

Regulating Emotion AI in the United States: Insights from Empirical Inquiry

Alexis Shore Ingber
University of Michigan
Ann Arbor, USA
ingber@umich.edu

Nazanin Andalibi
University of Michigan
Ann Arbor, USA
andalibi@umich.edu

Abstract

Emotion AI is increasingly deployed in high-stakes contexts, yet U.S. regulation does not adequately protect emotion data. Moreover, there is no representative evidence on public attitudes toward the regulation of emotion AI, including *how* it should be regulated and *who* should be responsible for its regulation. We present the first nationally representative survey on perceptions of emotion AI regulation. Findings suggest the U.S. public 1) believes emotion data should be regulated using existing data governance models and that emotion data should not be used as courtroom evidence; 2) supports banning emotion AI across contexts; and 3) emphasizes the responsibility of government collaboration with data subjects and experts. Our analysis demonstrates perceived ownership of emotion data, perceived accuracy and bias of emotion AI systems, and individual identity significantly shape support for banning emotion AI. We argue for a critical and collaborative approach to emotion AI regulation that recognizes the unique nature of emotion AI as compared to non-emotion AI systems.

CCS Concepts

• **Human-centered computing** → *Empirical studies in HCI*; • **Social and professional topics** → **Computing / technology policy**; **Government technology policy**.

Keywords

emotion recognition, affective computing, AI governance, technology policy, AI policy, Internet law

ACM Reference Format:

Alexis Shore Ingber and Nazanin Andalibi. 2025. Regulating Emotion AI in the United States: Insights from Empirical Inquiry. In *The 2025 ACM Conference on Fairness, Accountability, and Transparency (FAccT '25)*, June 23–26, 2025, Athens, Greece. ACM, New York, NY, USA, 14 pages. <https://doi.org/10.1145/3715275.3732014>

1 Introduction

Emotion artificial intelligence (emotion AI) promises to infer emotions based on data such as facial expression, language, body movements, or physiological signals [73]. Advocates claim this technology can improve qualities such as well being, safety, productivity, and consumer experiences [13]. With these promises, emotion AI has been deployed across high-stakes contexts [40, 74]. For example,

one AI hiring firm claimed they were able to predict job-seekers' likelihood to leave a job earlier than desired (i.e., "job hopping") based on emotion AI inferences [47]. Claimed benefits have precipitated a market for emotion AI expected to hit \$42.9 billion by 2027 [70]. That said, emotion AI has received widespread critique both with respect to its technical accuracy and scientific validity [9] and ethics [50, 105], including inherent issues with relying on technical classifiers for something as subjective as emotion, which people may not even know how to describe in their own words [3, 49]. Indeed, critics oppose claims that emotion AI can know humans better than they know themselves [91]. Emotion AI's invasion upon and usage of intimate data [3], as well as its potential undue influence on people subjected to it, have caused concerns about privacy and free expression, prompting a push for emotion AI governance [7, 103] globally [49] and in the United States (U.S.) [7].¹ In fact, the U.S. public has reported to be generally uncomfortable with the use of emotion AI across contexts [4].

Given the diverse data sources upon which emotion AI can build inferences, some of which have special legal protections (i.e., health, biometrics, children's data), this emerging technology has proven difficult for policymakers to address. U.S. legal scholar Jennifer Bard [7] reviews the complicated conceptualization of emotion AI given the variety of ways it *could* be regulated depending on *how* information about emotions and affect, as inferred by emotion AI (i.e., emotion data) is classified in the law. Of note, the European Union (EU) has recently banned emotion AI when used in workplace or education settings *except* when it involves a healthcare or safety need [109]. While the EU AI Act leaves much to be desired in terms of banning emotion recognition [116] neither similar progress nor advocacy around protecting emotion data is present in the U.S.² As U.S. policymakers tackle AI governance, including its discriminatory impact [71, 82], targeted insights on emotion AI is critical to capture its unique nature and implications. This effort builds on that which has fought for regulating facial recognition technology (FRT), an example of AI which can but does not always entail emotion inferences based on facial data [8]. Although FRT bans have been effective in some cities across the U.S., their scope is limited, do not cover non-facial emotion AI, and appear to have lost momentum [97].

¹It is significant to note that we designed this study and wrote this paper prior to changes in the U.S. political landscape which took place in early 2025. It is possible that the federal efforts to protect against AI's harms will no longer be a priority for the U.S. government by the time of this paper's publication.

²Glimmers of potential to regulate emotion AI appear at the state level *if* emotion data were to be classified as biometric data. For example, Illinois's Biometric Information Privacy Act (BIPA) is currently the only state-level legislation in the U.S. which may be able to safeguard individuals from the deployment of emotion AI technologies if only they collect and use biometric data [55].



This work is licensed under a Creative Commons Attribution 4.0 International License. FAccT '25, Athens, Greece

© 2025 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-1482-5/25/06

<https://doi.org/10.1145/3715275.3732014>

To this end, we assess the U.S. public opinion of *how* emotion data should be conceptualized, *what contexts* people believe emotion AI should be restricted, and *who* should be held responsible for its regulation. Public opinion on emotion AI can help inform how elected officials might govern this technology [19], building on more general AI research [119] and that regarding FRT bans [99]. We offer an investigation of 1) public opinion on banning emotion AI across contexts of deployment, taking into consideration relevant identity factors, perceptions of ownership over emotion data, and accuracy and bias of emotion AI systems; and 2) responsibility attribution for emotion AI governance. We also provide empirical support for Bard's [7] review of prospective emotion data conceptualizations, demonstrating how emotion AI regulation could parallel existing data governance models and if emotion data should be admissible in U.S. courts.

We draw on a sample representative of the U.S. population with an oversample of individuals with disabilities, people of color, and gender minorities ($n = 599$), who are more likely to experience emotion AI-inflicted harms [4, 69, 78, 87, 93].

We find the U.S. public conceptualizes the regulation of emotion data similar to health data, "sensitive" data, and manipulative web interfaces. Further, the public does not believe inferences about emotion should be admissible as evidence in court. Our results support a ban of emotion AI across contexts, with such being partially explained by perceptions of ownership over emotion data, perceptions of accuracy and bias of emotion AI, and identity factors. While a large portion of participants attribute responsibility to the government for regulating emotion AI, there is also a substantial degree of support—both quantitatively and qualitatively—for public power as well as interdisciplinary collaboration.

We argue for policymakers to consider how existing law and policy may be adapted to inform emotion AI governance and call for a dedicated focus on emotion AI in light of public opinion. Given the degree of ownership that people perceive over their emotion data, we contend others beyond the data subject should not be able to derive value from emotion data. In addition, we suggest regulators complicate assessing emotion AI for its accuracy and bias; although these factors explain much of the public support for emotion AI's ban across contexts of deployment, these factors may prove impossible to ever address. Finally, we argue the U.S. government should take responsibility for emotion AI through spearheading diverse and representative collaborative efforts.

2 Prior Work

2.1 Getting to Know "You": Emotion AI's Implications

Despite the surge in literature about perceptions of AI (e.g., [72, 80]), less attention has been given to AI claiming to infer emotions. An exception to this is FRT [94, 98, 99]. Selinger & Hartzog [98] argue FRT's detection of emotional states and private thoughts could give way to individuals losing control over their identity, as they could be falsely associated with groups to which they do not identify. Of note, the study of emotion AI should not be limited to FRT. Schneier [97] argues focusing on FRT bans is a mistake given the larger picture regarding ubiquitous surveillance making discriminatory inferences off peoples' personal data. Thus, an expanded

push against emotion AI is important as emotions are markedly sensitive data [3, 74]. Additionally, emotions play a critical role in everyday decision-making, granting emotion AI manipulative power [64, 96]. In fact, Donohue defines "the use of biometric data to identify, analyze, predict, and manipulate a person's beliefs, desires, emotions, cognitive processes, and/or behavior" as *biomanipulation* [33]. Past work explores how people perceive emotion AI across contexts, finding evidence of its harms [3, 27, 50, 75, 92, 93, 95], such as privacy harms and imposed emotional labor when deployed in the workplace [27]. Further, job-seekers perceive emotion AI use in job interviews as unjust [87]. Similar concerns exist regarding emotion AI use in mental healthcare [57] and on social media [3]. This evidence demonstrating emotion AI's harms across contexts, in addition to concerns surrounding its accuracy and validity [9, 73], qualifies an investment into its regulation.

2.2 Regulating Emotion AI in the United States

AI governance in the U.S. has received scholarly and political engagement over the past decade, leading to general guidance [112, 118] and creation of AI-related task forces at the federal level [111]. Notably, federal reports have outlined principles for developing and deploying AI while respecting civil liberties, aligning with U.S. democratic values [118] and recognizing the risks undermining human rights [112]. These efforts have led to progress on federal AI legislation. In September 2024, Senator Markey of Massachusetts introduced the AI Civil Rights Act, which would directly address the sinister side of AI across sectors by prioritizing 1) equity, 2) transparency, and 3) choice [71]. Notably, AI governance which highlights civil rights has also occurred at the state level. For instance, the Colorado AI Act places restrictions on "high risk" systems that "makes, or is a substantial factor in making, a consequential decision" [24]. In addition, California recently signed 18 AI-related bills into law, with implications across industries [107].

