*. 427–439.
- [40] Roxana Daneshjou, Mary P Smith, Mary D Sun, Veronica Rotemberg, and James Zou. 2021. Lack of transparency and potential bias in artificial intelligence data sets and algorithms: a scoping review. *JAMA dermatology* 157, 11 (2021), 1362–1369.
- [41] Ton De Jong and Monica GM Ferguson-Hessler. 1996. Types and qualities of knowledge. *Educational psychologist* 31, 2 (1996), 105–113.
- [42] Marc J De Vries. 2003. The nature of technological knowledge: Extending empirically informed studies into what engineers know. *Techne: Research in philosophy and technology* 6, 3 (2003), 117–130.
- [43] Fernando Delgado, Stephen Yang, Michael Madaio, and Qian Yang. 2023. The participatory turn in ai design: Theoretical foundations and the current state of practice. In *Proceedings of the 3rd ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization*. 1–23.

- [44] Dominik Dellermann, Adrian Calma, Nikolaus Lipusch, Thorsten Weber, Sascha Weigel, and Philipp Ebel. 2021. The future of human-AI collaboration: a taxonomy of design knowledge for hybrid intelligence systems. *arXiv preprint arXiv:2105.03354* (2021).
- [45] Wesley Hanwen Deng, Glen Berman, Harmanpreet Kaur, Reva Schwartz, Memunat A Ibrahim, Motahhare Eslami, Kenneth Holstein, and Michael Madaio. 2024. Collaboratively Designing and Evaluating Responsible AI Interventions. In *Companion Publication of the 2024 Conference on Computer-Supported Cooperative Work and Social Computing*. 658–662.
- [46] Wesley Hanwen Deng, Manish Nagireddy, Michelle Seng Ah Lee, Jatinder Singh, Zhiwei Steven Wu, Kenneth Holstein, and Haiyi Zhu. 2022. Exploring how machine learning practitioners (try to) use fairness toolkits. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*. 473–484.
- [47] Advait Deshpande and Helen Sharp. 2022. Responsible ai systems: who are the stakeholders?. In *Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society*. 227–236.
- [48] Kevin C Desouza and Yukika Awazu. 2006. Knowledge management at SMEs: five peculiarities. *Journal of knowledge management* 10, 1 (2006), 32–43.
- [49] Shipi Dhanorkar, Christine T Wolf, Kun Qian, Anbang Xu, Lucian Popa, and Yunyao Li. 2021. Who needs to know what, when?: Broadening the Explainable AI (XAI) Design Space by Looking at Explanations Across the AI Lifecycle. In *Proceedings of the 2021 ACM Designing Interactive Systems Conference*. 1591–1602.
- [50] Finale Doshi-Velez and Been Kim. 2017. Towards a rigorous science of interpretable machine learning. *arXiv preprint arXiv:1702.08608* (2017).
- [51] Lilian Edwards. 2022. Regulating AI in Europe: four problems and four solutions. *Ada Lovelace Institute* 15 (2022), 2022.
- [52] Upol Ehsan, Q Vera Liao, Michael Muller, Mark O Riedl, and Justin D Weisz. 2021. Expanding explainability: Towards social transparency in ai systems. In *Proceedings of the 2021 CHI conference on human factors in computing systems*. 1–19.
- [53] Upol Ehsan, Koustuv Saha, Munmun De Choudhury, and Mark O Riedl. 2023. Charting the sociotechnical gap in explainable ai: A framework to address the gap in XAI. *Proceedings of the ACM on human-computer interaction* 7, CSCW1 (2023), 1–32.
- [54] Shitao Fang, Koji Yatani, and Kasper Hornbæk. 2026. What We Talk About When We Talk About Frameworks in HCI. *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems* (2026).
- [55] Joan Fazey, Lukas Bunse, Joshua Msika, Maria Pinke, Katherine Preedy, Anna C Evely, Emily Lambert, Emily Hastings, Sue Morris, and Mark S Reed. 2014. Evaluating knowledge exchange in interdisciplinary and multi-stakeholder research. *Global Environmental Change* 25 (2014), 204–220.
