



# **“I Don’t Know If We’re Doing Good. I Don’t Know If We’re Doing Bad”: Investigating How Practitioners Scope, Motivate, and Conduct Privacy Work When Developing AI Products**

Hao-Ping (Hank) Lee, *Carnegie Mellon University*; Lan Gao and Stephanie Yang, *Georgia Institute of Technology*; Jodi Forlizzi and Sauvik Das, *Carnegie Mellon University*

<https://www.usenix.org/conference/usenixsecurity24/presentation/lee>

**This paper is included in the Proceedings of the 33rd USENIX Security Symposium.**

**August 14–16, 2024 • Philadelphia, PA, USA**

978-1-939133-44-1

**Open access to the Proceedings of the 33rd USENIX Security Symposium is sponsored by USENIX.**

# “I Don’t Know If We’re Doing Good. I Don’t Know If We’re Doing Bad”: Investigating How Practitioners Scope, Motivate, and Conduct Privacy Work When Developing AI Products

Hao-Ping (Hank) Lee<sup>1</sup>, Lan Gao<sup>2</sup>, Stephanie Yang<sup>2</sup>, Jodi Forlizzi<sup>1</sup>, Sauvik Das<sup>1</sup>  
<sup>1</sup>*Carnegie Mellon University*, <sup>2</sup>*Georgia Institute of Technology*

## Abstract

How do practitioners who develop consumer AI products scope, motivate, and conduct privacy work? Respecting privacy is a key principle for developing ethical, human-centered AI systems, but we cannot hope to better support practitioners without answers to that question. We interviewed 35 industry AI practitioners to bridge that gap. We found that practitioners viewed privacy as actions taken against pre-defined intrusions that can be exacerbated by the capabilities and requirements of AI, but few were aware of AI-specific privacy intrusions documented in prior literature. We found that their privacy work was rigidly defined and situated, guided by compliance with privacy regulations and policies, and generally demotivated beyond meeting minimum requirements. Finally, we found that the methods, tools, and resources they used in their privacy work generally did not help address the unique privacy risks introduced or exacerbated by their use of AI in their products. Collectively, these findings reveal the need and opportunity to create tools, resources, and support structures to improve practitioners’ awareness of AI-specific privacy risks, motivations to do AI privacy work, and ability to address privacy harms introduced or exacerbated by their use of AI in consumer products.

## 1 Introduction

Privacy is one of the five most commonly mentioned principles for human-centered AI (HAI) [33] — an approach to AI research and practice that aims to center human needs, societal good, and safety [35, 48, 51]. However, we know little about how practitioners who design and develop consumer-facing AI technologies define and scope privacy, what motivates and inhibits their privacy work, and the methods, tools, and resources they use in their work. Understanding these questions is essential because prior work suggests that there remains a substantial “gap between principle and practice” in HAI [51, 59]. While privacy is viewed as paramount to the development of human-centered AI technologies, we cannot

hope to adequately support AI practitioners in designing for privacy without first understanding their existing attitudes and workflows. Moreover, because AI technologies have the potential to pose unique privacy harms (e.g., facial recognition for police surveillance [30], deep fake pornography [12], training data reconstruction attacks [58]), and because the design pipeline for AI differs significantly from traditional software engineering [3, 18], there is reason to believe that the privacy challenges and processes faced by industry practitioners when developing consumer-facing AI products should differ from developing other products.

A recent broad survey of the usable privacy and security literature suggests that there are three broad barriers that users face in implementing privacy and security best practices: awareness of threats and mitigation measures, motivation to act, and ability to convert intention into action [21]. These barriers make up the Security and Privacy Acceptance Framework (SPAF). As a first step towards contextualizing the barriers practitioners face in AI privacy work, we pose three research questions corresponding to each of the SPAF barriers:

- RQ1** How well do AI practitioners’ definitions of privacy work reflect awareness of AI-exacerbated privacy threats?
- RQ2** What motivates and inhibits privacy work for consumer-facing AI products?
- RQ3** What constitutes privacy work for AI practitioners and what affects their ability to do this work?

To answer our research questions, we conducted semi-structured interviews with  $N = 35$  industry practitioners from 25 companies who engaged in privacy work for a consumer-facing AI product in some capacity. One of these interviews was a group interview with five practitioners who worked closely on a set of products. In our study, we define consumer-facing AI products as products that employ AI technologies that train on data from or about end-users and/or make inferences on data from or about end-users.

We found that practitioners viewed privacy as protecting users from intrusive or non-consensual uses of personal data (e.g., surveillance, secondary use). When reflecting on how they defined and/or situated privacy, these privacy harms were often introduced or exacerbated by the capabilities (e.g., the ability to identify individuals from their data) and/or requirements (e.g., the need to collect large stores of personal data) of AI. We also observed that practitioners primarily followed a compliance-centered approach in their privacy work. While prior work has noted that compliance requirements act as a forcing function for practitioners to practice and promote privacy [37, 54, 56], our findings further show how a compliance-centered approach: (i) allowed practitioners to prioritize privacy even if it was viewed as secondary to other design goals such as model performance; (ii) revealed the tensions between privacy values and other important objectives in AI product development; and, (iii) encouraged practitioners to conceive of privacy work as meeting minimum compliance standards with little-to-no end-user engagement, despite their defining privacy as minimizing harms to end-users. We also found that practitioners relied on design references and automated audits to help minimize privacy risks, but observed that the tools and artifacts they used were not specific to their product, or to the harms introduced or exacerbated by AI. As a result, practitioners felt ill-equipped to handle their privacy work and discussed the need for more product- and AI-specific guidance in the tools they employed in designing for privacy.

In summary, our work makes the following contributions:

- We extend the usable privacy literature on barriers developers face in privacy work to the context of consumer AI products, which can pose unique privacy risks. We provide clarity on how practitioners *define and scope* privacy work, what *motivates and inhibits* their work, and *what affects their ability* to do this work.
- We extend the human-centered AI (HAI) literature on developing ethical AI technologies by discussing how the principle-practice gap manifests for privacy in the development of consumer AI products.
- Drawing from our interview insights, we outline a vision for how we might better support practitioners by creating tool, artifacts, and support structures that help improve practitioners' awareness of AI-exacerbated privacy threats, motivation to address these threats, and ability to effectively address these threats.

## 2 Related Work

### 2.1 Human-Centered AI

Prior research has identified that AI technologies can generate results that many find intrusive, offensive, or unjust when

uncritically applied to high-stakes scenarios such as health-care and labor assessment [2, 19]. To resolve the negative externalities often entailed by AI, Human-centered AI (HAI) emphasizes that AI technologies should be created to be more socially responsible. Specifically, prior work has suggested building and testing reliable, safe, and trustworthy interactive AI systems by enhancing a team's awareness of HAI [35, 51]. In addition, Riedl et al. suggested that people should treat the AI algorithms as "part of a larger system consisting of humans" when ideating interactive AI systems [48].

Recent meta-reviews of guidelines seeking to operationalize HAI and ethical AI identified privacy as one of the most frequently mentioned principles [26, 33]. Yet, there remains a significant "gap between principles and practice" [59]: i.e., putting the principles into day-to-day practice remains a challenge for many AI practitioners.

To bridge this principle-practice gap, efforts have been made to model practitioner difficulties with applying HAI principles in practice and to provide practitioners with turnkey guidance for applying HAI principles to product design. Researchers have identified, for example, barriers to prioritizing human-centered approaches to AI design and development faced by AI practitioners in both industry and the public sector [32, 57]. Several major technology companies have aimed to reduce these barriers by publishing repositories to educate AI practitioners on how to build and test socially responsible AI technologies [4, 6, 45]. Governments have passed regulations for ethical AI usage and development that target AI practitioners [33]. Researchers have proposed checklists [28, 32, 39], guidelines [51], and practices such as data statements [10] and nutrition labels for publicly released datasets [31] to support stakeholders in mitigating ethical risks in AI system development.

Nevertheless, there has been little work seeking to model and improve how privacy is defined and designed for in AI development pipelines. Existing research and practice of privacy in AI focuses on protection, control, and agency over user data (e.g., consent for data usage, data protection, right to erase the data) [65], or the development of techniques that offer quantifiable privacy guarantees when processing data (e.g., differential privacy, federated learning). This work, however, takes a narrow view of what is helpful for practitioners to create privacy-preserving AI technologies.

To that end, we extend the literature on HAI by modeling how industry AI practitioners define and design for privacy.

### 2.2 Privacy in software engineering

Prior work has explored how practitioners incorporate privacy principles into software engineering more generally. Some research has explored and collected developers' attitudes toward privacy through surveys [50], interviews [7, 36, 38, 56], and sourcing comments from public online platforms [37, 56]. Rather than building conceptual frameworks (e.g., [42]) of



*how privacy should be*, these efforts have helped identify common patterns in how software engineers *actually think about and approach privacy*. Specifically, prior work suggests that privacy is often a secondary concern, which is seen as the main obstacle to improving privacy practices in software engineering [7, 36, 54, 55]. Li et al. conducted two studies with developers and found that privacy practices were seen to bring extra costs but generate few benefits [37, 38]. Prior work also identified other barriers to incorporating privacy in software engineering: i.e., privacy misconceptions [36, 38], knowledge gaps [36], and the lack of guidelines and regulations [36].

To help practitioners overcome these barriers, tools and methods have been proposed to facilitate incorporating privacy and security best practices in development, including development tools that detect privacy and security issues [36, 41, 66], and design artifacts (e.g., agendas, workbooks, and guidelines) that engage stakeholders in designing to mitigate privacy and security concerns [1, 60]. In addition, researchers have also studied the effectiveness of these interventions, and sought to understand how to improve the experience of using these tools. For example, Tahaei et al. conducted a study on developers' attitudes toward security notifications in the context of static analysis [55].

A recent survey on what leads to users accepting and rejecting best practices in security and privacy suggests that there are three key barriers users must overcome: awareness, motivation, and ability [21]. These barriers make up the Security and Privacy Acceptance Framework (SPAF) and can also help explain why developers and practitioners accept or reject privacy best practices [ibid].

We extend the prior literature on privacy in software engineering by empirically exploring and modeling the SPAF barriers practitioners face when developing consumer-facing AI products. Doing so is important because the unique capabilities of AI can introduce new conceptualizations of privacy [43], and because the design and deployment pipeline for AI products differs significantly from that of traditional software engineering [3, 18]. To our knowledge, our work provides the first in-depth insights into how practitioners define and scope privacy work for consumer-AI products (Section 4.1), what motivates and inhibits their privacy work (Section 4.2), and the methods, artifacts, and resources they use in their work (Section 4.3).