Despite efforts—though limited at the time of writing this paper—to effectively restrict AI in the U.S. across high-stakes contexts, existing law does not explicitly recognize emotion AI, with only some local restrictions specific to FRT use by government and police officials [26, 97] that have not been entirely effective [68]. Similarly, New York City Local Law 144, which requires audits on algorithmic hiring tools for bias and may implicate popular emotion AI hiring tools, has been criticized as ineffective given its inability to hold companies accountable. On one hand, the lack of explicit emotion AI policy is surprising given the advocacy for protection of individual thoughts and feelings [15, 54, 65], calls to ban emotion AI due to its questionable theoretical foundations [8, 28, 104, 113], and the Federal Trade Commission's (FTC) recognition of risks associated with profiling based on psychometric data following Cambridge Analytica. On the other hand, legal scholar Jennifer Bard [7] partially attributes this gap in technology policy to the lack of consensus around conceptualizing emotions as data and whether emotions' regulation should be similar to existing relevant data categories [7].

Within her framework, Bard [7] also questions if emotion data should be conceptualized as forensic science, serving as evidence in courtrooms. Although jurors and judges have historically looked to defendants' emotions as indicative of remorse [6] and are trusting in that produced by algorithms (i.e., one that detects lying) [114],

Barrett et al. [9] contend an incorrect inference made about emotions could lead to ill-informed decisions. This resembles concerns about using polygraph data as evidence [62]. If policymakers are to develop restrictions around emotion AI, it would be useful to understand if the public perceives its outputs as valid courtroom evidence. Building on Bard's review [7] and highlighting the importance of public opinion, we ask the following research question:

RQ1: How does the U.S. public think regulators should conceptualize emotion data when regulating emotion AI?

2.3 Context-Dependent Regulation

Public opinion about AI governance is context-dependent [100], which may extend to emotion AI. Quantitative scholarship has revealed public perceptions of FRT differ across contexts of deployment [20, 99]. For example, Engelman et al. [34] found individuals felt more positive about building inferences based on face data in advertising than hiring. Reflecting on these contextual differences, researchers have advocated for context-specific AI governance that considers the variety of stakeholders involved, including the public [12]. This sentiment has extended to emotion AI [51].

This study examines public opinion of regulating emotion AI across contexts which were addressed in the proposed AI Civil Rights Act—to the extent they are relevant to emotion AI—producing direct implications for AI governance in the U.S.: employment (workplace management, hiring); public spaces; education; border control; healthcare [71]. Though not explicitly mentioned in the AI Civil Rights Act, we also include consumer research, cars, social media, and children's toys as deployment contexts that raise implications for deception and insurance discrimination, which are referenced in the Act. Notably, when measuring perceptions about regulating an emerging technology, it is common to ask in terms of a ban [59, 99]. Further, outright "bans" of emotion AI were made by the EU AI Act, confirming the ecological validity of this operationalization [109]. This framing also establishes consistent mental models as opposed to a more ambiguous term like "regulate".

RQ2: How does the U.S. public opinion about banning emotion AI differ across contexts of deployment?

2.3.1 What Shapes Opinion on Emotion AI Bans? Ownership of Emotion Data and Perceptions of Accuracy and Bias. Opinions about banning emotion AI across contexts may depend on individual perceptions. Given its intimacy and sensitivity, individuals could feel a high sense of *psychological ownership* over their emotion data. Psychological ownership refers to the bond individuals create with objects, conceptualizing those objects as "theirs" [11]. The intimate connection individuals may have with their emotions could lead them to believe emotion data should not be handled by others without explicit permission. Of note, legal theorizing about data ownership has been largely struck down given the intangibility and relational nature of data [23]. More recently, however, legal scholars argued data *can* be considered tangible property when the subject is identifiable, ownership is clearly established, and misuse of the data can be determined [44]. We echo the sentiment of legal scholar Margot Kaminski, who argues data ownership is related to personal dignity in that individuals should be able to control the usage of something derived directly from them [56]. We predict:

H1: Perceptions of ownership over emotion data will positively predict support for banning emotion AI across contexts.

Individuals' opinions about banning emotion AI may vary based on their beliefs about the system's accuracy and bias. The accuracy and validity of emotion AI has prompted questions about reproducibility [81] as well as ethics [9, 58]. While some believe accuracy should be improved [38, 77], others note more accurate emotion AI would be more intrusive [43]. In the education context, for instance, classifying children as "bored" without providing an explanation conflicts with the tenets of pedagogy [31]. Additionally, bias is prevalent in emotion AI systems—stemming from a lack of diversity in the training data [1]—along the lines of gender [60, 69], race [69, 93], culture [69] and disability [78, 120]. This concern around bias has been reflected by U.S. policymakers. For example, the FTC addressed AI bias as unfair under Section 5 of the FTC Act in its settlement with Rite Aid over the company's discriminatory use of FRT [10]. It would be fruitful to understand, then, if opinions about emotion AI regulation depend on how well people think it works and if they perceive it to operate equally across individuals.

RQ3: How do perceptions about a) accuracy and b) bias in emotion AI impact the U.S. public's support for banning emotion AI across contexts?

2.3.2 Identity's Role in Attitudes Toward Emotion AI Regulation. Attitudes about AI adoption in the U.S. differ based on gender [53] as well as socio-economic status (SES) and education [16], with people of minoritized gender groups and lower SES and education having more negative views. In terms of race, recent work found individuals with darker skin tones (as proxy for race) were more likely to use FRT, with the authors explaining this finding as a strategy for challenging existing training data biases [83]. As a result of evident identity-based differences in opinions specific to emotion AI [87] and following prior work [63], we will control for gender, race, disability status, education, SES, and age. Additionally, our analysis will control for political orientation [53] and general privacy concerns as these variables have influence on the acceptance of AI in particular contexts [35, 84].

2.4 Attributing Responsibility for Emotion AI Regulation

In the absence of clear legal frameworks, emotion AI governance implicates a range of stakeholders, including those at the individual, organizational and governmental levels [30, 67]. The plethora of stakeholders who may be held accountable for technology has long been an issue in computing, which Nissenbaum defined as "the problem of many hands" [79]. The U.S. government has acknowledged the necessity for a diverse group of stakeholders responsible for AI more broadly through its creation of the National AI Advisory Committee (NAIAC) which is comprised interdisciplinary experts from the private sector, academia, non-profits, and civil society as well as the AI Safety Institute within the National Institute of Standards and Technology (NIST).³

³More information about the AI Safety Institute can be found here: <https://www.nist.gov/aisi>. It is unclear what the future of these federal government efforts will look like in the new Trump administration.

As U.S. lawmakers and civil rights advocates craft AI regulation and best practices, it is important to understand who the public deems responsible for emotion AI. With respect to AI more broadly, while individuals in the U.S. perceive AI as important for government and companies to manage, they do not necessarily trust them to do so [119]. However, another study found individuals trust corporate stakeholders significantly less than policymakers or citizens to govern AI [29]. In addition to governmental actors and technology developers, the public wants to be more involved in what they currently see as limited stake in AI regulation [117]. Focusing on emotion AI, we ask:

RQ4: To whom do the U.S. public attribute responsibility for emotion AI regulation?

3 Methodology

3.1 Participants

We conducted a survey ($n=599$) representative of the U.S. population by age, sex,⁴ race, and political orientation, oversampling for respondents with disabilities, and those of minoritized genders and races who may be more likely to experience (emotion) AI-inflicted harm [27, 71, 110, 118]. Recruitment was done through Prolific, a reliable recruitment platform [86]. We compensated participants \$12/hour. We used Prolific's inclusion feature to ensure participants were over the age of 18 and lived in the U.S., and manually reviewed all responses ($n = 649$). We excluded 50 responses due to failing attention checks⁵ or providing incomplete survey responses. This study was approved by our university's institutional review board. A breakdown of our sample across demographic variables is provided in Table 2 in Appendix A.1.

3.2 Procedure

This paper was part of a larger survey measuring public opinion on emotion AI in the U.S. Since participants may have lacked knowledge on what emotion AI is or their definition may not be consistent with our conceptualization, we provided a definition [27] at the beginning of the survey and as appropriate throughout.⁶ Participants then answered a series of quantitative items, as described below. Following best practices [46], we randomized items to reduce order effects. Finally, we asked participants a single open ended question, described in the section hereafter.

3.3 Measures

All quantitative measures were evaluated on a 7-point Likert scale ranging from 1 ("Strongly disagree") to 7 ("Strongly agree") unless otherwise noted.