- [56] Batya Friedman, Peter H Kahn, Alan Borning, and Alina Hultgren. 2013. Value sensitive design and information systems. *Early engagement and new technologies: Opening up the laboratory* (2013), 55–95.
- [57] Timnit Gebru, Jamie Morgenstern, Briana Vecchione, Jennifer Wortman Vaughan, Hanna Wallach, Hal Daumé Iii, and Kate Crawford. 2021. Datasheets for datasets. *Commun. ACM* 64, 12 (2021), 86–92.
- [58] Michael E Gorman. 2002. Types of knowledge and their roles in technology transfer. *The Journal of Technology Transfer* 27, 3 (2002), 219–231.
- [59] Robert M Grant. 1996. Toward a knowledge-based theory of the firm. *Strategic management journal* 17, S2 (1996), 109–122.
- [60] Kristina Groth and John Bowers. 2001. On finding things out: Situating organisational knowledge in CSCW. In *ECSCW 2001: Proceedings of the Seventh European Conference on Computer Supported Cooperative Work 16–20 September 2001, Bonn, Germany*. Springer, 279–298.
- [61] Sven Ove Hansson. 2013. What is technological knowledge? In *Technology teachers as researchers*. Brill, 15–31.
- [62] Amy K Heger, Liz B Marquis, Mihaela Vorvoreanu, Hanna Wallach, and Jennifer Wortman Vaughan. 2022. Understanding machine learning practitioners’ data documentation perceptions, needs, challenges, and desiderata. *Proceedings of the ACM on Human-Computer Interaction* 6, CSCW2 (2022), 1–29.
- [63] Dennis R Herschbach et al. 1995. Technology as knowledge: Implications for instruction. *Volume 7 Issue 1 (fall 1995)* (1995).
- [64] Merve Hickok. 2021. Lessons learned from AI ethics principles for future actions. *AI and Ethics* 1, 1 (2021), 41–47.
- [65] Martin Hitziger, Maurizio Aragrande, John A Berezowski, Massimo Canali, Victor Del Rio Vilas, Sabine Hoffmann, Gilberto Igrejas, Hans Keune, Alexandra Lux, Mieghan Bruce, et al. 2019. EVOLvINC: evaluating knowledge integration capacity in multistakeholder governance. *Ecology and Society* 24, 2 (2019), art36.
- [66] Robert R Hoffman, Nigel R Shadbolt, A Mike Burton, and Gary Klein. 1995. Eliciting knowledge from experts: A methodological analysis. *Organizational behavior and human decision processes* 62, 2 (1995), 129–158.
- [67] James Hollan, Edwin Hutchins, and David Kirsh. 2000. Distributed cognition: toward a new foundation for human-computer interaction research. *ACM Transactions on Computer-Human Interaction (TOCHI)* 7, 2 (2000), 174–196.

- [68] Sungsoo Ray Hong, Jessica Hullman, and Enrico Bertini. 2020. Human factors in model interpretability: Industry practices, challenges, and needs. *Proceedings of the ACM on Human-Computer Interaction* 4, CSCW1 (2020), 1–26.
- [69] Edwin Hutchins. 1995. *Cognition in the Wild*. MIT press.
- [70] Maurice Jakesch, Zana Bućinca, Saleema Amershi, and Alexandra Olteanu. 2022. How different groups prioritize ethical values for responsible AI. In *proceedings of the 2022 ACM conference on fairness, accountability, and transparency*. 310–323.
- [71] Anna Jobin, Marcello Ienca, and Effy Vayena. 2019. The global landscape of AI ethics guidelines. *Nature machine intelligence* 1, 9 (2019), 389–399.
- [72] Michael I Jordan and Tom M Mitchell. 2015. Machine learning: Trends, perspectives, and prospects. *Science* 349, 6245 (2015), 255–260.
- [73] Emma Kallina and Jatinder Singh. 2024. Stakeholder Involvement for Responsible AI Development: A Process Framework. In *Proceedings of the 4th ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization*. 1–14.
- [74] Holtzblatt Karen and Jones Sandra. 2017. Contextual inquiry: A participatory technique for system design. In *Participatory design*. CRC Press, 177–210.