### 3 Method

We conducted semi-structured interviews to inquire into participants' experiences with privacy work for consumer-facing AI products and services. The semi-structured approach allowed us to remain consistent across participants, while affording the flexibility to ask in-depth follow-up questions as necessary. For our interview study, we defined "consumer-facing AI" broadly as products that employ AI technologies that train on data from or about end-users and/or make infer-

ences on data from or about end-users. In total, we conducted 31 interviews (30 individuals, 1 group) with 35 participants working on products or services that involve consumer-facing AI as defined above. We refer to our participants as "AI practitioners," and their products as "consumer-facing AI products."

#### 3.1 Semi-structured Interviews

We started each interview by asking participants to think of a specific consumer-facing AI product for which they or their team had engaged in privacy work during the development process. We situated the questions we asked for that particular product, following our main research questions. We developed three sets of interview questions accordingly. To answer **RQ1**, we first asked practitioners how they defined and scoped privacy for the consumer-facing AI product in question. We then analyzed these practitioner-driven definitions to better understand their *awareness* of the AI-exacerbated privacy harms discussed in prior literature. To answer **RQ2**, we asked practitioners what motivated and inhibited their privacy work for the product in question. We then analyzed their responses to synthesize factors that address their *motivation* to do privacy work in the context of developing consumer-facing AI products. Finally, to answer **RQ3**, we asked practitioners about the actions, tools, artifacts, and resources they utilized in their privacy work, the challenges they faced, and what tools they could envision that, if provided, would have been helpful for their privacy work for the product in question. We then analyzed their responses to understand what affects practitioners' *ability* to perform privacy work for consumer-facing AI products. In sum, building off and extending the SPAF [21], our interviews helped us uncover the awareness, motivation, and ability barriers practitioners face in doing AI privacy work.

We first ran four pilot interviews to ensure that the questions we asked were easy to interpret and answer in a manner that provided us with data pertinent to answering our research questions. After concluding the pilot interviews, we found that only minor edits were necessary — mostly revising questions that were too generic and leaving some questions as "optional" for the interviewer to cover in the interview to ensure the interview session could be completed within 40-60 minutes.

#### 3.2 Study Procedure and Recruitment

During the scheduling process, we provided an option for participants to invite co-workers working on the same product(s) for a "group interview." Group interviews followed the same protocol as described in Section 3.1, though participants were also encouraged to engage in discussion with other participants in the same interview session. We completed one 90-minute group interview session with a team of five (P20-P24). The rest of our 30 participants were interviewed individually, with each session lasting between 40-60 minutes.

All interviews were conducted remotely. We compensated all participants with a \$100 USD gift card.

The interviews were conducted in English ( $n=23$ ) and Chinese ( $n=8$ )<sup>1</sup>. The first author conducted 27 of the 31 interviews; two other authors conducted two English interviews and two Chinese interviews, respectively. When possible, a second interviewer would also join the interview session to take notes and to occasionally ask follow-up questions.

We first explained the purpose of the study to each participant. We provided participants with a written consent form before the interview, and asked for verbal consent before we began the interview; the consent form also included a section for participants to consent to the audio and video recording of the session. We also provided participants with a preview of the interview protocol before the interview, in part so they could be prepared to answer our questions in a manner that did not reveal secret or proprietary information. We also informed participants that they could terminate their interview and withdraw their consent to participate in the study at any time. They were also informed that their audio recordings would be transcribed in a de-identified manner for further analysis, and that the raw audio content would be deleted post-transcription. Our study was approved by both the CMU and Georgia Tech IRBs.

We recruited industry practitioners who “have experience with designing and/or developing consumer-facing AI products,” and “have participated in discussions about end-user privacy as it pertains to consumer-facing AI products that they have helped build.” We reached out to potential participants via the authors’ professional networks, such as through social media platforms (9/35), i.e., Facebook, LinkedIn, Twitter, Slack, Discord, and the alumni networks of our institutions (18/35). We also recruited from the authors’ direct contacts who matched the criteria (8/35). In sum, our 35 participants shared their experiences on how privacy is defined, motivated, and practiced in designing and developing consumer-facing AI products across 20 technology companies and five start-up technology companies. Except for the group interview, those who worked at the same companies were from different product teams. Participants worked across a range of products and team roles (see Appendix Table 4). The most common four AI technologies our participants incorporated into their consumer-facing AI products were: recommender systems ( $n=14$ ), conversational AI/chatbots ( $n=10$ ), natural language processing tools ( $n=10$ ), and predictive analytics ( $n=10$ ). The most common three application domains for these products were: healthcare ( $n=8$ ), general-purpose machine learning tools ( $n=8$ ), and media and entertainment ( $n=7$ ). The three most common roles our participants identified with were researcher ( $n=16$ ), software engineer ( $n=13$ ), and designer ( $n=13$ ). Their ages ranged from 23 to 48 ( $M=31.82$ ,

<sup>1</sup>We informed participants that the questions would be asked in English, but that they could choose to reply in the language with which they felt more comfortable speaking.

$SD=6.71$ ); 16 identified as male, 14 identified as female, and five preferred not to disclose their gender identity.

### 3.3 Data Analysis

All interview sessions were first audio recorded and transcribed<sup>2</sup>. Then, we conducted an iterative, open coding process [17] on participants’ responses following our three research questions. The first author performed the initial coding on ten interview transcripts, and iteratively constructed a codebook in active discussion with three other authors. Another author joined the coding process when the initial codebook was constructed, and was trained with the codebook. The two coders coded the same six interviews individually and reached an agreement on all codes. They then split the interviews and individually coded, meeting regularly to discuss codes and themes. They regularly reviewed the other coder’s codes (i.e., interview snippets that were applied codes) to ensure they applied the codes similarly, and resolved all disagreements. All the authors also regularly met and discussed emerging themes during the coding process. We present the key themes, guided by our research questions, in Section 4, and include the codebook in Appendix Table 3.

### 3.4 Limitations

This study has several limitations. First, because our findings are qualitative and based on our participants’ actual experience and practice, they should not be interpreted as representative of all AI practitioners and industry contexts. Additionally, to make our interview protocol reflective of practitioners’ actual workflows, we consciously did not prime our participants about what are and are not AI-specific intrusions and practices. Thus, they might have responded differently to a protocol with a different focus. Our goal was to explore how practitioners defined AI privacy work, what motivated and inhibited this work, and what affected their ability to do their work when building AI products so that we might generate insights and hypotheses on how the community might improve existing practices. To that end, our method was appropriate for our goals. Second, in order to cover a wide range of privacy considerations and practices, we recruited practitioners who specifically had experience with doing privacy work in developing consumer-facing AI products. This inclusion criterion may have skewed our sample toward participants with higher privacy awareness. Finally, due to our recruiting strategies, our sample is not perfectly generalizable: most of our participants worked at North American and European companies. We note that there is an asymmetric risk for institutions to be upfront about privacy practices: poor privacy practices can result in a media firestorm, but good privacy practices are unlikely to be lauded and popularized except by niche audiences.

<sup>2</sup>The two coders engaged in the data analysis were both English-Chinese bilinguals. Thus, the interview data was not translated before the analysis.

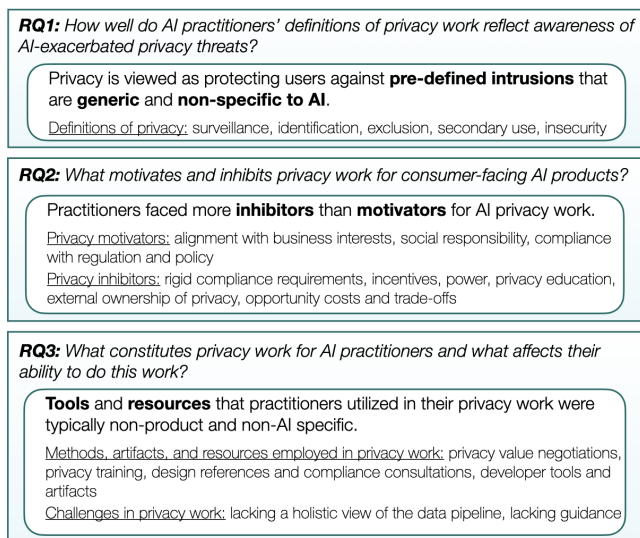


Figure 1: We answer our research questions by showing how practitioners define and scope privacy work (RQ1), what motivates and inhibits their work (RQ2), and what affects their ability to do this work (RQ3).

Thus, part of the challenge in unearthing how AI practitioners approach privacy is that institutional privacy practices may be intentionally opaque, similar to practices in pursuit of other HAI principles such as fairness [32, 57]. To that end, while we can say that our practitioners came from a wide range of organizations, worked on a broad variety of consumer-facing AI products, and hailed from different parts of the world, we cannot reveal exactly which organizations and products. Nevertheless, the practitioners we interviewed were forthcoming to the extent allowable by their employment agreements.

## 4 Results

We present our findings as they relate to three core research questions of interest (see Figure 1). We report on how frequently participants discussed the identified themes.

### 4.1 RQ1: How well do AI practitioners' definitions of privacy work reflect awareness of AI-exacerbated privacy threats?

The first barrier outlined in the SPAF is the awareness barrier — people's understanding of context-specific threats and the mitigation measures thereof [21]. To model awareness barriers practitioners faced in AI privacy work, we first analyzed what participants thought of as privacy work when designing and developing consumer-facing AI products. We then explored to what extent participants' definitions reflected an awareness of the unique privacy harms entailed by AI technologies that have been discussed in prior literature (e.g.,

memorization leaks [13] and membership inference attacks [52, 67]).

#### 4.1.1 How practitioners defined privacy

From the use of facial recognition technologies to surveil minorities [20], to issues of knowledge and consent in how one's personal data is used to train large models [9, 44], to the amplification of embarrassing content in recommender systems [15], prior literature has documented a number of consumer privacy concerns uniquely created or exacerbated by AI. To identify and understand awareness barriers in AI privacy work, we first explored how practitioners defined privacy in the development of consumer AI products and analyzed how their definitions aligned with the emergent privacy risks arising from the unique capabilities and requirements of AI. Although our participants worked across diverse roles and application domains (see Appendix Table 4), the common denominator was that they viewed privacy as the need to protect consumers when creating consumer-facing AI technologies. To formally categorize practitioners' views of potential privacy intrusions that their AI products could entail, we mapped our findings to the broader taxonomy of privacy harms proposed by Solove [53]. The taxonomy comprises intrusions pertaining to collecting and processing of personal data, which is closely aligned with how practitioners conceived of privacy work when building consumer-facing AI products (see Table 1). Our participants identified the need to address some of the privacy harms entailed by the unique capabilities and requirements of building AI systems: e.g., the collection of personal data to train effective machine learning models, and the processing of personal data to make predictions about user preferences and actions. Many other documented AI-entailed privacy intrusions, however, did not appear to factor into how participants defined privacy in the context of AI.