⁴Prolific provides screening based on sex, not gender.

⁵We included two attention checks that were presented to participants at a random point during the survey: "To ensure you are paying attention, please select 'Very Uncomfortable' for this statement," and "To ensure you are paying attention, please select 'Strongly agree' for this statement."

⁶Participants were told: "Please keep in mind the following definition of emotion AI: systems that promise to automatically detect your emotions and moods (e.g., stress, boredom, calmness, fear, fatigue, attentiveness, happiness, sadness, disgust, surprise, anger) based on various kinds of data (e.g., facial expressions, voice, text/language, biometric information)".

Emotion Data Conceptualization. We measured this with reference to Bard's [7] classifications of how emotion data may be regulated. We provided participants with the following directions:

"U.S. law provides different levels of protection for different types of data. One of the challenges of regulating emotion AI is placing emotion data (i.e., data about emotions) in an existing category. Please complete the following sentence: I believe emotion data should be protected like [BLANK] by regulators. Note: You don't need to know how these other data types are regulated. We want to know how similar or different you think emotion data is to these other data types."

Participants could select among the following options: thoughts and beliefs, bodily integrity, wiretapping, protected health information, biometric data, "sensitive" information (i.e., personally identifiable data that, if exposed, could lead to harm to an individual or organization) or web interface designs that manipulate users into making decisions they may otherwise not make. Participants could also indicate if they thought emotion AI was unique and needed its own policy or that it did not require any regulation.

Emotion Data Admissibility in Courtrooms. We measured this using a single item that participants rated their level of agreement with: "I think data created by emotion AI (i.e., people's inferred emotions), should be allowed to be used as evidence in U.S. courts" ($M = 2.61$; $SD = 1.74$).

Support for Banning Emotion AI. We measured support for banning emotion AI across the following contexts: job interviews, the workplace, social media, children's toys, consumer research, public spaces, education, border control, and healthcare. This was informed by the AI Civil Rights Act [71], previous literature, and existing systems [13, 14, 57, 73, 75, 76, 108]. We told participants the following: "We are interested in your perspectives about how, if at all, you think emotion AI use should be regulated in the U.S. Within the questions to follow, please indicate your level of agreement with the following policies." With this framing in mind, participants completed the following sentence "Emotion AI use should be banned [in/for]...". Scale items and corresponding reliability metrics are included in Table 1.

Perceived Accuracy of Emotion AI. We measured this through a single item assessing participants' agreement with the following statement: "Emotion AI would make accurate inferences about me" ($M = 3.19$, $SD = 1.66$).

Perceived Bias of Emotion AI. We measured this through a single item assessing participants' level of agreement with the following statement: "Emotion AI would accurately infer everyone's emotions, regardless of their gender, race, or other attributes." ($M = 3.04$, $SD = 1.87$).

Responsibility Attribution. We measured this across U.S. government regulators ($M = 5.41$; $SD = 1.74$), emotion AI developers ($M = 4.55$; $SD = 1.95$), and data subjects (i.e., self regulation) ($M = 6.33$; $SD = 1.08$) using items adapted from [66]. All items are provided in the Appendix. We also measured this qualitatively through asking participants: "Who should decide what restrictions (if any) should be imposed on emotion AI?"

General Privacy Concerns. We measured this using a validated privacy scale [61]. For example, "All things considered, the internet causes serious privacy problems." ($M = 5.42$, $SD = 0.96$; $\alpha = .70$). All items are provided in Appendix A.2.

Perceived Ownership of Emotion Data. We measured this using three items adapted from previous literature [85] including "I feel a personal ownership over data about my emotions", "I feel like I own my emotions", "I feel like my emotions are my own" ($M = 6.24$, $SD = .95$; $\alpha = .82$).

3.4 Limitations

This study contains several limitations. First, our procedure was limited in framing emotion AI broadly as opposed to focusing on different data inputs through which emotions could be inferred (e.g., vocal tone, speed of digital speech, FRT, heart rate) as well as emotion AI producing specific emotion labels (e.g., sadness) versus those that infer valence and intensity. It is possible that public opinion on banning or attributing responsibility for emotion AI regulation would vary across these criteria. The public opinion may also vary depending on what specific inferences are being made by the emotion AI. Future research may inquire about different stages of the emotion data lifecycle to gain refined insights for law and policy. Additionally, as the media continues to cover emotion AI's development, future research may build on previous work [99] by understanding if the framing of emotion AI would influence perceptions about regulation. Other methodological approaches would also yield useful insights for emotion AI regulation, such as an interview study with U.S. policymakers or a content analysis of news media related to emotion AI deployment. Another limitation of our methodology is the framing around "banning" emotion AI. While this is a reliable way to measure individuals' attitudes toward regulating emerging technologies [59, 99] future research may validate our findings through asking individuals' perceptions on emotion AI regulation using alternative terminology. Finally, our analysis uses a binary racial and disability classification due to limited sample sizes across specific races and disabilities. Future work may address this through targeted sampling of specific demographic groups of interest.

4 Results

4.1 Conceptualizing Emotion Data in Law and Policy

RQ1 asked how the U.S. public conceptualizes emotion data when considering how to regulate emotion AI. Directly applying Bard's framework [7], participants drew conceptual parallels between the regulation of emotion data and other data types (See Figure 1a).

Most participants reported emotion data should be regulated like protected health information ($n = 458$) or manipulative web design ($n = 447$). Many also expressed emotion data should be regulated akin to information classified as "sensitive" ($n = 384$), thoughts and beliefs ($n = 332$), wiretapping ($n = 329$), or bodily integrity ($n = 301$). A smaller sample felt there would be no need for emotion AI-specific regulations ($n = 278$), and an even smaller group saw emotion as requiring its own distinct policy ($n = 244$). A co-occurrence matrix illustrating patterns in participants' selection of conceptualizations for emotion data regulation is in Appendix A.3.

We also found participants generally did not think emotion data should be admissible as evidence in a court of law, with nearly half

Table 1: Measurement items regarding banning emotion AI

Context	Measurement Items	α	M (SD)
Job interviews	...banned in employment interviews when inferring job-seekers' emotions.	–	5.03 (1.53)
Workplace	...banned in workplaces when inferring workers' emotions.	–	5.03 (1.50)
Social media	...banned in social media inferring social media users' emotions.	–	4.89 (1.53)
Children's toys	...banned in children's toys that infer children's emotions.	–	4.70 (1.57)
Consumer research	...banned in understanding consumer reactions to advertisements.	–	4.69 (1.59)
Cars	...banned in cars when inferring drivers' emotions. ...banned in cars when inferring passengers' emotions.	0.92	4.63 (1.90)
Public spaces	...banned in public entertainment venues like parks and stadiums. ...banned in airport security. ...banned in train/bus/public transportation stations.	0.90	4.60 (1.63)
Education	...banned in educational settings when inferring students' emotions.	–	4.49 (1.59)
Border control	...banned in immigration and border protection.	0.81	4.31 (1.80)
Healthcare	...banned for mental healthcare diagnosis. ...banned for mental healthcare treatment and intervention.	0.92	3.96 (1.67)

of the sample reporting to "Strongly disagree" with this application ($n = 242$).⁷

4.2 Perceptions on Banning Emotion AI Across Contexts

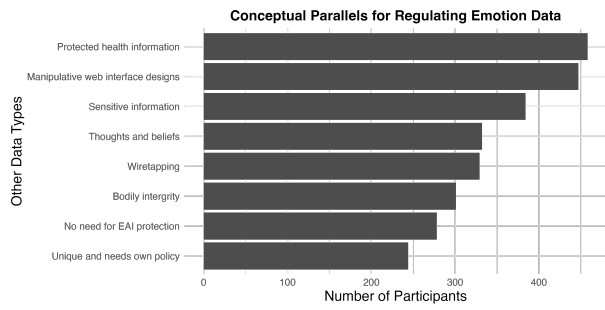
RQ2 investigated how public opinion of banning emotion AI differs across deployment contexts. Descriptively, participants were most supportive of banning emotion AI in work-related contexts (i.e., job interviews and the workplace), whereas this support—while still relatively high—dwindled for its use in healthcare (See Figure 1b).

We sought to understand if opinions about banning emotion AI were dependent on perceptions of ownership over emotion data (H1), accuracy (RQ2a), and bias (RQ2b) of emotion AI, and other individual-level factors. We ran a series of regression models across contexts of deployment, entering the same predictor variables into each model.⁸ We evaluated all necessary assumptions of linear regression a priori.⁹ The results reported hereafter demonstrate

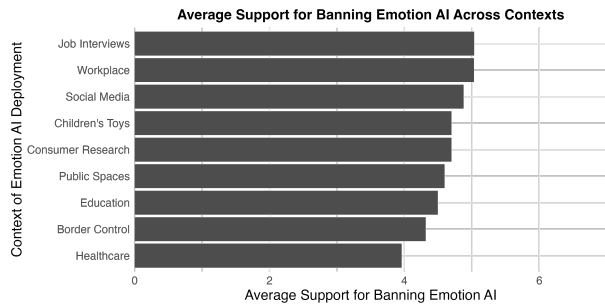
⁷Participants also reported to "Disagree" ($n = 98$), "Somewhat disagree" ($n = 69$), "Neither agree nor disagree" ($n = 91$), "Somewhat agree" ($n = 57$), "Agree" ($n = 23$), or "Strongly agree" ($n = 19$) that emotion data should be allowed to be used as evidence in U.S. courts.