- [75] Harmanpreet Kaur, Harsha Nori, Samuel Jenkins, Rich Caruana, Hanna Wallach, and Jennifer Wortman Vaughan. 2020. Interpreting interpretability: understanding data scientists’ use of interpretability tools for machine learning. In *Proceedings of the 2020 CHI conference on human factors in computing systems*. 1–14.
- [76] Anna Kawakami, Amanda Coston, Haiyi Zhu, Hoda Heidari, and Kenneth Holstein. 2024. The Situate AI Guidebook: Co-Designing a Toolkit to Support Multi-Stakeholder, Early-stage Deliberations Around Public Sector AI Proposals. In *Proceedings of the CHI Conference on Human Factors in Computing Systems*. 1–22.
- [77] Samuel Kernan Freire. 2023. The Human Factors of AI-Empowered Knowledge Sharing. In *Extended Abstracts of the 2023 CHI Conference on Human Factors in Computing Systems*. 1–5.
- [78] Samuel Kernan Freire, Chaofan Wang, Santiago Ruiz-Arenas, and Evangelos Niforatos. 2023. Tacit knowledge elicitation for shop-floor workers with an intelligent assistant. In *Extended Abstracts of the 2023 CHI Conference on Human Factors in Computing Systems*. 1–7.
- [79] Hassan Khosravi, Simon Buckingham Shum, Guanliang Chen, Cristina Conati, Yi-Shan Tsai, Judy Kay, Simon Knight, Roberto Martinez-Maldonado, Shazia Sadiq, and Dragan Gašević. 2022. Explainable artificial intelligence in education. *Computers and Education: Artificial Intelligence* 3 (2022), 100074.
- [80] John Kidd. 2005. Knowledge management tools and techniques: practitioners and experts evaluate KM solutions. *Knowledge Management Research & Practice* 3, 2 (2005), 117–119.
- [81] Sunnie SY Kim, Elizabeth Anne Watkins, Olga Russakovsky, Ruth Fong, and Andrés Monroy-Hernández. 2023. "Help Me Help the AI": Understanding How Explainability Can Support Human-AI Interaction. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. 1–17.
- [82] Nils Knoth, Antonia Tolzin, Andreas Janson, and Jan Marco Leimeister. 2024. AI literacy and its implications for prompt engineering strategies. *Computers and Education: Artificial Intelligence* 6 (2024), 100225.
- [83] Rebecca Krosnick, Fraser Anderson, Justin Matejka, Steve Oney, Walter S. Lasecki, Tovi Grossman, and George Fitzmaurice. 2021. Think-aloud computing: Supporting rich and low-effort knowledge capture. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*. 1–13.
- [84] Patrick Lambe. 2014. *Organising knowledge: taxonomies, knowledge and organisational effectiveness*. Elsevier.
- [85] Michelle Seng Ah Lee and Jat Singh. 2021. The landscape and gaps in open source fairness toolkits. In *Proceedings of the 2021 CHI conference on human factors in computing systems*. 1–13.
- [86] Seongmin Lee, Zijie J Wang, Aishwarya Chakravarthy, Alec Helbling, ShengYun Peng, Mansi Phute, Duen Horng Chau, and Minsuk Kahng. 2024. LLM Attributor: Interactive Visual Attribution for LLM Generation. *arXiv preprint arXiv:2404.01361* (2024).
- [87] Claire Liang, Julia Proft, Erik Andersen, and Ross A Knepper. 2019. Implicit communication of actionable information in human-ai teams. In *Proceedings of the 2019 CHI conference on human factors in computing systems*. 1–13.
- [88] Q Vera Liao, Daniel Gruen, and Sarah Miller. 2020. Questioning the AI: informing design practices for explainable AI user experiences. In *Proceedings of the 2020 CHI conference on human factors in computing systems*. 1–15.
- [89] Q Vera Liao, Hariharan Subramonyam, Jennifer Wang, and Jennifer Wortman Vaughan. 2023. Designerly understanding: Information needs for model transparency to support design ideation for AI-powered user experience. In *Proceedings of the 2023 CHI conference on human factors in computing systems*. 1–21.
- [90] Q Vera Liao and Jennifer Wortman Vaughan. [n. d.]. AI transparency in the age of llms: A human-centered research roadmap. ([n. d.]).