**Surveillance** (3/35) refers to the automated monitoring and collection of personal data, but not necessarily its direct use [53]. Many of our practitioners alluded to how AI products can help create a surveillance infrastructure due to 1) the AI-afforded capability to monitor specific individuals in vast data streams, or 2) the requirement to train effective AI models on large-scale personal data. P7, a designer working on a B2B consumer-facing AI product that aims to improve the well-being of company employees, discussed how this product could also potentially raise surveillance concerns: *"they actually collect data on people who come into the office, and they can even answer questions: how long are you working based off your work laptop... and could even get a lot of people in trouble by maybe... their employer saying you're working less than everyone else."* Practitioners noted that AI products could exacerbate surveillance intrusions by incentivizing and encouraging the harvesting of personal data to improve model performance. They surfaced tensions between the *utility* of the AI products they created and the *intrusive-*



Table 1: Participants expressed the following privacy concerns that are exacerbated by the capabilities and requirements of AI in consumer-facing AI products: surveillance, identification, exclusion, secondary use, and insecurity [53].

| <i>Privacy intrusions</i>                                                                                                                                                                | <i>Practitioners' concerns about privacy intrusions for consumer-facing AI products</i>                                                                                                                                                                                |
|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>Surveillance (3/35)</b><br>Privacy intrusions resulting from watching, listening to, or recording individuals' activities                                                             | Surveillance intrusions are exacerbated by creating a surveillance infrastructure due to the <i>AI systems' ability to monitor individuals' activities</i> and the <i>requirement to collect large-scale personal data to train effective AI models</i> .              |
| <b>Identification (10/35)</b><br>Privacy intrusions resulting from linking information to particular individuals                                                                         | Identification intrusions are exacerbated through the presence of <i>personally identifying information (PII) about users in the machine learning data pipeline</i> and are directly introduced by the <i>capability of AI models to make inferences about users</i> . |
| <b>Exclusion (4/35)</b><br>Privacy intrusions resulting from failing to allow the data subject to know about the data that others have about her and participate in its handling and use | Exclusion intrusions are exacerbated by a <i>lack of awareness of how personal data is being used by AI</i> and a <i>lack of agency over how personal data is being used by AI</i> .                                                                                   |
| <b>Secondary use (2/35)</b><br>Privacy intrusions resulting from using information collected for one purpose for a different purpose without the end-user's consent                      | Secondary use intrusions are exacerbated by the <i>(re-)use of users' personal data to train new AI models</i> without securing consent for the new use.                                                                                                               |
| <b>Insecurity (14/35)</b><br>Privacy intrusions resulting from leaks and unauthorized access of personal data                                                                            | Insecurity intrusions are exacerbated by <i>poor security practices</i> that can lead to, e.g., unauthorized access, personal data leaks, or personally identifiable data that is collected and/or used to train consumer-facing AI products.                          |

ness of collecting the data necessary to unlock that utility. A researcher working on a health-care/well-being system (P18) stated “*So definitely... in terms of research, the more data, the merrier*”. However, when they tried to productionize the product — which collected users’ browser activity such as page views, keyboard, and mouse interactions — P18 expressed that they did not collect as much data as might have been useful for improving model performance in order to respect “*how much data participants will be willing to contribute*.”

**Identification** (10/35) refers to the threat of being able to link specific data points to an individual [53]. Practitioners noted that AI products might 1) exacerbate this threat through the presence of personally identifying information (PII) about users in the machine learning data pipeline, and 2) directly introduce the threats through the capability of AI models to make inferences about users. To mitigate this threat, practitioners shared their attempts to “sanitize” the data fed into the data pipeline to reduce the presence of PII. For example, a designer and content expert (P9) working on a chatbot manually de-identified the dataset used to train the model, “*we get our clients, customer service records, we usually would remove some personal data from those records... like names and phone numbers*.” When a technical lead (P8) discussed his team’s approach to training privacy-respecting machine learning models, he alluded to concepts from differential privacy [23]. He viewed privacy as “*not being able to pinpoint behavior to a single origin*,” and noted that “*you should only be able to analyze things in aggregate manners, and not be able to do that root cause to a single point that’s potentially causing a behavior*.”

**Exclusion** (4/35) takes place when user awareness and agency over the use of personal data are limited [53]. Some practitioners defined privacy as ensuring that their users were aware of how their data was being used by AI. Other practi-

tioners defined privacy as affording users agency over *how* their data is being used by AI. A product director (P34) discussed how they address such concerns by “*deleting that information, eventually, when the customer stops being a customer or upon request for some reason*.”

**Secondary use** (2/35) threats encompass the (re-)use of personal data that are used to train consumer-facing AI products without consent for other purposes [53]. For example, a software engineer (P3) stressed the importance of not repurposing a dataset collected from users: “*we want people’s information that they give us to be safe and not used for anything else other than actually recommending them clothes*.”

**Insecurity** (14/35) threats result from poor security practices that can lead to, e.g., personal data leaks [53]. This privacy concern is associated with the rich personally identifiable data that practitioners collect and/or use to train the AI they use in their products. For example, a researcher (P22) working on customer service/management technology described the need to “*put borders around certain data, but also internally share the data and understand what’s happening*”. Others expressed the need to ensure appropriate and secure data storage and adhere to data retention policies.

#### 4.1.2 Summary: Awareness barriers in AI privacy work

According to the SPAF [21], our participants exhibited limited awareness of how the capabilities and requirements of AI might affect the privacy threats entailed by a product. While a number of the privacy harms and intrusions that practitioners discussed were not specific or exclusive to AI technologies, AI technologies can exacerbate those threats in a way that practitioners did not specifically highlight or mention.

Consider, for example, exclusion threats in which users have limited awareness and agency over how their personal

data is used. Large Language Models (LLMs), which are trained on massive corpora of textual information scraped from the web [11], can significantly exacerbate exclusion threats. Indeed, as noted in prior work, it can be difficult for any individual to exercise control over how data they have shared online can and cannot be used by such models [9, 44]. Similarly, while secondary use threats long predate the use of AI technologies, AI technologies can again exacerbate these threats. For example, transfer learning techniques, in which a pre-trained foundation model is fine-tuned to new contexts of use, allow for rapid prototyping of context-specific ML models. However, uncritical use of transfer learning can lead to secondary use threats — consent acquired for using personal data in the original model may not necessarily translate to the new model and some privacy risks have been shown to carry forward into derivative models (e.g., [67]). Prior academic literature on the security of machine learning has highlighted that AI can also exacerbate privacy threats that result from poor operational security — e.g., adversarial attacks that can reconstruct the raw personal data on which commercially deployed models were trained [52]. However, our participants did not discuss these AI-exacerbated privacy threats — either because they were unaware of these threats or because the structures in place to think about privacy for AI products remain generic and non-specific to AI.

## 4.2 RQ2: What motivates and inhibits privacy work for consumer-facing AI products?

The second barrier described in the SPAF is motivation — whether or not people *want* to act in accordance with best practices for security and privacy [21]. To model motivational considerations in AI privacy work, we analyzed participants' interview responses around decision-making processes and the reasons participants engaged in or deprioritized privacy work for their consumer-facing AI products. We found that practitioners' privacy work was primarily motivated by compliance requirements, alignment with business interests, and social responsibility. In contrast, we found that practitioners' privacy work could also be inhibited by the rigidity of compliance requirements, incentives, power, education, ownership, and opportunity cost.

### 4.2.1 Privacy motivators

We first highlight the three motivators for privacy work in the development of consumer AI products: alignment with business interests, internal motivation to build socially responsible AI, and compliance with privacy regulations.

**Alignment with business interests:** Some participants (9/35) discussed competitive differentiation and client concerns as key motivators for privacy-related design deliberations. For example, a technical lead (P28) in an AI startup company noted: *“privacy can be a differentiator. And when*

*you're doing a startup, especially if you're in a crowded space, you're looking for any way, any angle that you have to say that you're different from other things that are out there.”* A researcher (P22) discussed mitigating client concerns as a key motivator for privacy deliberation in their customer relationship management product: *“we deal with tickets in... HR issues. Sometimes they're very sensitive topics... in proposing a new idea for... use of AI for helping agents to maintain and see through tickets quickly... we saw lots of concerns from customers on that.”* Given the media and public relations risk associated with privacy mistakes in AI technologies (e.g., [5, 30]), building a privacy-respecting business model is of increasing interest to many organizations.

**Social responsibility:** Some participants (5/35) discussed a personal desire to build socially responsible AI as their motivation for considering privacy when designing consumer-facing AI products. For example, a researcher (P35) working on machine learning building tools noted: *“the people that tend to come here that are building new ML features are aware of bad cases of ML being... inappropriately applied. And no one wants to have that happen... I don't think there need to be models that decide whether or not people are going to recommit a crime.”* Likewise, when discussing why they abandoned a particular product direction that would involve secondary use of personal data, P25 said: *“such an act was not clearly defined, and because the user does not give us consent [for other purposes]. So later we figured our morals won't let us do this; we must get user consent.”*

**Compliance with regulation and policy:** Finally, echoing results from prior work on general developer motivations to act on privacy [54], complying with external regulatory mandates, such as those imposed by the EU General Data Protection Regulations (GDPR), was a key catalyst for our participants' privacy work on consumer AI products (19/35). For example, a product director (P34) listed several regulatory requirements with which their products had to comply: *“we are completely GDPR compliant. We're also very involved in the privacy frameworks that are in Canada... there's the CCPA in California, as well as the new Virginia law... because we deal with... PII [personally identifiable information] in all of these locations, including European customers. So we have compliance in all of those, and we strive to maintain that compliance.”* P31 went a step further, suggesting that their privacy work *“are considered compliance... we don't do it by choice, like it's always enforced.”* Some participants also discussed privacy reviews conducted by *external* teams (e.g., the privacy team) as a catalyst for privacy work: *“if you didn't give it enough consideration, chances are, it's not gonna pass the privacy review...”* (P30).

### 4.2.2 Privacy inhibitors

We next discuss the six factors that inhibited practitioners' privacy work. Factors that impact practitioners' privacy and



security practices have been studied more broadly in software engineering [29]. In this section, we build and extend on these findings by highlighting how AI products can create new and exacerbate known inhibitors that hinder privacy work.