⁸Predictor variables included in the regression models were gender, race, disability status, age, socio-economic status (SES), education, perceived accuracy of emotion AI, perceived bias of emotion AI, perceived ownership over emotion data, and general privacy concerns.

⁹Specifically, we assessed normality of residuals using Q-Q plots and multicollinearity via the variance inflation factor (VIF). To ensure the robustness of our model, we



(a) U.S. public conceptualizations of emotion data regulation compared to other data types



(b) U.S. public support for banning emotion AI across contexts

Figure 1: Descriptive overview of emotion data conceptualization and support for banning emotion AI across contexts

the influence of individual-level variables on support for banning emotion AI in job interviews (Adjusted R^2 : 0.28), the workplace (Adjusted R^2 : 0.32), social media (Adjusted R^2 : 0.23), children's toys (Adjusted R^2 : 0.25), consumer research (Adjusted R^2 : 0.20), cars (Adjusted R^2 : 0.35), public spaces (Adjusted R^2 : 0.35), education (Adjusted R^2 : 0.28), border control (Adjusted R^2 : 0.40), and healthcare (Adjusted R^2 : 0.26). Results reported hereafter include significant findings ($p < .05$) and contexts are reported in order from highest to lowest ban support. Tables detailing these models in full can be found in Appendix B. Some of our findings had very small effect sizes (partial $\eta^2 < .01$) which should be interpreted with caution.

H1 predicted perceptions of ownership over emotion data would be a positive predictor of support for banning emotion AI across contexts. This hypothesis was supported in most of the analyzed contexts except job interviews, border control, or healthcare: workplace ($b = 0.15$, $SE = 0.06$, $p = .006$, partial $\eta^2 = .01$), social media ($b = 0.13$, $SE = 0.06$, $p = .04$, partial $\eta^2 < .01$), children's toys ($b = 0.29$, $SE = 0.06$, $p < .001$, partial $\eta^2 = .03$), consumer research ($b = 0.13$, $SE = 0.07$, $p = .04$, partial $\eta^2 < .01$), cars ($b = 0.24$, $SE = 0.07$, $p = .001$, partial $\eta^2 = .02$), public spaces ($b = 0.22$, $SE = 0.07$, $p = .003$, partial $\eta^2 = .01$), and education ($b = 0.20$, $SE = 0.06$, $p = .001$, partial $\eta^2 = .02$).

RQ2a sought to understand the role of **perceived accuracy of emotion AI** on support for emotion AI bans. When respondents perceived emotion AI to be *less* accurate, they were significantly

identified and removed influential data points using Cook's distance, resulting in a marginally smaller sample size ($n = 566$). Additionally, the Breusch-Pagan test revealed the presence of heteroscedasticity. To correct for this and ensure the accuracy of our statistical inferences, we applied robust standard errors across all models.

more supportive of banning its deployment in all contexts: job interviews ($b = -0.16$, $SE = 0.04$, $p < .001$, partial $\eta^2 = .16$), the workplace ($b = -0.11$, $SE = 0.04$, $p = .01$, partial $\eta^2 = .18$), social media ($b = -0.10$, $SE = 0.05$, $p = .04$, partial $\eta^2 = .13$), children's toys ($b = -0.11$, $SE = 0.05$, $p = .03$, partial $\eta^2 = .15$), consumer research ($b = -0.14$, $SE = 0.05$, $p = .007$, partial $\eta^2 = .13$), cars ($b = -0.34$, $SE = 0.06$, $p < .001$, partial $\eta^2 = .27$), public spaces ($b = -0.32$, $SE = 0.06$, $p < .001$, partial $\eta^2 = .26$), education ($b = -0.21$, $SE = 0.05$, $p < .001$, partial $\eta^2 = .20$), border control ($b = -0.23$, $SE = 0.07$, $p < .001$, partial $\eta^2 = .22$), and healthcare ($b = -0.33$, $SE = 0.05$, $p < .001$, partial $\eta^2 = .21$).

Addressing **RQ2b**, we found **perceptions of bias** were significantly predictive of support for emotion AI bans across all contexts. In other words, when individuals expressed a lower level of agreement that emotion AI systems would accurately infer emotions regardless of gender, race, or other attributes, they had higher support for banning emotion AI across all contexts: workplace ($b = -0.25$, $SE = 0.04$, $p < .001$, partial $\eta^2 = .07$), job interviews ($b = 0.19$, $SE = 0.04$, $p < .001$, partial $\eta^2 = .04$), social media ($b = -0.23$, $SE = 0.04$, $p < .001$, partial $\eta^2 = .05$), children's toys ($b = -0.23$, $SE = 0.04$, $p < .001$, partial $\eta^2 = .05$), consumer research ($b = -0.19$, $SE = 0.05$, $p < .001$, partial $\eta^2 = .03$), cars ($b = -0.22$, $SE = 0.05$, $p < .001$, partial $\eta^2 = .03$), public spaces ($b = -0.26$, $SE = 0.05$, $p < .001$, partial $\eta^2 = .05$), education ($b = -0.23$, $SE = 0.05$, $p < .001$, partial $\eta^2 = .05$), border control ($b = -0.26$, $SE = 0.05$, $p < .001$, partial $\eta^2 = .06$) and healthcare ($b = -0.13$, $SE = 0.05$, $p = .005$, partial $\eta^2 = .02$).

Individual-level variables were predictive of support for banning particular deployments of emotion AI.

Gender. Compared to cisgender men, cisgender women were significantly more likely to support a ban on emotion AI in job interviews ($b = 0.27$, $SE = 0.11$, $p = .02$, partial $\eta^2 = .06$), the workplace ($b = 0.28$, $SE = 0.11$, $p = .01$, partial $\eta^2 = .08$) social media ($b = 0.25$, $SE = 0.12$, $p = .04$, partial $\eta^2 = .07$) and healthcare ($b = 0.36$, $SE = 0.14$, $p = .008$, partial $\eta^2 = .07$).

Race. People of color were significantly less supportive of a ban on emotion AI in the workplace ($b = -0.21$, $SE = 0.11$, $p = .04$, partial $\eta^2 = .02$), social media ($b = -0.31$, $SE = 0.04$, $p < .001$, partial $\eta^2 = .03$), and children's toys ($b = -0.42$, $SE = 0.12$, $p < .001$, partial $\eta^2 = .03$) than white participants.

Disability. Disabled respondents were significantly more supportive of banning emotion AI than those who were not disabled in the context of children's toys ($b = 0.24$, $SE = 0.12$, $p = .04$, partial $\eta^2 < .01$) and consumer research ($b = -0.25$, $SE = 0.13$, $p = .04$, partial $\eta^2 < .01$).

Age. Age was a significant, negative predictor of support for banning emotion AI in public spaces ($b = -0.01$, $SE = 0.00$, $p = .01$, partial $\eta^2 < .01$) and border control ($b = -0.02$, $SE = 0.00$, $p < .001$, partial $\eta^2 = .03$).

SES. There was a positive relationship between SES and support for banning emotion AI in the context of the workplace ($b = 0.08$, $SE = 0.03$, $p = .006$, partial $\eta^2 = .01$).

Political orientation. Compared to those who were more conservative, more liberal participants were significantly more supportive of a ban on emotion AI for use in border control ($b = 0.28$, $SE = 0.05$, $p < .001$, partial $\eta^2 = .14$) and job interviews ($b = 0.14$, $SE = 0.05$, $p = .003$, partial $\eta^2 = .06$).

General privacy concerns. Higher privacy concerns were significantly predictive of support for banning emotion AI in the workplace ($b = 0.11$, $SE = 0.06$, $p = .04$, partial $\eta^2 = .01$), and cars ($b = 0.18$, $SE = 0.06$, $p < .001$, partial $\eta^2 = .02$).

4.3 Attributing Responsibility for Emotion AI Regulation

RQ4 examined to whom the public attributed responsibility for regulating emotion AI. To address this, we conducted quantitative and qualitative analyses. Paired-samples t-tests show respondents attributed a significantly higher level responsibility to the government than developers for regulating emotion AI, $t(564) = 8.29$, $p < .001$. However, respondents also felt *they* should have the authority to influence decisions made based on emotion data, significantly more so than the government ($t(564) = 11.47$, $p < .001$) or developers ($t(564) = 19.54$, $p < .001$).