- [91] Q Vera Liao and Jennifer Wortman Vaughan. 2024. AI Transparency in the Age of LLMs: A Human-Centered Research Roadmap. *Harvard Data Science Review Special Issue* 5 (2024).

- [92] Zhengzhong Liu, Aurick Qiao, Willie Neiswanger, Hongyi Wang, Bowen Tan, Tianhua Tao, Junbo Li, Yuqi Wang, Suqi Sun, Omkar Pangarkar, et al. [n. d.]. LLM360: Towards Fully Transparent Open-Source LLMs. In *First Conference on Language Modeling*.
- [93] Duri Long and Brian Magerko. 2020. What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI conference on human factors in computing systems*. 1–16.
- [94] Ana Lucic, Madhulika Srikumar, Umang Bhatt, Alice Xiang, Ankur Taly, Q Vera Liao, and Maarten de Rijke. 2021. A multistakeholder approach towards evaluating AI transparency mechanisms. *arXiv preprint arXiv:2103.14976* (2021).
- [95] Bengt-Ake Lundvall. 2016. From the economics of knowledge to the learning economy. *The learning economy and the economics of hope* 133 (2016).
- [96] Marco Lünich and Birte Keller. 2024. Explainable Artificial Intelligence for Academic Performance Prediction. An Experimental Study on the Impact of Accuracy and Simplicity of Decision Trees on Causability and Fairness Perceptions. In *The 2024 ACM Conference on Fairness, Accountability, and Transparency*. 1031–1042.
- [97] Michael Madaio, Lisa Egede, Hariharan Subramonyam, Jennifer Wortman Vaughan, and Hanna Wallach. 2022. Assessing the fairness of ai systems: Ai practitioners' processes, challenges, and needs for support. *Proceedings of the ACM on Human-Computer Interaction* 6, CSCW1 (2022), 1–26.
- [98] Michael Madaio, Shivani Kapania, Rida Qadri, Ding Wang, Andrew Zaldivar, Remi Denton, and Lauren Wilcox. 2024. Learning about responsible AI on-the-job: Learning pathways, orientations, and aspirations. In *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency*. 1544–1558.
- [99] Tambiana Madiega. 2021. Artificial intelligence act. *European Parliament: European Parliamentary Research Service* (2021).
- [100] Yao Li Mao, Dakuo Wang, Michael Muller, Kush R Varshney, Ioana Baldini, Casey Dugan, and Aleksandra Mojsilović. 2019. How data scientists work together with domain experts in scientific collaborations: To find the right answer or to ask the right question? *Proceedings of the ACM on Human-Computer Interaction* 3, GROUP (2019), 1–23.
- [101] David W McDonald and Mark S Ackerman. 1998. Just talk to me: a field study of expertise location. In *Proceedings of the 1998 ACM conference on Computer supported cooperative work*. 315–324.
- [102] Marina Micheli, Isabelle Hupont, Blagoj Delipetrev, and Josep Soler-Garrido. 2023. The landscape of data and AI documentation approaches in the European policy context. *Ethics and Information Technology* 25, 4 (2023), 56.
- [103] Gloria J Miller. 2022. Stakeholder roles in artificial intelligence projects. *Project Leadership and Society* 3 (2022), 100068.
- [104] Marvin Minsky et al. 1974. A framework for representing knowledge.
- [105] Margaret Mitchell, Simone Wu, Andrew Zaldivar, Parker Barnes, Lucy Vasserman, Ben Hutchinson, Elena Spitzer, Inioluwa Deborah Raji, and Timnit Gebru. 2019. Model cards for model reporting. In *Proceedings of the conference on fairness, accountability, and transparency*. 220–229.
- [106] Brent Mittelstadt. 2019. Principles alone cannot guarantee ethical AI. *Nature machine intelligence* 1, 11 (2019), 501–507.
- [107] Michael J Muller and Sarah Kuhn. 1993. Participatory design. *Commun. ACM* 36, 6 (1993), 24–28.