**Rigid compliance requirements:** While compliance with privacy regulations often motivated privacy work, we observed that compliance-driven approaches to privacy sometimes also inhibited practitioners (6/35) from going beyond minimum requirements to engage in more human-centered and product-specific conceptions of privacy. Privacy, in other words, simply meant compliance. For example, a lead designer (P16) discussed experiences where she advocated for privacy in product meetings, but would be met with resistance from co-workers because the product was already “compliant”: *“most of the time, especially in general product development, and what the engineers are doing, it’s so standardized, that’s not really a conversation, because there’s nothing to be done about it. It just is the way that it is.”* A technical lead (P8) mentioned that they strictly followed compliance and customer requirements as guiding principles in designing for privacy, but this approach also precluded them from taking a broader view of privacy in design: *“we don’t self impose like, hey, this is a gray area. So... we think about this, but we don’t have that as the main driver of requirements.”*

**Incentives:** Practitioners (7/35) discussed how advocating for privacy might be indirectly misaligned with career incentives. As a technical lead (P11) pointed out: *“people are not really incentivized to do this correctly. And if they wanted to do things correctly, it becomes extra effort, and influences their completed work, fewer results, and as a result they get promoted slower than their peers.”* A lead designer (P16) further stressed that advocating for privacy can lead to tension, which can be perceived as dampening excitement: *“I like to joke that my job is to be the buzzkill who stops other people from doing things...I’m like... you need to consider that to consider this or else you can’t do it... So, there’s also a lot of tension.”* P31, an engineer, discussed how the prioritization of speed in job assessments further factors into how incentive structures might inhibit privacy work: *“we’re always rewarded for delivering things on time, and as fast as we can. So that’s probably one of the biggest reasons where privacy and to large extent even security becomes an afterthought, because we just want to get things done.”*

**Power:** Other practitioners (3/35) discussed feeling powerless due to organizational structures. This sense of powerlessness negatively impacted how practitioners viewed their individual efforts toward designing for privacy because of *“a disconnect between whatever they’re [individual contributors] able to say versus who’s making decisions”*, as noted by a designer (P7).

**Privacy education:** While some participants saw privacy as a competitive differentiator and motivator for privacy work (Section 4.2.1), others (6/35) acknowledged that the low visibility of privacy in a company could be the result of an

organizational-wide lack of privacy education. For example, a UX researcher (P17) shared that *“[the team is] not really well versed in these things. So they are more just let’s just do something and ask questions later.”*

**External ownership of privacy:** Some other participants (5/35) shared that because “privacy” was owned by a different team in their organization, they did not consider privacy at all in their workflow. For example, a software engineer (P15) noted: *“there’s a legal team totally... for that purpose. When we do our job, we don’t know who this user is, and we don’t really look at their privacy.”* Externalizing privacy “as the responsibility of others” has been noted in prior work [34,63]. Moreover, as Wong and Mulligan allude, offloading privacy as a matter strictly for legal counsel limits the role that design teams have in advocating for the needs of their users [60].

**Opportunity costs and trade-offs:** More than any other inhibitor, our participants discussed how privacy work comes with a number of opportunity costs and trade-offs that make it difficult to prioritize. As prior work has shown in the context of end-user attitudes and behaviors towards security and privacy, privacy is often a *secondary* concern that comes after other priorities like usability and functionality [22]. P16 captured this ethos in describing what inhibited their privacy work: *“we have bigger fish to fry, you know, the higher priorities of things we got to do”*. Development resources for AI products are limited and schedules are tight, and negotiating the value of privacy relative to other design goals under these constraints was complicated by the fact that while the benefits of designing for privacy were abstract, its costs were easier to make tangible. We observed this de-prioritization of privacy work in a number of ways.

Practitioners (9/35) discussed prioritizing **functionality and user experience** the AI-infused features offer over their potential privacy harms. For example, a software engineer (P2) mentioned: *“since what we’re working on directly faces our users, the first thing that comes to our minds is not privacy, but user’s experience. More specifically: user search quality, and all other aspects of their experience. Privacy will not be considered until we know our users have a positive experience.”*

Other participants (7/35) brought up trade-offs between privacy and AI-fueled **business objectives**. As a UX researcher (P27) working on advertisement recommendation put it: *“if we weren’t tracking people’s behaviors there, there’s no way that we could provide insights for advertisers.”* A product director (P34) for a job matching tool discussed how their “aggressive” privacy policy of deleting data for customers who had not used their system for a while impacted their ability *“to do data science because we’re losing historical data that can be used as a baseline.”* Some participants (3/35) explicitly expressed their view that designing for privacy can inhibit product innovation. For example, a research director (P6) discussed how privacy *“slows down and tampers some of the creative aspects of projects”*. A software engineer (P25)

further discussed how more “conservative” co-workers could weaponize the use of privacy to inhibit more progressive design explorations, because it is hard to argue against the need for privacy compliance: *“When the [conservatives] bring up...their more conventional ways, you come to think: if you want to break the existing framework or push those who’d like to move forward, do you have any kind of “weapons?” or tools that balance the two sides? Because once [conservatives] talk about their experience or what will happen next, it is not easy to fight against them.”*

Practitioners (7/35) also described trade-offs in **model performance** entailed by privacy compliance. Indeed, understanding how to train machine learning models that are privacy-preserving while maintaining high performance remains an open and active research question [46, 64]. Some of our participants (3/35) explicitly noted that this trade-off stems from the fact that the data to which a model has access may be more coarse-grained owing to privacy compliance: *“we are getting less and less idea about, for example, what an end-user is like, if having a higher standard of privacy”* (P5). Others (4/35) mentioned the performance trade-off stems from the fact that a model has access to less data, making training effective models more challenging. A machine learning engineer (P31) noted: *“taking care of privacy means we might use less data, or we might remove some information from our data, which might degrade the model performance... It’s a challenge to have a privacy-preserving model, and also achieve the same performance.”*

Many practitioners (19/35) also discussed how designing for privacy entailed additional **engineering costs** during development. These included the need for a more complex data pipeline when designing an AI system. A software engineer (P5) stated: *“[privacy compliance] might make the system or the design more complicated, instead of a more straightforward idea, saying like, hey, we could just leverage the data of something else that will have to go way around.”* Prolonged development time was also a common “cost” of privacy. Participants mentioned that privacy requirements created bottlenecks during the development process, which were mostly caused by the time and effort spent on validating that their products met compliance requirements (7/35), or review processes conducted by external privacy teams to get access to training datasets (P25, P30) or getting approval for a product proposal (P6, P10, P33).

#### 4.2.3 Summary: Motivation barriers in AI privacy work

In the context of the motivation barrier described by the SPAF [21], we found that practitioners face many more inhibitors than motivators for privacy work in developing consumer AI products, and thus exhibit low overall motivation to engage in privacy work beyond minimum compliance requirements. These compliance requirements, in turn, were generally non-AI specific. Moreover, those who advocated

for privacy work beyond the minimum often did so with trepidation and with the knowledge that this work was not well aligned with existing incentive structures for performance assessment and promotion. We also found that while compliance requirements are necessary and useful as forcing functions for privacy work, they can also inhibit more creative explorations of how to design privacy-respectful products.

Indeed, many of these inhibitors and motivators we discussed in this section are not unique to the workflow of AI products — e.g., limited development resources, tensions between different product objectives, and barriers to communicating with different stakeholders have been found to hinder security and privacy best practices in software engineering more broadly [29, 54]. Nevertheless, the capabilities and requirements of AI can exacerbate these inhibitors. For example, costs to model performance in the name of privacy can measurably degrade user experience. Collecting less personal data can reduce the utility of advertiser tooling, reducing the monetizability of a product. The technical and organizational barriers may be greater for privacy work in AI product development because AI-specific privacy-preserving solutions vary across the product life-cycle [47]. Moreover, each stage of AI/ML workflow, which can be either data-oriented or model-oriented, has different requirements and objectives that can conflict with privacy [3]. All the while, the “benefits” of improving privacy are more abstract and often simply boil down to compliance requirements imposed by external teams and regulation. These inhibitors complicate how practitioners prioritize privacy with respect to other product objectives, and how they communicate and coordinate privacy work with other teams [29].

### 4.3 RQ3: What constitutes privacy work for AI practitioners and what affects their ability to do this work?

The final barrier discussed in the SPAF is the *ability* barrier — i.e., the challenges that people face when translating intention into action in the context of security and privacy [21]. To understand the ability barriers practitioners face in their AI privacy work, we analyzed responses to questions about the specific activities they considered as a part of their privacy work, as well as the artifacts they envisioned might help them with their privacy work. We also summarize the ways in which practitioners require more support in their AI privacy work.

#### 4.3.1 Methods, artifacts, and resources employed in privacy work

Our participants described using an assortment of methods, artifacts, and resources in their privacy work for consumer-facing AI products, including ad-hoc privacy value negotiations, privacy trainings, design reference materials & compliance consultations, and developer tools & artifacts. We

Table 2: AI practitioners rely on various sources, tools, and artifacts when designing for privacy: i.e., privacy training, design documentation, privacy & legal experts, developer tools, and privacy checklists.

| <i>Sources, Tools, and Artifacts</i>     | <i>Descriptions</i>                                                                                            | <i>Examples</i>                                                   |
|------------------------------------------|----------------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------|
| Privacy training (18/35)                 | <i>Generic</i> but <i>mandatory</i> training that practitioners' employers used to educate them about privacy. | New-hire privacy training, employee annual training               |
| Company-wise design documentation (9/35) | <i>Task-specific</i> and <i>on-demand</i> design documentation about privacy.                                  | Company-wise design references, prior designs from the company    |
| Privacy & legal experts (11/35)          | <i>Task-specific</i> and <i>on-demand</i> privacy/legal consultant.                                            | Privacy and legal teams in the company                            |
| Developer tools (3/35)                   | <i>Task-specific</i> and <i>mandatory</i> tools for product development.                                       | Azure DevOps, IDE plugin, privacy notifications/ pop-up questions |
| Privacy checklists (3/35)                | <i>Task-specific</i> and <i>mandatory</i> procedures and processes to ensure privacy compliance.               | Risk assessment checklist, privacy review form                    |

summarize sources, tools, and artifacts that practitioners rely on when doing privacy work in Table 2.