To account for other attributions of responsibility for regulating emotion AI, we analyzed participants' responses to an open-ended question asking the same. The first author conducted open coding [106] on each response focusing on *who* participants believed should be held responsible for regulating emotion AI. Each response could be coded with multiple codes. The authors discussed cases that were unclear or left room for interpretation to achieve consensus. The authors then grouped codes into larger categories. For example, responses mentioning medical professionals, ethicists, scientists, or human rights advocates were grouped into the "experts" category. We removed 35 responses that did not answer the prompt or provided illegible responses. Our qualitative analysis identified actors that participants deemed relevant in emotion AI regulation: government, the public/data subjects, developers, deployers (e.g., actors who use the technology to evaluate data subjects) or experts.

A large number of participants attributed responsibility to the government ($n = 296$), exemplified by statements such as "The government should decide." However, some recognized the limitations of this attribution, noting, for instance, "most people in Congress are out of touch regarding these new emerging technologies." Conversely, a significant number of participants believed that the public—people subjected to emotion AI—should play a role in regulating its use ($n = 198$). These responses often emphasized the importance of data subjects' consent. As one participant articulated, emotion AI "should only be used with full consent and understanding of the person monitored, if at all."

Some participants thought the responsibility to regulate emotion AI should lie with those developing the technology ($n = 55$) or deciding to deploy it on data subjects ($n = 19$). These participants attributed responsibility to "the creators" or said rules around emotion AI should be "outlined in policies with all applicable private companies," respectively. Additionally, participants mentioned a variety of other experts ($n = 70$) as relevant to emotion AI regulation. Some noted these experts should collaborate ($n = 41$). Within this subsample, several participants viewed the government as part of these collaborative efforts. For example, one respondent suggested collaboration should inform legislative reform: "The restrictions should be a collaborative effort. The creators should find out what people want or need and work from that before connecting with the government about regulation." Another participant noted "the

government needs to listen to the people," implying the public's input should be part of the government's regulatory approach. Others stated collaboration should occur independently of the government, advocating for entities like "a third party ethics board to determine if emotion AI companies are using their products for personal gain or compromising bodily autonomy and sensitive information." Interestingly, some respondents did not attribute responsibility to any actor, arguing instead that emotion AI be completely banned, making the question irrelevant to them ($n = 43$). As one participant stated: "First of all, it shouldn't exist. I don't think anyone should use it, period." Finally, some participants reported to be "not sure," ($n = 25$) in many cases due to feeling insufficiently informed about the technology.

5 Discussion

This study contributes a comprehensive understanding of how the U.S. public conceives of emotion data, how emotion AI should be regulated across contexts, and who should be held responsible for how emotion AI gets used. We discuss the implications of our findings for U.S. policymakers and civil rights advocates.

5.1 How does the U.S. public conceive of "emotion data"?

Our analysis reveals how the U.S. public conceives of emotion data. This is significant, as legal scholar Bard [7] suggests a lack of consensus and clarity on how to conceptualize emotion data is a key challenge in regulating emotion AI. Our results show the U.S. public does not think the government needs to develop a new policy for regulating emotion data; rather, they predominately recognize parallels between emotion data regulation and that of personal health data, "sensitive" data, and manipulative web interfaces. Legal scholar Daniel Solove similarly suggests despite a lack of AI law in the U.S., privacy law could address the harms of AI if "properly conceptualized and constituted" [102].

Many participants felt emotion data exhibited conceptual similarities to health data, which is regulated under the Health Insurance Portability and Accountability Act (HIPAA). However, HIPAA is confined to "covered healthcare entities," meaning it would fail to acknowledge emotion data as health data when used outside of the traditional healthcare system. As a result, HIPAA would insufficiently regulate emotion AI given its deployment in contexts outside of healthcare. Rather than solely protecting data which may be used to infer emotion in HIPAA-covered contexts, we argue U.S. policymakers should give emotion data special protections similar to the GDPR's protection of heart rate, skin conductance, and blood pressure [49] which is not bound to their usage in healthcare. In sum, our findings underscore regulatory measures should acknowledge health data's relevance beyond the traditional healthcare industry. Participants also believed emotion data could be governed similarly to how the law has previously classified other "sensitive" data. As previously mentioned, we defined "sensitive" data to participants as "personally identifiable data that, if exposed, could harm an individual or organization." While sensitivity classifications have been popular in U.S. law and policy (e.g., AB-1008 of California [17]), scholars have recently pushed against them [101]. Quinn & Maltieri argue against sensitivity classifications given

that rapidly advancing computing power may generate data more "sensitive" than what could ever be defined by law [88]. Rather than taking our results as a push for sensitivity classifications, we suggest policymakers recognize emotion data similar to that which is personally identifiable and, if released, could engender harm. These harms include several identified in Citron & Solove's typology [22] such as chilling effects upon recognition of emotion surveillance or discrimination as a result of biased and inaccurate inferences.

Participants also drew a parallel between the regulation of emotion data and other manipulative web interfaces, demonstrating recognition of biomanipulation [33] in that data could be leveraged to manipulate subjects into making decisions they would not otherwise make. Of note, Citron & Solove [22] argue autonomy harms are a privacy harm which occur as a result of people being "directly denied free will to decide or are tricked into thinking that they are freely making choices when they are not." This further supports that the deployment of emotion AI would precipitate a privacy harm which policymakers must recognize. In fact, the use of emotion data for manipulative web design may act more powerfully than traditionally recognized dark patterns [42] given its claim to be personalized to an individuals' emotional state. Our results support a push for the FTC to include emotion AI in their discussion of manipulative dark patterns [25], which would help to address the law's broader inadequacy in addressing pervasive surveillance [48].

Finally, we find the U.S. public does not believe the data generated by emotion AI should be allowed as evidence in criminal investigations. This reflects a lack of public trust toward emotion data as criminal evidence, echoing broader sentiments about the use of any AI-generated outcomes in the courtroom [62]. Despite methodological concerns about algorithmic decision-making in forensic science [62] the American Polygraph Association's recent guidance, which assigns a role to algorithms in polygraph examinations, demonstrates inconsistency between science and practice [2]. Rather than limiting the use of emotion AI inferences, our findings suggest public endorsement of prohibition in the courtroom.

5.2 Context-Dependent Regulation of Emotion AI and Determinants

Overall, participants were supportive of banning emotion AI across contexts. It is notable that the highest support for banning emotion AI was in job interviews and the workplace, contexts which have recently received policy attention in the EU [109]. However, the U.S. public also supports banning emotion AI in children's toys and consumer research, for example, which are currently not present in the proposed AI Civil Rights Act nor other AI governance efforts. That being said, the FTC has taken steps to address deceptive AI practices, filing formal complaints against AI applications designed to mislead consumers [37]. Moreover, an FTC Commissioner recently voiced concerns about AI's role in collecting data from children [41]. While emotion inference has not been specifically captured in these discourses, communications from the FTC signal a gradual expansion of AI governance beyond the most prominent use cases. Our results provide evidence for civil rights advocates and policymakers to push against emotion AI not only in known contexts of deployment, but also those which are emerging.

5.2.1 Sense of Ownership Over Emotion Data. Results suggest people feel a sense of ownership over their emotion data and that this contributes a significant portion of the variance in supporting a ban of emotion AI across most of the analyzed contexts. This sense of ownership was corroborated by our quantitative and qualitative findings around responsibility attribution, where participants indicated a desire to make decisions about their own emotion data. These results provide a theoretical extension of psychological ownership to emotion data (i.e., the ability to create a bond with something inanimate, perceiving it to be "theirs") [11] and contributes to the debate around data as capable of being owned [23]. Notably, previous research has extended psychological ownership to personal data, finding ownership to have a negative relationship with willingness to disclose personal data [21]. Our findings extend this research to emotion data, further highlighting ownership as an influential psychological mechanism over technology use and acceptance. Emotion AI regulation should reflect data subjects' desire to match perceptions of emotion data ownership with reality, ensuring technologies are not intruding upon data perceived to be *theirs* to make decisions about. That ownership was not predictive of banning emotion AI in contexts such as border control and healthcare suggest that perhaps these contexts of deployment felt more justified despite building inferences off of emotion data. This provides evidence of the context-dependent nature of perceptions about emotion AI.

5.2.2 A Need for Diverse Perspectives in Regulating Emotion AI. Our results demonstrate that individual factors shape the U.S. public opinion on emotion AI regulation. Recent research similarly suggests Japanese people's attitudes toward emotion AI across contexts is dependent on variables such as gender, age, and education [52]. These results have broader implications for AI policy and development, highlighting the need to consider how individuals with varying demographic attributes view AI. Additionally, these findings can guide future research on emotion AI by suggesting the influence individual factors have on its acceptance. We highlight the implications of some of these key findings.