- [108] Meena Devii Muralikumar and David W McDonald. 2024. Analyzing Collaborative Challenges and Needs of UX Practitioners when Designing with AI/ML. *Proceedings of the ACM on Human-Computer Interaction* 8, CSCW2 (2024), 1–25.
- [109] Davy Tsz Kit Ng, Jac Ka Lok Leung, Samuel Kai Wah Chu, and Maggie Shen Qiao. 2021. Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence* 2 (2021), 100041.
- [110] Donald A Norman. 2014. Some observations on mental models. In *Mental models*. Psychology Press, 7–14.
- [111] Chinasa T Okolo and Hongjin Lin. 2024. "You can't build what you don't understand": Practitioner Perspectives on Explainable AI in the Global South. In *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems*. 1–10.
- [112] Mohamad Amin Osman, Shahrul Azman Mohd Noah, and Saidah Saad. 2022. Ontology-based knowledge management tools for knowledge sharing in organization—a review. *IEEE access* 10 (2022), 43267–43283.
- [113] Soya Park, April Yi Wang, Ban Kwas, Q Vera Liao, David Piorkowski, and Marina Danilevsky. 2021. Facilitating knowledge sharing from domain experts to data scientists for building nlp models. In *Proceedings of the 26th International Conference on Intelligent User Interfaces*. 585–596.
- [114] Jaana Parviainen and Marja Eriksson. 2006. Negative knowledge, expertise and organisations. *International Journal of Management Concepts and Philosophy* 2, 2 (2006), 140–153.
- [115] Samir Passi and Phoebe Sengers. 2020. Making data science systems work. *Big data & society* 7, 2 (2020), 2053951720939605.
- [116] Marc Pinski, Martin Adam, and Alexander Benlian. 2023. AI knowledge: Improving AI delegation through human enablement. In *Proceedings of the 2023 CHI conference on human factors in computing systems*. 1–17.

- [117] David Piorkowski, Soya Park, April Yi Wang, Dakuo Wang, Michael Muller, and Felix Portnoy. 2021. How ai developers overcome communication challenges in a multidisciplinary team: A case study. *Proceedings of the ACM on Human-Computer Interaction* 5, CSCW1 (2021), 1–25.
- [118] Volkmar Pipek and Volker Wulf. 2009. Infrastructuring: Toward an integrated perspective on the design and use of information technology. *Journal of the Association for Information Systems* 10, 5 (2009), 1.
- [119] Laurence Prusak. 2009. *Knowledge in organisations*. Routledge.
- [120] Inioluwa Deborah Raji, Andrew Smart, Rebecca N White, Margaret Mitchell, Timnit Gebru, Ben Hutchinson, Jamila Smith-Loud, Daniel Theron, and Parker Barnes. 2020. Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. In *Proceedings of the 2020 conference on fairness, accountability, and transparency*. 33–44.
- [121] Bogdana Rakova, Jingying Yang, Henriette Cramer, and Rumman Chowdhury. 2021. Where responsible AI meets reality: Practitioner perspectives on enablers for shifting organizational practices. *Proceedings of the ACM on Human-Computer Interaction* 5, CSCW1 (2021), 1–23.
- [122] Shalaleh Rismani and AJung Moon. 2023. What does it mean to be a responsible AI practitioner: An ontology of roles and skills. In *Proceedings of the 2023 AAAI/ACM Conference on AI, Ethics, and Society*. 584–595.
- [123] Anna Rogers, Timothy Baldwin, and Kobi Leins. 2021. ‘Just What do You Think You’re Doing, Dave?’ A Checklist for Responsible Data Use in NLP. In *Findings of the Association for Computational Linguistics: EMNLP 2021*. 4821–4833.
- [124] Günter Ropohl. 1997. Knowledge types in technology. *International Journal of technology and design education* 7 (1997), 65–72.
- [125] Jennifer Rowley. 2007. The wisdom hierarchy: representations of the DIKW hierarchy. *Journal of information science* 33, 2 (2007), 163–180.
- [126] Gilbert Ryle. 1949. Descartes’ myth: The concept of mind.
- [127] Nithya Sambasivan and Rajesh Veeraraghavan. 2022. The deskilling of domain expertise in AI development. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*. 1–14.