**Privacy value negotiations.** Our participants described situations in which privacy compliance conflicted with other important AI system design considerations, such as product objectives, model performance, and development resources. Often, tensions between these design objectives and privacy largely resulted from privacy harms rooted in the capabilities and requirements of the AI technologies employed in a product. When such conflicts arose, there was a need to negotiate the value of privacy relative to other goals. Participants (5/35) sometimes engaged in **ad-hoc risk assessments** to rationalize away the need to center privacy in design deliberations for AI-infused features. For example, a data scientist (P32) working on recommendation system asserted they did not engage in privacy discussion because their developments were based on a mature product: *“the frequency [for discussion around user privacy] is not high in part because so much of it is derivative of the same type of work, and so you just get to this point of like, you’re not trying to add anything or adjust anything in a way.”* In contrast, they also noted *“anytime there’s a new kind of feature instead of recommender... some other AI-based feature, then all the band-aids get ripped off, and all the discussion starts all over again. Because now you’re moving into new territory.”* A researcher (P30) working on machine learning optimization tools that were used by many consumer-facing AI products asserted that there was little privacy concern in their tools since they abstracted user data to the extent that they believed the data utilized was not sensitive: *“often the data we use in our insights data is a typical mathematical synthesized test function. So in that case, I don’t think there’s any privacy concern is... that we should think too much about”*. Other participants described needing to “advocate” for privacy beyond meeting minimum compliance requirements, as described by P16: *“It’s really just like going an extra mile of like, is this meeting users’ concerns there? And they’ll be like, Oh, well, it’s meeting the requirements. And I’d be like, okay, the bare minimum is not very good”*.

**Privacy training.** Many participants (18/35) mentioned

that the companies in which they were employed had company-wide trainings to educate them about privacy. This training was typically mandatory, and happened several times a year or when employees were newly hired. However, these trainings were often described as generic and not directly useful for developing consumer-facing AI products. For example, a software engineer (P2) commented that their required training provided *“nothing but a general concept.”* Unsurprisingly, many of our participants (11/35) expressed a desire for improved education about privacy requirements, regulations, and design references that pertained specifically to their work. A product director (P34), for example, wanted education to help clarify misconceptions and ambiguity about privacy-pertinent regulations that can affect design: *“you end up talking to people where you quickly realize that they have a misconception about GDPR or the California Act, or alien privacy regulations. And then that’s actually affecting the way that they think about the design of something.”*

**Design references and compliance consultations.** Practitioners (16/35) also referenced design documentation and consulted privacy/legal experts when designing for privacy. These events were on-demand, and our participants (9/35) shared some use cases for these internal design references. For example, a UX researcher (P33) working on a chatbot discussed referring to design examples in these internal references: *“some of it is examples of what other teams have done... there’s like learnings from other groups that we can take advantage of... like, how do other teams collect terms of service, or how do other teams do platform agreements?”* Some participants (11/35) mentioned dedicated privacy and legal teams that are responsible for company-wide privacy compliance with regulations and policies, as resources to reduce privacy risks and resolve regulatory confusion. For example, referring to training and fitting new machine learning consumer-facing models, a machine learning engineer (P31) said: *“we refer to our legal experts whenever we are confused or when we feel we don’t know if we’re doing the right thing.”*

While these design references and compliance consultations were not tailored toward the use of AI, they provided



ad-hoc guidance for privacy concerns that may be exacerbated by AI. For example, a compliance consultation can help clarify privacy regulations and reduce the likelihood of implementing an insecure or leaky AI data pipeline. Nevertheless, these ad-hoc privacy tools rely on practitioners' individual awareness of AI-exacerbated privacy intrusions to "find the right reference" or "ask the right question," which itself can be challenging when designing for privacy. Accordingly, participants (4/35) expressed a desire for more support to find appropriate design references emblematic of *"those best practices around the ways in which you can ship things without requiring end-user data, or other strategies around training models on end-user datasets... in a distributed way..."* (P35).

**Developer tools and artifacts.** Our participants (5/35) also expressed using a range of developer tools and/or artifacts to support their privacy work. These tools were task-specific and their use was mandatory, but they were not tailored toward the use of AI nor did they help sharpen focus on the privacy challenges entailed by AI.

Nevertheless, practitioners still found the tools useful, even if insufficient. Some of these tools performed automated reviews of source code to flag potential insecure data practices: e.g., a machine learning engineer (P31) noted: *"we get flagged for any privacy violations...if we use like Azure DevOps to check in our code."* Participants also mentioned tools that encouraged reflection on the use of user data in development. P30, a researcher working on machine learning optimization tools stated: *"[there] are checkboxes or pop-up questions [that ask] you to confirm whether the current query or the current piece of code you're using is touching user data or not."* Other artifacts helped practitioners follow specific procedures to ensure compliance: e.g., checklists and forms. These artifacts helped practitioners ensure secure stewardship of user data — e.g., to draft a review procedure (P30) — and to mitigate non-consensual secondary use intrusions — e.g., to review the privacy policy of third-party vendors (P34).

However, practitioners requested tools and artifacts that provided more product- and AI-specific guidance. Some participants (4/35) expressed the need for a checklist to inform standard practices in designing for privacy, and to allow practitioners to individually assess privacy compliance for their AI products. As a data scientist (P32) elaborated: *"the checklist of private means location, and you have to change it to no location [data], no email, no whatever... like we certify each of these things through checklist."* A UX researcher (P17) further stressed that it could be "really impactful" to have an artifact that helps practitioners assess whether or not they have *"met the minimum standards that we have for privacy... it goes all the way straight through to QA [quality assurance]."*

### 4.3.2 Summary: Ability barriers in AI privacy work

In the context of the SPAF [21], we found that practitioners faced a number of significant ability barriers in their pri-

vacuity work. Specifically, practitioners lacked a holistic view of the data pipeline for their AI products, and "last-mile" guidance for how to best approach their privacy work for their specific product and their use of AI. These challenges revealed that practitioners were not readily equipped to design for privacy. While their privacy efforts were mostly compliance-centered, individually, they expended significant effort in learning how compliance requirements applied to individual products. Moreover, we found that many of these challenges manifested when practitioners attempted to be proactive and initiate privacy practices that went beyond minimum compliance requirements.

**Practitioners lack a holistic view of the data pipeline, which makes privacy work difficult.** AI systems are complex, and practitioners (4/35) found it difficult to reason about the downstream privacy implications of data they collected in light of the varied data policies, systems, and AI/ML models that are utilized in their organizations and the interactions therein. As a researcher (P35) working on machine learning building tools stated: *"we collect something about use case three [that] could impact use case two or one... Because there's just so many different moving parts... Any one person probably doesn't know every single model that's happening."* A product director (P34) for a job matching tool further explained how these complex data flows made it difficult to assess the privacy risks of AI products: *"The technology is often really sophisticated, and so sometimes the data is leaving your AWS account, sometimes it's not. All kinds of AI and policies control, like who can and can't see that data... And so it becomes difficult [to] tease out the true risk."*

**Practitioners lack guidance and need to rely on individual judgment in their privacy work.** We found our participants (9/35) often found it challenging to navigate and interpret privacy requirements in the context of the specific products they were developing and their specific uses of AI in those products. This difficulty stemmed from the fact that privacy regulations are complicated (P1, P4, P11), and the techniques and procedures used to assure compliance for AI products were too new to have established knowledge in the organization (P6, P16). As P16 noted: *"I'm more doing computer vision stuff. It's pretty new, and so not a lot of people have the answer... it kind of comes down to making my own [decision], and to know what's going to be good, or risk compliance issues."* Some participants (3/35) discussed the insufficiency of blanket privacy requirements to be directly applied to the context of AI, and highlighted the need for individual judgment. As noted by a technical lead (P8): *"there are probably many features that you may consider innocuous... but actually may contain some privacy-related things. For example... zip code can be proxy for race, because... people of a certain race live in a certain neighborhood... sometimes it's actually... super hard to build a system that satisfies all the kinds of tricky privacy requirements."* Additionally, our participants (3/35) found it challenging to evaluate the ef-

fectiveness of their privacy practices in AI products, in part because privacy practices in industry are generally opaque: “I don’t know if we’re doing good. I don’t know if we’re doing bad... I’d have no clue. I have no visibility into what other people are doing in the space” (P16).

Indeed, lacking clear guidance creates obstacles to promoting privacy and security practices in software development [29, 54]. Partially owing to this lack of guidance and reliance on individual judgment, participants described engaging in ad-hoc risk assessments for their AI products: for example, shirking privacy work because a product is a derivative of another product that already went through privacy review. There is a danger that practitioners who are not trained in privacy may underestimate risk without guidance. For example, even if a product is “derivative” of another, it can entail secondary use intrusions. In sum, our findings further reveal that AI can exacerbate these issues in an organization, because AI development contexts are relatively new with few established best practices and standards for privacy.

## 5 Discussion

In summary, for **RQ1**, we found that practitioners viewed privacy as protecting users against pre-defined intrusions that *could* be exacerbated by AI, but that they were not fully aware of how the capabilities and requirements of AI related to their privacy work. For **RQ2**, we found that practitioners faced more inhibitors than motivators for AI privacy work; thus, many practitioners approached privacy work beyond meeting minimum compliance requirements with trepidation. For **RQ3**, we found that the myriad tools and resources that practitioners utilized in their privacy work were typically non-product and non-AI specific, hampering their ability to do AI-specific privacy work. Practitioners felt ill-equipped to handle privacy work and desired more product- and AI-specific privacy guidance. We next synthesize promising avenues to better enhance industry practitioners’ awareness of, motivation to act on, and ability to address AI-exacerbated privacy intrusions when developing consumer-facing AI products, adapting prior successful approaches documented in the SPAF [21].

### 5.1 Improving practitioners’ awareness of AI-exacerbated privacy threats

We found that practitioners remain largely unaware of privacy threats either newly introduced or exacerbated by incorporating AI technologies into consumer products and services. Part of the reason for this low awareness may be because, as our participants discussed, existing educational resources for privacy are overly general and contain little AI-specific information. Building off prior efforts at addressing the awareness barrier in usable privacy as described in the SPAF paper [21], we envision two potential methods to improve practitioner

awareness of AI-exacerbated privacy threats: AI-specific privacy educational materials and simulated attacks.

*AI-specific training campaigns:* Awareness campaigns are commonly used to increase awareness of privacy and security threats more generally [21]. Similar campaigns may be effective at raising practitioners’ awareness of how AI technologies might change the landscape of privacy threats for a consumer product or service by, e.g., mapping the unique capabilities (e.g., prediction, profiling) and data requirements (e.g., curation of personally identifiable data) of AI onto the privacy risks those capabilities and requirements could entail (e.g., non-consensual identification of users, secondary use). Future work can take a holistic view of how the capabilities and requirements of AI might affect privacy work, and help practitioners identify these AI-exacerbated risks.

*Simulated attacks of design concepts:* Simulated attacks can create teachable moments where individuals are more receptive to privacy and security training [21]. For example, red teaming exercises simulate how adversaries might break or compromise a product to proactively identify and fix vulnerabilities and have already been adapted in AI contexts to reduce harms [27]. By forcing practitioners to articulate how they intend to use AI in their consumer product or service, as well as the data requirements to unlock that functionality, structured red teaming exercises may help them become aware of the potential privacy intrusions and the downstream consequences/harms (e.g., [16]) they must address.