Compared to white participants, people of color were less supportive of banning emotion AI across a number of contexts. This aligns with recent research finding people of color to be more interested in engaging with AI [83] and more comfortable with emotion AI deployment [4] than white people. It is possible people of color believe emotion AI systems are less biased than humans as a result of previous discriminatory experiences. In terms of gender, cisgender women were significantly more likely to support a ban of emotion AI than cisgender men in job interviews, the workplace, social media and healthcare. Women are more likely to experience discrimination in hiring and the workplace [5, 32] as well as healthcare [39], which would explain why they would not want algorithms deciding their competence or medical diagnoses, respectively. Cisgender women may also feel they could be manipulated by emotion AI on social media, particularly targeted advertisements which threaten their decision-making freedom [36]. Across disability status, disabled participants were more supportive of banning emotion AI in children's toys and consumer research than participants who were not disabled. Scholars [78, 120] and global policymakers [110] have questioned emotion AI's ability to

accurately infer disabled individuals' emotions. The relationship between identity factors (i.e., gender, race and disability status) and public opinion on banning emotion AI supports continued anti-discriminatory AI policy, which we argue should extend to emotion AI. Further, this evidence bolsters support for the AI Civil Rights Act's discrimination clause [71],¹⁰ as algorithms frequently use non-protected identifiers to make decisions about people [115].

5.3 Accuracy and Bias of Emotion AI Matter to the U.S. Public

Public opinion around banning emotion AI was consistently dependent on perceptions around emotion AI's bias and accuracy. Our analysis suggests if emotion AI were to be deployed, the public would be more supportive if they believed the systems were accurate and not biased across all contexts. We caution against using this finding as motive to develop more accurate emotion AI systems or that accuracy in emotion recognition is a desirable goal. Rather, we urge developers to dissent from this simplistic assertion and question the possibility of ever achieving accuracy of emotion recognition. In practice, policymakers are currently trying to address AI accuracy through auditing requirements [45, 71]. However, audits are both challenging [18, 90] and raise ethical questions, particularly when they concern biometric data [89]. Interestingly, although the NAIAC Charter includes bias as a core responsibility of the committee, accuracy is not noted [111]. Given the importance of emotion AI's bias and accuracy to the American public, policymakers should better understand if emotion AI systems *could* ever be accurate. There is evidence emotion AI will likely never be accurate [9, 74, 105], echoing sentiments that it is impossible for technology to know someone better than they know themselves [91]. It is possible that accuracy is not spotlighted by the NAIAC Charter because they similarly think this is not realistically achievable [111]. Policymakers should use our results as impetus to better understand if the systems they are attempting to regulate could or *should* ever make accurate predictions about human emotion.

5.4 Placing Responsibility for Emotion AI: Looking Forward in Emotion AI Regulation

Although the U.S. public would prefer to leave regulation of emotion AI to the government than technology developers, many participants thought *they*, the subjects of emotion AI, should hold power over their emotion data. Notably, some participants believed it would be fruitful for emotion AI to be governed by a collaborative group of experts. The U.S. government has already taken steps toward this through the NAIAC. That being said, NAIAC members lean toward technical expertise, with only 3 of its 23 members working on ethics in their private sector roles.¹¹ Our findings suggest the public believes regulating emotion AI must include those beyond technical experts such as those who are engaged with the social scientific and ethical aspects of emotion AI use. A comprehensive approach to emotion AI regulation should incorporate diverse

expertise and prioritize public involvement to ensure ethical and socially responsible outcomes.

6 Conclusion

This study provides a comprehensive account of U.S. public opinion on emotion AI. As emotion AI permeates across sectors, policy-makers must consider the nuanced drivers of public opinion. Our findings suggest rather than developing entirely new policies, it could be effective to adapt existing data governance frameworks to encompass emotion data. That being said, the U.S. public is overall supportive of a ban of emotion AI across all contexts we examined. Contextual- and identity-driven differences in this support for regulation highlight the necessity for policymakers and advocates to address the specific concerns and expectations of diverse groups.

Acknowledgments

We are grateful for the reviewers' constructive feedback on this work. We would also like to thank members of the Privacy Research Group at New York University's School of Law for providing feedback on an earlier draft of this paper. This work was sponsored by NSF CAREER award 2236674.

References

- [1] American Bar Association. 2024. The Price of Emotion: Privacy, Manipulation, and Bias in Emotional AI. *Business Law Today* (Sept. 2024). https://www.americanbar.org/groups/business_law/resources/business-law-today/2024-september/price-emotion-privacy-manipulation-bias-emotional-ai/
- [2] American Polygraph Association. 2024. Model Policy for Algorithm Use in Evidentiary Exams. https://www.polygraph.org/policies_and_acts.php. Available as a PDF document accessed from the American Polygraph Association website.
- [3] Nazanin Andalibi and Justin Buss. 2020. The human in emotion recognition on social media: Attitudes, outcomes, risks. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. 1–16.
- [4] Nazanin Andalibi and Alexis Shore Ingber. 2025. Public Perceptions About Emotion AI Use Across Contexts in the United States. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. Association for Computing Machinery. <https://doi.org/10.1145/3706598.3713501>
- [5] Lori Andrews and Hannah Bucher. 2022. Automating Discrimination: AI Hiring Practices and Gender Inequality. *Cardozo L. Rev.* 44 (2022), 145.
- [6] Susan A. Bandes. 2014. Remorse, demeanor, and the consequences of misinterpretation: The limits of law as a window into the soul. *Journal of Law, Religion and State* 3, 2 (2014), 170–199.
- [7] Jennifer S. Bard. 2021. Developing legal framework for regulating emotion AI. *BUJ Sci. & Tech. L.* 27 (2021), 271.
- [8] Lindsey Barrett. 2020. Ban facial recognition technologies for children-and for everyone else. *BUJ Sci. & Tech. L.* 26 (2020), 223.
- [9] Lisa Feldman Barrett, Ralph Adolphs, Stacy Marsella, Aleix Martinez, and Seth D. Pollak. 2019. Emotional Expressions Reconsidered: Challenges to Inferring Emotion From Human Facial Movements. *Psychological science in the public interest: a journal of the American Psychological Society* 20, 1 (2019), 1.
- [10] Alvaro M. Bedoya. 2023. Statement of Commissioner Alvaro M. Bedoya On FTC v. Rite Aid Corporation & Rite Aid Headquarters Corporation. Commission File No. 202-3190. Available at: <https://www.ftc.gov/legal-library/browse/cases-proceedings/public-statements/statement-commissioner-alvaro-m-bedoya-ftc-v-rite-aid-corporation>.
- [11] James K. Beggan. 1992. On the social nature of nonsocial perception: The mere ownership effect. *Journal of personality and social psychology* 62, 2 (1992), 229.
- [12] Teemu Birkstedt, Matti Minkkinen, Anushree Tandon, and Matti Mäntymäki. 2023. AI governance: themes, knowledge gaps and future agendas. *Internet Research* 33, 7 (2023), 133–167.
- [13] Karen L. Boyd and Nazanin Andalibi. 2023. Automated emotion recognition in the workplace: How proposed technologies reveal potential futures of work. *Proceedings of the ACM on human-computer interaction* 7, CSCW1 (2023), 1–37.
- [14] Matt Burgess. 2024. Amazon-Powered AI Cameras Used to Detect Emotions of Unwitting UK Train Passengers. *Wired* (17 June 2024). <https://www.wired.com/story/amazon-ai-cameras-emotions-uk-train-passengers/>. Accessed: 2024-09-05.

¹⁰This clause proposes to protect against technology that "causes a disparate impact" or "discriminates" in "the equal enjoyment of goods, services, or other activities or opportunities, related to a consequential action, on the basis of a protected characteristic."

¹¹An archived list of the NAIAC members can be found here: <https://web.archive.org/web/20250115061539/https://ai.gov/naiac/>