- [128] Kjeld Schmidt. 2012. The trouble with ‘tacit knowledge’. *Computer supported cooperative work (CSCW)* 21 (2012), 163–225.
- [129] Kjeld Schmidt and Liam Bannon. 1992. Taking CSCW seriously: Supporting articulation work. *Computer supported cooperative work (CSCW)* 1, 1 (1992), 7–40.
- [130] Renee Shelby, Shalaleh Rismani, Kathryn Henne, AJung Moon, Negar Rostamzadeh, Paul Nicholas, N’Mah Yilla-Akbari, Jess Gallegos, Andrew Smart, Emilio Garcia, et al. 2023. Sociotechnical harms of algorithmic systems: Scoping a taxonomy for harm reduction. In *Proceedings of the 2023 AAAI/ACM Conference on AI, Ethics, and Society*. 723–741.
- [131] Wina Smeenk. 2023. The empathic Co-Design Canvas: A tool for supporting multi-stakeholder co-design processes. *International Journal of Design* 17, 2 (2023), 81–98.
- [132] Daniel Smilkov, Nikhil Thorat, Been Kim, Fernanda Viégas, and Martin Wattenberg. 2017. Smoothgrad: removing noise by adding noise. *arXiv preprint arXiv:1706.03825* (2017).
- [133] Kacper Sokol and Peter Flach. 2020. Explainability fact sheets: a framework for systematic assessment of explainable approaches. In *Proceedings of the 2020 conference on fairness, accountability, and transparency*. 56–67.
- [134] Jaemarie Solyst, Lauren Wilcox, and Michael Madaio. 2025. "The Conduit by which Change Happens": Processes, Barriers, and Support for Interpersonal Learning about Responsible AI. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. 1–15.
- [135] Madhulika Srikumar, Jiyoo Chang, and Kasia Chmielinski. 2024. Risk mitigation strategies for the open foundation model value chain. *Partnership on AI* 19 (2024).
- [136] Thilo Stadelmann, Julian Keuzenkamp, Helmut Grabner, and Christoph Würsch. 2021. The AI-atlas: didactics for teaching AI and machine learning on-site, online, and hybrid. *Education Sciences* 11, 7 (2021), 318.
- [137] Susan Leigh Star. 1989. The structure of ill-structured solutions: Boundary objects and heterogeneous distributed problem solving. In *Distributed artificial intelligence*. Elsevier, 37–54.
- [138] Peter Steinbach, Heidi Seibold, and Oliver Guhr. 2021. Teaching machine learning in 2020. In *European Conference on Machine Learning and Principles and Practice of Knowledge Discovery in Databases*. PMLR, 1–6.
- [139] Karin Stolpe and Jonas Hallström. 2024. Artificial intelligence literacy for technology education. *Computers and Education Open* 6 (2024), 100159.
- [140] Hariharan Subramonyam, Jane Im, Colleen Seifert, and Eytan Adar. 2022. Solving separation-of-concerns problems in collaborative design of human-AI systems through leaky abstractions. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*. 1–21.
- [141] Harini Suresh, Steven R Gomez, Kevin K Nam, and Arvind Satyanarayan. 2021. Beyond expertise and roles: A framework to characterize the stakeholders of interpretable machine learning and their needs. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*. 1–16.

- [142] Ningjing Tang, Megan Li, Amy A Winecoff, Michael Madaio, Hoda Heidari, and Hong Shen. 2025. Navigating Uncertainties: Understanding How GenAI Developers Document Their Models on Open-Source Platforms. *CoRR* (2025).
- [143] Jake B Telkamp and Marc H Anderson. 2022. The implications of diverse human moral foundations for assessing the ethicality of artificial intelligence. *Journal of Business Ethics* 178, 4 (2022), 961–976.
- [144] Gareth Terry, Nikki Hayfield, Victoria Clarke, Virginia Braun, et al. 2017. Thematic analysis. *The SAGE handbook of qualitative research in psychology* 2, 17-37 (2017), 25.
- [145] Violet Turri, Katelyn Morrison, Katherine-Marie Robinson, Collin Abidi, Adam Perer, Jodi Forlizzi, and Rachel Dzombak. 2024. Transparency in the Wild: Navigating Transparency in a Deployed AI System to Broaden Need-Finding Approaches. In *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency*. 1494–1514.