### 5.2 Improving practitioners’ motivation to address AI-exacerbated privacy threats

We found that practitioners face more inhibitors than motivators for AI privacy work, especially for work that goes beyond meeting minimum compliance requirements. Many participants felt disengaged from privacy decision-making, expressing a sense of resignation and even resentment about the specific actions they needed to take for privacy when privacy was not uniformly valued by their team. Indeed, some participants were actively discouraged from doing privacy work beyond the minimum (Section 4.2.2). Building off methods to improve motivation in usable privacy more generally [21], we envision two broad approaches for improving practitioners’ motivation to engage in AI privacy work: pro-social design and transformational games.

*Pro-social design:* Prior work has shown how developers’ security and privacy practices are largely socially driven [62], and that pro-social design can improve security and privacy practices more generally [61]. Indeed, our participants expressed a desire for indexed repositories of privacy best practices in the context of AI both within and outside their organizations (Section 4.3.1). A shared repository of best practices for AI privacy work, with examples of how product concepts changed before and after applying the practice, and social proof that others value privacy work could increase motiva-

tion in turn.

*Transformational games:* Practitioners had trouble prioritizing privacy work because its benefits were abstract but its costs were concrete. Transformational games — which aim to change players beyond the experience of the game itself [8] — have been shown to be effective educational tools for abstract concepts in security and privacy [21]. One example is Hacked Time [14], which allows users to travel between two distinct points in time — the past and present — to help game players observe how action or inaction in the past can lead to good or bad security outcomes in the future. When incorporated into mandatory training programs, educational games that help practitioners understand how proactive privacy work can avoid media firestorms, unforeseen surveillance, or illicit uses of generative AI may help motivate privacy work.

### 5.3 Improving practitioners’ ability to address AI-exacerbated privacy threats

A recent meta review on the manifestation of values in AI ethics [65] suggests that few tools have been created to help practitioners in AI-specific privacy work. Our findings provide empirical evidence to support this claim. The practitioners we interviewed were ill-equipped for the privacy work expected of them and did not have tools to help mitigate privacy threats and design needs introduced by AI. Rather, practitioners were required to adapt generic tools and methods, despite wanting more AI- and product-specific privacy guidance. This mismatch hampered practitioners’ ability to evaluate the effectiveness of their privacy practices (Section 4.3.2), and sometimes led them to view their privacy work as a confusing hindrance relative to other, more well-understood goals like improving model performance (Section 4.2.2). Building off prior work in ethical AI and usable privacy and security, we envision the need for AI-specific privacy design assessment and design tools.

For example, work in human-centered AI has already begun to explore approaches like checklists [28, 32, 39], value cards [40], and impact assessments [25, 49] that help practitioners explicitly foreground principles like fairness, accountability, and transparency in the design of AI systems. Extant resources for ethical AI like the Google PAIR guidebook [45], however, discuss privacy generically and do not specifically map the capabilities and data requirements of AI technologies to the threats they create and exacerbate. AI-specific privacy developer tools, checklists, and other such artifacts should help address the ability barriers our participants described.

### 5.4 Addressing all barriers with a more human-centered design process

One of the key upshots of the SPAF paper was that *all three* barriers must be addressed to change behavior [21]. Thus,

there is a need for integrative approaches that address awareness, motivation, and ability together. We envision the creation of a turnkey, human-centered design process for improving practitioners’ awareness of AI-specific privacy threats, motivation to address those threats, and ability to address those threats. Compliance mandates can then be more fluid and process-based [42], rather than rigid and inhibiting, by requiring practitioners to engage in and report on this process.

One approach that is promising includes methods to explore the utility versus intrusiveness of competing design concepts for consumer-facing AI products, as demonstrated in prior work [24]. Our own findings suggest that practitioners often view privacy as coming into conflict with other design goals, such as model performance. Thus, early in the needs-finding process, one can imagine presenting a range of potential design concepts, in storyboards or as low-fidelity prototypes, that highlight use case-specific tensions between privacy and other design goals with respect to the appropriateness/suitability and data flow [42]. For example, in some storyboards, model performance should be weighted higher than privacy and vice versa. Do stakeholders feel more strongly about privacy or about model utility in the scenarios presented? Do different stakeholders feel differently?

With this empirical data in hand, we hypothesize that: (i) practitioners will become more *aware* of the privacy risks their use of AI may entail, (ii) practitioners will feel more *motivated* to take ownership over what privacy means for their products, still guided by what regulation and policy require, and (iii) practitioners will have the *ability* to assess how effectively they are addressing the privacy harms uniquely entailed or exacerbated by their use of AI in those designs.

## 6 Conclusion

We interviewed 35 industry practitioners who develop consumer-facing AI products to model how they defined and scoped AI privacy work, what motivated and inhibited this work, as well as what affected their ability to do their work. Practitioners showed limited awareness of AI-exacerbated privacy threats, faced more inhibitors than motivators to go beyond minimum compliance requirements, and had few tools and resources that provided AI-specific privacy guidance. We also found that while regulatory compliance is still necessary and helpful to get practitioners to prioritize privacy at all, this compliance-centered approach can inhibit formative design explorations of what privacy should mean for a specific product when AI privacy work is operationalized as generic and outcome-based, rather than AI-specific and process-based. To that end, our work suggests that there is a strong need for more turnkey design tools and artifacts that help practitioners address awareness, motivation, and ability barriers to AI privacy work when developing consumer-facing AI products.



## Acknowledgments

This work was generously supported by the National Science Foundation through SaTC grants 2126058, 2126066, and 2316768. We thank the members of the SPUD lab who provided valuable feedback throughout our research. We also thank our anonymous reviewers for their feedback.

## References

- [1] Yasemin Acar, Sascha Fahl, and Michelle L Mazurek. You are not your developer, either: A research agenda for usable security and privacy research beyond end users. *2016 IEEE Cybersecurity Development (SecDev)*, pages 3–8, 2016.
- [2] Ifeoma Ajunwa, Kate Crawford, and Jason Schultz. Limitless worker surveillance. *Calif. L. Rev.*, 105:735, 2017.
- [3] Saleema Amershi, Andrew Begel, Christian Bird, Robert DeLine, Harald Gall, Ece Kamar, Nachiappan Nagappan, Besmira Nushi, and Thomas Zimmermann. Software engineering for machine learning: A case study. In *2019 IEEE/ACM 41st International Conference on Software Engineering: Software Engineering in Practice (ICSE-SEIP)*, pages 291–300. IEEE, 2019.
- [4] Saleema Amershi, Dan Weld, Mihaela Vorvoreanu, Adam Fourney, Besmira Nushi, Penny Collisson, Jina Suh, Shamsi Iqbal, Paul N Bennett, Kori Inkpen, et al. Guidelines for human-ai interaction. In *Proceedings of the 2019 chi conference on human factors in computing systems*, pages 1–13, 2019.
- [5] P Anand and M Bergen. Big teacher is watching: How ai spyware took over schools. bloomberg, 2021.
- [6] Apple. *Human Interface Guidelines — Machine Learning*. 2020.
- [7] Rebecca Balebako, Abigail Marsh, Jialiu Lin, Jason I Hong, and Lorrie Faith Cranor. The privacy and security behaviors of smartphone app developers. 2014.
- [8] Sasha A Barab, Melissa Gresalfi, and Adam Ingram-Goble. Transformational play: Using games to position person, content, and context. *Educational researcher*, 39(7):525–536, 2010.
- [9] Solon Barocas and Helen Nissenbaum. Big data’s end run around procedural privacy protections. *Communications of the ACM*, 57(11):31–33, 2014.
- [10] Emily M Bender and Batya Friedman. Data statements for natural language processing: Toward mitigating system bias and enabling better science. *Transactions of the Association for Computational Linguistics*, 6:587–604, 2018.
- [11] Emily M Bender, Timnit Gebru, Angelina McMillan-Major, and Shmargaret Shmitchell. On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM conference on fairness, accountability, and transparency*, pages 610–623, 2021.
- [12] M. Burgess. *The Biggest Deepfake Abuse Site Is Growing in Disturbing Ways*. WIRED. 2021.
- [13] Nicholas Carlini, Florian Tramer, Eric Wallace, Matthew Jagielski, Ariel Herbert-Voss, Katherine Lee, Adam Roberts, Tom Brown, Dawn Song, Ulfar Erlingsson, et al. Extracting training data from large language models. In *30th USENIX Security Symposium (USENIX Security 21)*, pages 2633–2650, 2021.
- [14] Tianying Chen, Margot Stewart, Zhiyu Bai, Eileen Chen, Laura Dabbish, and Jessica Hammer. Hacked time: Design and evaluation of a self-efficacy based cybersecurity game. In *Proceedings of the 2020 ACM Designing Interactive Systems Conference*, pages 1737–1749, 2020.
- [15] Ben C. F. Choi, Zhenhui (Jack) Jiang, Bo Xiao, and Sung S. Kim. Embarrassing Exposures in Online Social Networks: An Integrated Perspective of Privacy Invasion and Relationship Bonding. *Information Systems Research*, 26(4):675–694, 2015. Publisher: INFORMS.
- [16] Danielle Keats Citron and Daniel J Solove. Privacy harms. *BUL Rev.*, 102:793, 2022.
- [17] Juliet Corbin and Anselm Strauss. Basics of qualitative research. 3rd edn thousand oaks, 2008.
- [18] Henriette Cramer, Jenn Wortman Vaughan, Ken Holstein, Hanna Wallach, Jean Garcia-Gathright, Hal Daumé III, Miroslav Dudík, and Sravana Reddy. Challenges of incorporating algorithmic fairness into industry practice. *FAT\* Tutorial*, 2019.
- [19] Kate Crawford, Meredith Whittaker, Madeleine Clare Elish, Solon Barocas, Aaron Plasek, and Kadija Ferryman. The ai now report. *The Social and Economic Implications of Artificial Intelligence Technologies in the Near-Term*, 2016.
- [20] W Cunrui, Q Zhang, W Liu, Y Liu, and L Miao. Facial feature discovery for ethnicity recognition. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, 9(2):e1278, 2019.
- [21] Sauvik Das, Cori Faklaris, Jason I Hong, Laura A Dabbish, et al. The security & privacy acceptance framework (spaf). *Foundations and Trends® in Privacy and Security*, 5(1-2):1–143, 2022.
- [22] Paul Dourish, Rebecca E Grinter, Jessica Delgado De La Flor, and Melissa Joseph. Security in the wild: user