- [15] Federico Cabitza, Andrea Campagner, and Martina Mattioli. 2022. The unbearable (technical) unreliability of automated facial emotion recognition. *Big data & society* 9, 2 (2022), 20539517221129549.
- [16] M Calice, L Bao, T Newman, D Scheufele, D Brossard, and M Xenos. 2022. US Public attitudes on artificial intelligence. *OSF Home* (2022).
- [17] California State Legislature. 2024. California Assembly Bill No. 1008 (2023–2024). https://leginfo.ca.gov/faces/billTextClient.xhtml?bill_id=202320240AB1008 Accessed: 2024-10-25.
- [18] Stephen Casper, Carson Ezell, Charlotte Siegmman, Noam Kolt, Taylor Lynn Curtis, Benjamin Bucknall, Andreas Haupt, Kevin Wei, J  r  my Scheurer, Marius Hobbhahn, et al. 2024. Black-box access is insufficient for rigorous ai audits. In *The 24th ACM Conference on Fairness, Accountability, and Transparency*. 2254–2272.
- [19] Devin Caughey and Christopher Warshaw. 2018. Policy preferences and policy change: Dynamic responsiveness in the American states, 1936–2014. *American Political Science Review* 112, 2 (2018), 249–266.
- [20] Hyesun Choung, Prabu David, and Tsai-Wei Ling. 2024. Acceptance of AI-Powered Facial Recognition Technology in Surveillance Scenarios: Role of Trust, Security, and Privacy Perceptions. *Technology in Society* (2024), 102721.
- [21] Patrick Cichy, Torsten-Oliver Salge, and Rajiv Kohli. 2014. Extending the privacy calculus: the role of psychological ownership. (2014).
- [22] Danielle Keats Citron and Daniel J Solove. 2022. Privacy harms. *Boston University Law Review* 102 (2022), 793.
- [23] Julie E Cohen. 2017. Examined lives: Informational privacy and the subject as object. In *Law and Society Approaches to Cyberspace*. Routledge, 473–538.
- [24] Colorado General Assembly. 2024. Concerning the Implementation of an Automated Traffic Enforcement System. <https://leg.colorado.gov/bills/sb24-205> Accessed: 2024-09-24.
- [25] Federal Trade Commission. 2023. Bringing Dark Patterns to Light. <https://www.ftc.gov/reports/bringing-dark-patterns-light> Accessed: 2024-10-10.
- [26] Kate Conger, Richard Fausset, and Serge F Kovaleski. 2019. San Francisco bans facial recognition technology. *The New York Times* 14, 1 (2019).
- [27] Shanley Corvite, Kat Roemmich, Tillie Ilana Rosenberg, and Nazanin Andalibi. 2023. Data Subjects’ Perspectives on Emotion Artificial Intelligence Use in the Workplace: A Relational Ethics Lens. *Proceedings of the ACM on Human-Computer Interaction* 7, CSCW1 (2023), 1–38.
- [28] Kate Crawford, Roel Dobbe, Theodora Dryer, Genevieve Fried, Ben Green, Elizabeth Kazunas, Amba Kak, Varoon Mathur, Erin McElroy, Andrea Nill S  nchez, Deborah Raji, Joy Lisi Rankin, Rashida Richardson, Jason Schultz, Sarah Myers West, and Meredith Whittaker. 2019. AI Now 2019 Report. https://ainowinstitute.org/AI_Now_2019_Report.html
- [29] Prabu David, Hyesun Choung, and John S Seberger. 2024. Who is responsible? US Public perceptions of AI governance through the lenses of trust and ethics. *Public Understanding of Science* (2024), 09636625231224592.
- [30] Advait Deshpande and Helen Sharp. 2022. Responsible ai systems: who are the stakeholders?. In *Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society*. 227–236.
- [31] Nathalie DiBerardino and Luke Stark. 2023. (Anti-)Intentional Harms: The Conceptual Pitfalls of Emotion AI in Education. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*. 1386–1395.
- [32] Adio-Adet Dinika and Mona Sloane. 2023. AI and Inequality in Hiring and Recruiting: A Field Scan. In *Weizenbaum Conference" AI, Big Data, Social Media, and People on the Move"*. DEU, 1–13.
- [33] Laura K Donohue. 2024. Biomaniplulation. (2024).
- [34] Severin Engelmann, Chiara Ullstein, Orestis Papakyriakopoulos, and Jens Grossklags. 2022. What people think AI should infer from faces. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*. 128–141.
- [35] Pouyan Esmaeilzadeh, Tala Mirzaei, and Spurthy Dharanikota. 2021. Patients’ perceptions toward human–artificial intelligence interaction in health care: experimental study. *Journal of medical Internet research* 23, 11 (2021), e25856.
- [36] Lisa Farman, Maria Leonora Comello, and Jeffrey R Edwards. 2020. Are consumers put off by retargeted ads on social media? Evidence for perceptions of marketing surveillance and decreased ad effectiveness. *Journal of Broadcasting & Electronic Media* 64, 2 (2020), 298–319.
- [37] Federal Trade Commission. 2024. FTC Announces Crackdown on Deceptive AI Claims and Schemes. <https://www.ftc.gov/news-events/news/press-releases/2024/09/ftc-announces-crackdown-deceptive-ai-claims-schemes> Accessed: 2024-12-06.
- [38] Carmen Fern  ndez-Mart  n  ez and Alberto Fern  ndez. 2020. AI and recruiting software: Ethical and legal implications. *Paladyn, Journal of Behavioral Robotics* 11, 1 (2020), 199–216.
- [39] Benedetta Giovanola and Simona Tiribelli. 2023. Beyond bias and discrimination: redefining the AI ethics principle of fairness in healthcare machine-learning algorithms. *AI & society* 38, 2 (2023), 549–563.
- [40] Laurence Goasd  uff. 2018. Emotion AI Will Personalize Interactions. <https://www.gartner.com/smarterwithgartner/emotion-ai-will-personalize-interactions> Accessed: 2024-11-13.
- [41] Jody Godoy. 2024. FTC’s Holyoak concerned AI collecting children’s data. *Reuters* (2024). <https://www.reuters.com/technology/artificial-intelligence/ftcs-holyoak-concerned-ai-collecting-childrens-data-2024-11-14/> Accessed: 2024-12-06.
- [42] Colin M Gray, Nataliia Bielova, Cristiana Santos, and Thomas Mildner. 2024. An Ontology of Dark Patterns: Foundations, Definitions, and a Structure for Transdisciplinary Action. *arXiv preprint arXiv:2309.09640* (2024).
- [43] Gabriel Grill and Nazanin Andalibi. 2022. Attitudes and folk theories of data subjects on transparency and accuracy in emotion recognition. *Proceedings of the ACM on Human-Computer Interaction* 6, CSCW1 (2022), 1–35.
- [44] James Grimmelmann and Christina Mulligan. 2022. Data Property. *Am. UL Rev* 72 (2022), 829.
- [45] Lara Groves, Jacob Metcalf, Alayna Kennedy, Briana Vecchione, and Andrew Strait. 2024. Auditing Work: Exploring the New York City algorithmic bias audit regime. In *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency* (Rio de Janeiro, Brazil) (FAccT ’24). Association for Computing Machinery, New York, NY, USA, 1107–1120. <https://doi.org/10.1145/3630106.3658959>
- [46] RM Groves. 2011. *Survey methodology*. Vol. 198. John Wiley & Sons.
- [47] Karen Hao. 2020. An AI hiring firm says it can predict job hopping based on your interviews.
- [48] Woodrow Hartzog, Evan Selinger, and Johanna Gunawan. 2023. Privacy Nicks: How the Law Normalizes Surveillance. *Washington University Law Review* (2023).
- [49] Andreas H  uselmann, Alan M Sears, Lex Zard, and Eduard Fosch-Villaronga. 2023. EU law and emotion data. In *2023 11th International Conference on Affective Computing and Intelligent Interaction (ACII)*. IEEE, 1–8.
- [50] Javier Hernandez, Josh Lovejoy, Daniel McDuff, Jina Suh, Tim O’Brien, Arathi Sethumadhavan, Gretchen Greene, Rosalind Picard, and Mary Czerwinski. 2021. Guidelines for assessing and minimizing risks of emotion recognition applications. In *2021 9th International conference on affective computing and intelligent interaction (ACII)*. IEEE, 1–8.
- [51] Manh-Tung Ho, Peter Mantello, and Manh-Toan Ho. 2023. An analytical framework for studying attitude towards emotional AI: The three-pronged approach. *MethodsX* 10 (2023), 102149.
- [52] Manh-Tung Ho, Peter Mantello, and Quan-Hoang Vuong. 2024. Emotional AI in education and toys: Investigating moral risk awareness in the acceptance of AI technologies from a cross-sectional survey of the Japanese population. *Heliyon* 10, 16 (2024).
- [53] Michael C Horowitz and Lauren Kahn. 2021. What influences attitudes about artificial intelligence adoption: Evidence from US local officials. *PLoS One* 16, 10 (2021), e0257732.
- [54] Marcello Ienca and Gianclaudio Malgieri. 2022. Mental data protection and the GDPR. *Journal of Law and the Biosciences* 9, 1 (2022), Isac006.
- [55] Illinois General Assembly. 2008. Illinois Biometric Information Privacy Act. https://www.ilga.gov/legislation/ilcs/ilcs3.asp?ActID=3004_740_ILCS_14.
- [56] Margot E Kaminski. 2022. The Case for Data Privacy Rights (or, Please, a Little Optimism). *Notre Dame Law Review Reflection* 97 (2022), 385.
- [57] Nadia Karizat, Alexandra H Vinson, Shobita Parthasarathy, and Nazanin Andalibi. 2024. Patent Applications as Glimpses into the Sociotechnical Imaginary: Ethical Speculation on the Imagined Futures of Emotion AI for Mental Health Monitoring and Detection. *Proceedings of the ACM on Human-Computer Interaction* 8, CSCW1 (2024), 1–43.
- [58] Harmanpreet Kaur, Daniel McDuff, Alex C Williams, Jaime Teevan, and Shamsi T Iqbal. 2022. “I didn’t know I looked angry”: Characterizing observed emotion and reported affect at work. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*. 1–18.
- [59] Kimon Kieslich and Marco L  n  ch. 2024. Regulating AI-Based Remote Biometric Identification. Investigating the Public Demand for Bans, Audits, and Public Database Registrations. In *The 2024 ACM Conference on Fairness, Accountability, and Transparency*. 173–185.
- [60] Lauren Klein and Catherine D’Ignazio. 2024. Data Feminism for AI. In *The 2024 ACM Conference on Fairness, Accountability, and Transparency*. 100–112.
- [61] Bart P Knijnenburg, Alfred Kobsa, and Hongxia Jin. 2013. Dimensionality of information disclosure behavior. *International Journal of Human-Computer Studies* 71, 12 (2013), 1144–1162.
- [62] Kyriakos N Kotsoglou and Alex Biedermann. 2024. Polygraph-based deception detection and Machine Learning. Combining the Worst of Both Worlds? , 100479 pages.
- [63] Sarah Kreps, Julie George, Paul Lushenko, and Adi Rao. 2023. Exploring the artificial intelligence “Trust paradox”: Evidence from a survey experiment in the United States. *Plos one* 18, 7 (2023), e0288109.
- [64] Jennifer S. Lerner. 2003. The role of affect in decision making. In *Handbook of Affective Sciences*, Richard J. Davidson, Klaus R. Scherer, and H. Hill Goldsmith (Eds.). Oxford University Press, 619–642.
- [65] Sjors Ligthart. 2024. Mental privacy as part of the human right to freedom of thought? In *The law and ethics of freedom of thought: Cognitive liberty and privacy*. Palgrave Macmillan.