- [146] Fernando van der Vlist, Anne Helmond, and Fabian Ferrari. 2024. Big AI: Cloud infrastructure dependence and the industrialisation of artificial intelligence. *Big Data & Society* 11, 1 (2024), 20539517241232630.
- [147] Jerra Veeger and Michelle Westermann-Behaylo. 2022. Knowledge exchange within multi-stakeholder initiatives: tackling the Sustainable Development Goals. In *Research Handbook on Knowledge Transfer and International Business*. Edward Elgar Publishing, 108–135.
- [148] M Veldhuizen, V Blok, and D Dentoni. 2013. Organisational drivers of capabilities for multi-stakeholder dialogue and knowledge integration. *Journal on Chain and Network Science* 13, 2 (2013), 107–118.
- [149] James P Walsh and Gerardo Rivera Ungson. 2009. Organizational memory. In *Knowledge in Organisations*. Routledge, 177–212.
- [150] Jennifer Wang, Andrew D Selbst, Suresh Venkatasubramanian, and Solon Barocas. 2025. Distinguishing Predictive and Generative AI in Regulation. *UCLA School of Law, Public Law Research Paper Forthcoming* (2025).
- [151] Qiaosi Wang, Michael Madaio, Shaun Kane, Shivani Kapania, Michael Terry, and Lauren Wilcox. 2023. Designing responsible ai: Adaptations of ux practice to meet responsible ai challenges. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. 1–16.
- [152] Shenao Wang, Yanjie Zhao, Xinyi Hou, and Haoyu Wang. 2024. Large language model supply chain: A research agenda. *arXiv preprint arXiv:2404.12736* (2024).
- [153] James Wexler, Mahima Pushkarna, Tolga Bolukbasi, Martin Wattenberg, Fernanda Viégas, and Jimbo Wilson. 2019. The what-if tool: Interactive probing of machine learning models. *IEEE transactions on visualization and computer graphics* 26, 1 (2019), 56–65.
- [154] Jess Whittlestone, Rune Nyrup, Anna Alexandrova, and Stephen Cave. 2019. The role and limits of principles in AI ethics: Towards a focus on tensions. In *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*. 195–200.
- [155] David Gray Widder and Dawn Nafus. 2023. Dislocated accountabilities in the “AI supply chain”: Modularity and developers’ notions of responsibility. *Big Data & Society* 10, 1 (2023), 20539517231177620.
- [156] David Gray Widder, Sarah West, and Meredith Whittaker. 2023. Open (for business): Big tech, concentrated power, and the political economy of open AI. *Concentrated Power, and the Political Economy of Open AI (August 17, 2023)* (2023).
- [157] Maximiliane Windl, Sebastian S Feger, Lara Zijlstra, Albrecht Schmidt, and Pawel W Wozniak. 2022. ‘It Is Not Always Discovery Time’: Four Pragmatic Approaches in Designing AI Systems. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*. 1–12.
- [158] Shixian Xie, John Zimmerman, and Motahhare Eslami. 2025. Exploring What People Need to Know to be AI Literate: Tailoring for a Diversity of AI Roles and Responsibilities. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. 1–16.
- [159] Qian Yang, Aaron Steinfeld, Carolyn Rosé, and John Zimmerman. 2020. Re-examining whether, why, and how human-AI interaction is uniquely difficult to design. In *Proceedings of the 2020 chi conference on human factors in computing systems*. 1–13.
- [160] Mireia Yurrita, Dave Murray-Rust, Agathe Balayn, and Alessandro Bozzon. 2022. Towards a multi-stakeholder value-based assessment framework for algorithmic systems. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*. 535–563.
- [161] Douglas Zytke, Pamela J. Wisniewski, Shion Guha, Eric PS Baumer, and Min Kyung Lee. 2022. Participatory design of AI systems: opportunities and challenges across diverse users, relationships, and application domains. In *CHI Conference on Human Factors in Computing Systems Extended Abstracts*. 1–4.

Received May 13, 2025; revised January 13, 2026; accepted April 9, 2026