- strategies for managing security as an everyday, practical problem. *Personal and Ubiquitous Computing*, 8:391–401, 2004.
- [23] Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. Calibrating noise to sensitivity in private data analysis. In *Theory of cryptography conference*, pages 265–284. Springer, 2006.
  - [24] Sindhu Kiranmai Ernal, Stephanie S. Yang, Yuxi Wu, Rachel Chen, Kristen Wells, and Sauvik Das. Exploring the utility versus intrusiveness of dynamic audience selection on facebook. *Proc. ACM Hum.-Comput. Interact.*, 5(CSCW2), oct 2021.
  - [25] Casey Fiesler and Nicholas Proferes. “participant” perceptions of twitter research ethics. *Social Media+ Society*, 4(1):2056305118763366, 2018.
  - [26] Jessica Fjeld, Nele Achten, Hannah Hilligoss, Adam Nagy, and Madhulika Srikumar. Principled artificial intelligence: Mapping consensus in ethical and rights-based approaches to principles for ai. *Berkman Klein Center Research Publication*, (2020-1), 2020.
  - [27] Deep Ganguli, Liane Lovitt, Jackson Kernion, Amanda Askell, Yuntao Bai, Saurav Kadavath, Ben Mann, Ethan Perez, Nicholas Schiefer, Kamal Ndousse, et al. Red teaming language models to reduce harms: Methods, scaling behaviors, and lessons learned. *arXiv preprint arXiv:2209.07858*, 2022.
  - [28] Timnit Gebru, Jamie Morgenstern, Briana Vecchione, Jennifer Wortman Vaughan, Hanna Wallach, Hal Daumé Iii, and Kate Crawford. Datasheets for datasets. *Communications of the ACM*, 64(12):86–92, 2021.
  - [29] Marco Guttfleisch, Jan H Klemmer, Niklas Busch, Yasemin Acar, M Angela Sasse, and Sascha Fahl. How does usable security (not) end up in software products? results from a qualitative interview study. In *43rd IEEE Symposium on Security and Privacy, IEEE S&P*, pages 22–26, 2022.
  - [30] Kashmir Hill. The secretive company that might end privacy as we know it. In *Ethics of Data and Analytics*, pages 170–177. Auerbach Publications, 2020.
  - [31] Sarah Holland, Ahmed Hosny, Sarah Newman, Joshua Joseph, and Kasia Chmielinski. The dataset nutrition label. *Data Protection and Privacy, Volume 12: Data Protection and Democracy*, 12:1, 2020.
  - [32] Kenneth Holstein, Jennifer Wortman Vaughan, Hal Daumé III, Miro Dudik, and Hanna Wallach. Improving fairness in machine learning systems: What do industry practitioners need? In *Proceedings of the 2019 CHI conference on human factors in computing systems*, pages 1–16, 2019.
  - [33] Anna Jobin, Marcello Ienca, and Effy Vayena. The global landscape of ai ethics guidelines. *Nature Machine Intelligence*, 1(9):389–399, 2019.
  - [34] Brittany Johnson, Yoonki Song, Emerson Murphy-Hill, and Robert Bowdidge. Why don’t software developers use static analysis tools to find bugs? In *2013 35th International Conference on Software Engineering (ICSE)*, pages 672–681. IEEE, 2013.
  - [35] Fei-Fei Li. How to make ai that’s good for people. *The New York Times*, 7, 2018.
  - [36] Tianshi Li, Yuvraj Agarwal, and Jason I Hong. Coconut: An ide plugin for developing privacy-friendly apps. *Proceedings of the ACM on Interactive, Mobile, Wearable and Ubiquitous Technologies*, 2(4):1–35, 2018.
  - [37] Tianshi Li, Elizabeth Louie, Laura Dabbish, and Jason I Hong. How developers talk about personal data and what it means for user privacy: A case study of a developer forum on reddit. *Proceedings of the ACM on Human-Computer Interaction*, 4(CSCW3):1–28, 2021.
  - [38] Tianshi Li, Kayla Reiman, Yuvraj Agarwal, Lorrie Faith Cranor, and Jason I Hong. Understanding challenges for developers to create accurate privacy nutrition labels. In *CHI Conference on Human Factors in Computing Systems*, pages 1–24, 2022.
  - [39] Michael A Madaio, Luke Stark, Jennifer Wortman Vaughan, and Hanna Wallach. Co-designing checklists to understand organizational challenges and opportunities around fairness in ai. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, pages 1–14, 2020.
  - [40] Margaret Mitchell, Simone Wu, Andrew Zaldivar, Parker Barnes, Lucy Vasserman, Ben Hutchinson, Elena Spitzer, Inioluwa Deborah Raji, and Timnit Gebru. Model cards for model reporting. In *Proceedings of the conference on fairness, accountability, and transparency*, pages 220–229, 2019.
  - [41] Duc Cuong Nguyen, Dominik Wermke, Yasemin Acar, Michael Backes, Charles Weir, and Sascha Fahl. A stitch in time: Supporting android developers in writing secure code. In *Proceedings of the 2017 ACM SIGSAC Conference on Computer and Communications Security*, pages 1065–1077, 2017.
  - [42] Helen Nissenbaum. Privacy as contextual integrity. *Wash. L. Rev.*, 79:119, 2004.
  - [43] Helen Nissenbaum, Sebastian Benthall, Anupam Datta, Michael C Tschantz, and Piot Mardziel. Origin privacy: Protecting privacy in the big-data era. Technical report, NEW YORK UNIVERSITY, 2018.

- [44] Jonathan A Obar. Sunlight alone is not a disinfectant: Consent and the futility of opening big data black boxes (without assistance). *Big Data & Society*, 7(1):2053951720935615, 2020.
- [45] Google PAIR. People + ai guidebook. 2019.
- [46] Nicolas Papernot, Shuang Song, Ilya Mironov, Ananth Raghunathan, Kunal Talwar, and Úlfar Erlingsson. Scalable private learning with pate. In *International Conference on Learning Representations (ICLR)*, 2018.
- [47] Charith Peris, Christophe Dupuy, Jimit Majmudar, Rahul Parikh, Sami Smaili, Richard Zemel, and Rahul Gupta. Privacy in the time of language models. In *Proceedings of the Sixteenth ACM International Conference on Web Search and Data Mining*, pages 1291–1292, 2023.
- [48] Mark O Riedl. Human-centered artificial intelligence and machine learning. *Human Behavior and Emerging Technologies*, 1(1):33–36, 2019.
- [49] Stuart Schechter and Cristian Bravo-Lillo. Using ethical-response surveys to identify sources of disapproval and concern with facebook’s emotional contagion experiment and other controversial studies. 2014.
- [50] Swapneel Sheth, Gail Kaiser, and Walid Maalej. Us and them: a study of privacy requirements across north america, asia, and europe. In *Proceedings of the 36th International Conference on Software Engineering*, pages 859–870, 2014.
- [51] Ben Shneiderman. Bridging the gap between ethics and practice: guidelines for reliable, safe, and trustworthy human-centered ai systems. *ACM Transactions on Interactive Intelligent Systems (TiiS)*, 10(4):1–31, 2020.
- [52] Reza Shokri, Marco Stronati, Congzheng Song, and Vitaly Shmatikov. Membership inference attacks against machine learning models. In *2017 IEEE symposium on security and privacy (SP)*, pages 3–18. IEEE, 2017.
- [53] Daniel J Solove. A taxonomy of privacy. *University of Pennsylvania law review*, pages 477–564, 2006.
- [54] Mohammad Tahaei, Alisa Frik, and Kami Vaniea. Privacy champions in software teams: understanding their motivations, strategies, and challenges. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*, pages 1–15, 2021.
- [55] Mohammad Tahaei, Kami Vaniea, Konstantin Beznosov, and Maria K Wolters. Security notifications in static analysis tools: Developers’ attitudes, comprehension, and ability to act on them. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*, pages 1–17, 2021.
- [56] Mohammad Tahaei, Kami Vaniea, and Naomi Saphra. Understanding privacy-related questions on stack overflow. In *Proceedings of the 2020 CHI conference on human factors in computing systems*, pages 1–14, 2020.
- [57] Michael Veale, Max Van Kleek, and Reuben Binns. Fairness and accountability design needs for algorithmic support in high-stakes public sector decision-making. In *Proceedings of the 2018 chi conference on human factors in computing systems*, pages 1–14, 2018.
- [58] Ryan Webster, Julien Rabin, Loic Simon, and Frederic Jurie. This person (probably) exists. identity membership attacks against gan generated faces. *arXiv preprint arXiv:2107.06018*, 2021.
- [59] Alan FT Winfield and Marina Jirotko. Ethical governance is essential to building trust in robotics and artificial intelligence systems. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 376(2133):20180085, 2018.
- [60] Richmond Y Wong, Deirdre K Mulligan, Ellen Van Wyk, James Pierce, and John Chuang. Eliciting values reflections by engaging privacy futures using design workbooks. *Proceedings of the ACM on Human-Computer Interaction*, 1(CSCW):1–26, 2017.
- [61] Yuxi Wu, W Keith Edwards, and Sauvik Das. Sok: Social cybersecurity. In *2022 IEEE Symposium on Security and Privacy (SP)*, pages 1863–1879. IEEE, 2022.
- [62] Shundan Xiao, Jim Witschey, and Emerson Murphy-Hill. Social influences on secure development tool adoption: why security tools spread. In *Proceedings of the 17th ACM conference on Computer supported cooperative work & social computing*, pages 1095–1106, 2014.
- [63] Jing Xie, Heather Richter Lipford, and Bill Chu. Why do programmers make security errors? In *2011 IEEE symposium on visual languages and human-centric computing (VL/HCC)*, pages 161–164. IEEE, 2011.
- [64] Chao-Han Huck Yang, Sabato Marco Siniscalchi, and Chin-Hui Lee. PATE-AAE: Incorporating Adversarial Autoencoder into Private Aggregation of Teacher Ensembles for Spoken Command Classification. In *Proc. Interspeech 2021*, pages 881–885, 2021.
- [65] Mireia Yurrita, Dave Murray-Rust, Agathe Balayn, and Alessandro Bozzon. Towards a multi-stakeholder value-based assessment framework for algorithmic systems. In *2022 ACM Conference on Fairness, Accountability, and Transparency*, pages 535–563, 2022.
- [66] Sebastian Zimmeck, Rafael Goldstein, and David Baraka. Privacyflash pro: Automating privacy policy generation for mobile apps. In *NDSS*, 2021.



- [67] Yang Zou, Zhikun Zhang, Michael Backes, and Yang Zhang. Privacy analysis of deep learning in the wild: Membership inference attacks against transfer learning. *arXiv preprint arXiv:2009.04872*, 2020.

## A Appendix

Table 3: Codebook for the data analysis.

| <b>RQ1: How well do AI practitioners’ definitions of privacy work reflect awareness of AI-exacerbated privacy threats?</b>                          |                                                                                   |
|-----------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------|
| <b>Definitions of privacy</b>                                                                                                                       | <b>Definitions of privacy (cont.)</b>                                             |
| Surveillance                                                                                                                                        | Exclusion                                                                         |
| Identification                                                                                                                                      | Secondary use                                                                     |
| - <i>when curating dataset</i>                                                                                                                      | Insecurity                                                                        |
| - <i>when training models</i>                                                                                                                       | - <i>proper data access</i>                                                       |
| - <i>after deploying models</i>                                                                                                                     | - <i>protection of stored personal information</i>                                |
| <b>RQ2: What motivates and inhibits privacy work for consumer-facing AI products?</b>                                                               |                                                                                   |
| <b>Privacy motivators</b>                                                                                                                           | <b>Privacy inhibitors</b>                                                         |
| Alignment with business interests                                                                                                                   | Rigid compliance requirements                                                     |
| - <i>competitive differentiation</i>                                                                                                                | Incentives                                                                        |
| - <i>clients’ privacy concerns</i>                                                                                                                  | Power                                                                             |
| Social responsibility                                                                                                                               | Privacy education                                                                 |
| Compliance with privacy regulations                                                                                                                 | External ownership of privacy                                                     |
|                                                                                                                                                     | Opportunity costs and trade-offs                                                  |
|                                                                                                                                                     | - <i>UX and functionality</i>                                                     |
|                                                                                                                                                     | - <i>business objectives</i>                                                      |
|                                                                                                                                                     | - <i>model performance</i>                                                        |
|                                                                                                                                                     | - <i>development resources</i>                                                    |
| <b>RQ3: What constitutes privacy work for practitioners who develop consumer-facing AI products and what affects their ability to do this work?</b> |                                                                                   |
| <b>Methods, tools, artifacts, and resources to support designing for privacy</b>                                                                    | <b>Practitioner desires of product- and AI-specific privacy guidance</b>          |
| Privacy value negotiations                                                                                                                          | knowledge of specific privacy requirements, regulations, and references           |
| Privacy training                                                                                                                                    | checklist/standard for privacy practices                                          |
| Design references and compliance consultations                                                                                                      | <b>Challenges in designing for privacy</b>                                        |
| - <i>referencing design documents</i>                                                                                                               | Lacking a holistic view of the data pipeline                                      |
| - <i>consulting privacy/legal experts</i>                                                                                                           | Lacking guidance                                                                  |
| Developer tools and artifacts                                                                                                                       | - <i>privacy regulations are complicated</i>                                      |
|                                                                                                                                                     | - <i>privacy requirements are insufficient</i>                                    |
|                                                                                                                                                     | - <i>procedures used to assure compliance are new and underdeveloped</i>          |
|                                                                                                                                                     | - <i>to ensure/evaluate the effectiveness of privacy practices is challenging</i> |

Table 4: Participant Demographic.n/a = prefer not to say

| #   | Age | Gender | Application domain of the product                                                                                          | AI technology of the product                                                                                                                     | Role                                                                                                                      | Company Size |
|-----|-----|--------|----------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------------------|--------------|
| P1  | 25  | Male   | Process Mining                                                                                                             | Predictive Analytics                                                                                                                             | Software Engineer                                                                                                         | 25000+       |
| P2  | 26  | Male   | Media & Entertainment                                                                                                      | Information Retrieval, Speech and Voice                                                                                                          | Software Engineer                                                                                                         | 25000+       |
| P3  | 35  | Male   | Retail                                                                                                                     | Computer Vision/Image Analysis, Recommender Systems                                                                                              | Data Scientist, Software Engineer                                                                                         | 1000-4999    |
| P4  | 26  | Female | Holiday Rental Search                                                                                                      | Recommender Systems                                                                                                                              | Designer, Researcher                                                                                                      | 1- 9         |
| P5  | 26  | Male   | Marketing                                                                                                                  | Predictive Analytics, Recommender Systems                                                                                                        | Software Engineer                                                                                                         | 25000+       |
| P6  | 42  | Male   | Enterprise Software                                                                                                        | Conversational AI/Chatbots                                                                                                                       | Product Manager, Researcher, Technical Lead/Manager                                                                       | not sure     |
| P7  | 23  | Female | Defense/Military, Education, Public Sector, General-purpose ML Tools, Healthcare                                           | Information Retrieval, Predictive Analytics, Recommender Systems, Decision Support                                                               | Designer                                                                                                                  | 250-999      |
| P8  | 35  | Male   | General-purpose ML Tools, Financial: Other, Retail, Financial: Lending/Mortgage                                            | Decision Support, NLP, Predictive Analytics                                                                                                      | Technical Lead/Manager                                                                                                    | 10-49        |
| P9  | 26  | Female | Healthcare, Hiring/Recruiting, Financial: Lending/Mortgage, Financial: Other, Media & Entertainment, Public transportation | Conversational AI/Chatbots, NLP, Speech and Voice                                                                                                | Designer, Domain/Content Expert                                                                                           | 250-999      |
| P10 | n/a | n/a    | General-purpose ML Tools                                                                                                   | Computer Vision/Image Analysis, Recommender Systems                                                                                              | Software Engineer                                                                                                         | 25000+       |
| P11 | n/a | n/a    | Marketing, Media & Entertainment                                                                                           | Recommender Systems, User Modeling/Adaptive Hypermedia                                                                                           | Technical Lead/Manager                                                                                                    | 5000-24999   |
| P12 | 25  | Female | Healthcare                                                                                                                 | Computer Vision/Image Analysis                                                                                                                   | Product Manager, Researcher, Software Engineer                                                                            | 1- 9         |
| P13 | 40  | Female | Healthcare, Financial: Other                                                                                               | Computer Vision/Image Analysis, NLP, Information Retrieval, Predictive Analytics, Conversational AI/Chatbots, Decision support, Speech and Voice | Designer, Researcher, UX Design Director                                                                                  | 1000-4999    |
| P14 | 28  | Male   | Education, Healthcare                                                                                                      | Conversational AI/Chatbots, NLP, Speech and Voice                                                                                                | Product Manager, Project/Program Manager                                                                                  | 1- 9         |
| P15 | 37  | Male   | Marketing                                                                                                                  | Recommender Systems                                                                                                                              | Software Engineer                                                                                                         | 1000-4999    |
| P16 | 25  | Female | General-purpose ML Tools                                                                                                   | Computer Vision/Image Analysis                                                                                                                   | Researcher                                                                                                                | 25000+       |
| P17 | 45  | Female | Hiring/Recruiting, Utilities                                                                                               | Conversational AI/Chatbots, Predictive Analytics, Recommender Systems                                                                            | Researcher                                                                                                                | 50-249       |
| P18 | 31  | Male   | Healthcare                                                                                                                 | n/a                                                                                                                                              | Data Scientist, Designer, Software Engineer, Researcher                                                                   | 1000-4999    |
| P19 | 33  | Male   | Media & Entertainment                                                                                                      | Decision Support, NLP, Recommender Systems, Speech and Voice                                                                                     | Designer, Software Engineer                                                                                               | 25000+       |
| P20 | n/a | Female | Education, Healthcare, Hiring/Recruiting, Financial: Lending/Mortgage,                                                     | Conversational AI/Chatbots, NLP, Decision Support, Predictive Analytics, Text-based Clustering, Information Retrieval, Recommender Systems       | Product Manager, Designer                                                                                                 | 5000-24999   |
| P21 | n/a | n/a    | Financial: Other, Telecom, IT Services                                                                                     |                                                                                                                                                  | Designer                                                                                                                  |              |
| P22 | 48  | Female | Public Sector, Retail, Automotive, Insurance,                                                                              |                                                                                                                                                  | Researcher                                                                                                                |              |
| P23 | 30  | Female | Managed Service Providers,                                                                                                 |                                                                                                                                                  | Designer                                                                                                                  |              |
| P24 | n/a | n/a    |                                                                                                                            |                                                                                                                                                  | Researcher                                                                                                                |              |
| P25 | 30  | Male   | Hiring/Recruiting                                                                                                          | Decision support, NLP                                                                                                                            | Software Engineer, Product Owner                                                                                          | 25000+       |
| P26 | 35  | Male   | Cybersecurity                                                                                                              | Predictive Analytics                                                                                                                             | Researcher                                                                                                                | 1000-4999    |
| P27 | 31  | Female | Media & Entertainment, Advertising                                                                                         | User Modeling/Adaptive Hypermedia                                                                                                                | Researcher                                                                                                                | 250-999      |
| P28 | 40  | Male   | General-purpose ML Tools                                                                                                   | Conversational AI/Chatbots, Information Retrieval                                                                                                | Data Scientist, Designer, Product Manager, Project/Program Manager, Researcher, Software Engineer, Technical Lead/Manager | 1 - 9        |
| P29 | 26  | Female | Defense/Military, General-purpose ML Tools, Financial: Lending/Mortgage, Public Sector                                     | Decision Support, Predictive Analytics, Recommender Systems                                                                                      | Designer                                                                                                                  | 250-999      |
| P30 | 28  | Male   | General-purpose ML Tools, Media & Entertainment                                                                            | Recommender Systems, Decision Support, Predictive Analytics, Information Retrieval, User Modeling/Adaptive Hypermedia                            | Researcher                                                                                                                | 25000+       |
| P31 | 29  | Female | Cybersecurity                                                                                                              | Conversational AI/Chatbots, NLP, Recommender Systems                                                                                             | Data Scientist, Software Engineer                                                                                         | 25000+       |
| P32 | n/a | n/a    | Media & Entertainment                                                                                                      | Recommender Systems                                                                                                                              | Researcher, Software Engineer                                                                                             | 5000-24999   |
| P33 | n/a | Female | IT Services                                                                                                                | Conversational AI/Chatbots, NLP                                                                                                                  | Researcher                                                                                                                | 50-249       |
| P34 | 38  | Male   | Healthcare, Hiring/Recruiting                                                                                              | Conversational AI/Chatbots, Decision support                                                                                                     | Executive/General Manager, Technical Lead/Manager                                                                         | 50-249       |
| P35 | 28  | Male   | Education, General-purpose ML Tools, Consumer Technology                                                                   | Computer Vision/Image Analysis, NLP, Conversational AI/Chatbots, Speech and Voice                                                                | Designer, Researcher                                                                                                      | 25000+       |