- [66] Joon Soo Lim. 2017. The third-person effect of online advertising of cosmetic surgery: A path model for predicting restrictive versus corrective actions. *Journalism & Mass Communication Quarterly* 94, 4 (2017), 972–993.
- [67] Qinghua Lu, Liming Zhu, Xiwei Xu, Jon Whittle, Didar Zowghi, and Aurelie Jacquet. 2024. Responsible AI pattern catalogue: A collection of best practices for AI governance and engineering. *Comput. Surveys* 56, 7 (2024), 1–35.
- [68] Douglas MacMillan. 2024. Police in Austin, San Francisco secretly skirted facial recognition bans. *The Washington Post* (May 2024). <https://www.washingtonpost.com/business/2024/05/18/facial-recognition-law-enforcement-austin-san-francisco/#selection-473.0-477.17> Accessed: 2024-05-18.
- [69] Peter Mantello, Manh-Tung Ho, Minh-Hoang Nguyen, and Quan-Hoang Vuong. 2023. Bosses without a heart: socio-demographic and cross-cultural determinants of attitude toward Emotional AI in the workplace. *AI & society* 38, 1 (2023), 97–119.
- [70] MarketsandMarkets. 2024. Emotion Detection and Recognition Market by Component, Technology, Application, and Region - Global Forecast to 2024. <https://www.marketsandmarkets.com/Market-Reports/emotion-detection-recognition-market-23376176.html> Accessed: 2024-09-05.
- [71] Edward Markey. 2024. AI Civil Rights Act. <https://www.markey.senate.gov/view/senator-markey-hosts-press-conference-to-introduce-his-ai-civil-rights-act>
- [72] Kate K Mays, Yiming Lei, Rebecca Giovanetti, and James E Katz. 2022. AI as a boss? A national US survey of predispositions governing comfort with expanded AI roles in society. *AI & society* (2022), 1–14.
- [73] Andrew McStay. 2016. Empathic media and advertising: Industry, policy, legal and citizen perspectives (the case for intimacy). *Big data & society* 3, 2 (2016), 2053951716666868.
- [74] Andrew McStay. 2020. Emotional AI and EdTech: serving the public good? *Learning, Media and Technology* 45, 3 (2020), 270–283.
- [75] Andrew McStay. 2024. *Automating empathy: Decoding technologies that gauge intimate life*. Oxford University Press.
- [76] Andrew McStay and Gilad Rosner. 2021. Emotional artificial intelligence in children's toys and devices: Ethics, governance and practical remedies. *Big Data & society* 8, 1 (2021), 2053951721994877.
- [77] Samaneh Mohammadi, Mohammadreza Mohammadi, Sima Sinaei, Ali Balador, Ehsan Nowroozi, Francesco Flammini, and Mauro Conti. 2023. Balancing privacy and accuracy in federated learning for speech emotion recognition. In *2023 18th Conference on Computer Science and Intelligence Systems (FedCSIS)*. IEEE, 191–199.
- [78] Scott Monteith, Tasha Glenn, John Geddes, Peter C Whybrow, and Michael Bauer. 2022. Commercial use of emotion artificial intelligence (AI): implications for psychiatry. *Current Psychiatry Reports* 24, 3 (2022), 203–211.
- [79] Helen Nissenbaum. 1996. Accountability in a computerized society. *Science and engineering ethics* 2 (1996), 25–42.
- [80] Ekaterina Novozhilova, Kate Mays, and James E Katz. 2024. Looking towards an automated future: US attitudes towards future artificial intelligence instantiations and their effect. *Humanities and Social Sciences Communications* 11, 1 (2024), 1–11.
- [81] Access Now. 2023. Bodily Harms: Mapping the Risks of Emerging Biometric Tech. <https://www.accessnow.org/bodily-harms-mapping-the-risks-of-emerging-biometric-tech> PDF.
- [82] Partnership on Employment & Accessible Technology (PEAT). 2024. AI Inclusive Hiring Framework: Assess Impacts. <https://www.peatworks.org/ai-inclusive-hiring-framework/ten-focus-areas/assess-impacts/> Accessed: 2024-09-25.
- [83] Hannah Overbye-Thompson, Kristy A. Hamilton, and Dana Mastro. 2024. Reinvention mediates impacts of skin tone bias in algorithms: implications for technology diffusion. *Journal of Computer-Mediated Communication* 29, 5 (2024), zmae016. <https://doi.org/10.1093/jcmc/zmae016>
- [84] Matthew R O'Shaughnessy, Daniel S Schiff, Lav R Varshney, Christopher J Rozell, and Mark A Davenport. 2023. What governs attitudes toward artificial intelligence adoption and governance? *Science and Public Policy* 50, 2 (2023), 161–176.
- [85] Joann Peck, Victor A Barger, and Andrea Webb. 2013. In search of a surrogate for touch: The effect of haptic imagery on perceived ownership. *Journal of Consumer Psychology* 23, 2 (2013), 189–196.
- [86] Eyal Peer, Laura Brandimarte, Sonam Samat, and Alessandro Acquisti. 2017. Beyond the Turk: Alternative platforms for crowdsourcing behavioral research. *Journal of Experimental Social Psychology* 70 (2017), 153–163. <https://doi.org/10.1016/j.jesp.2017.01.006>
- [87] Cassidy Pyle, Kat Roemmich, and Nazanin Andalibi. 2024. US Job-Seekers' Organizational Justice Perceptions of Emotion AI-Enabled Asynchronous Interviews. *Proceedings of the ACM on Human-Computer Interaction* (2024).
- [88] Paul Quinn and Gianclaudio Malgieri. 2021. The difficulty of defining sensitive data—The concept of sensitive data in the EU data protection framework. *German Law Journal* 22, 8 (2021), 1583–1612.
- [89] Inioluwa Deborah Raji, Timnit Gebru, Margaret Mitchell, Joy Buolamwini, Joon-seok Lee, and Emily Denton. 2020. Saving face: Investigating the ethical concerns of facial recognition auditing. In *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, 145–151.
- [90] Inioluwa Deborah Raji, Andrew Smart, Rebecca N White, Margaret Mitchell, Timnit Gebru, Ben Hutchinson, Jamila Smith-Loud, Daniel Theron, and Parker Barnes. 2020. Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. In *Proceedings of the 2020 conference on fairness, accountability, and transparency*, 33–44.
- [91] Liran Razinsky. 2023. Better than they know themselves? Algorithms and subjectivity. *Subjectivity* 30, 4 (2023), 394–416.
- [92] Lauren Rhue. 2018. Racial influence on automated perceptions of emotions. Available at SSRN 3281765 (2018).
- [93] Lauren Rhue. 2019. Understanding the hidden bias in emotion-reading AIs.
- [94] Evan Ringel and Amanda Reid. 2023. Regulating Facial Recognition Technology: A Taxonomy of Regulatory Schemata and First Amendment Challenges. *Communication Law and Policy* 28, 1 (2023), 3–46.
- [95] Kat Roemmich and Nazanin Andalibi. 2021. Data subjects' conceptualizations of and attitudes toward automatic emotion recognition-enabled wellbeing interventions on social media. *Proceedings of the ACM on Human-Computer Interaction* 5, CSCW2 (2021), 1–34.
- [96] James A Russell, José-Miguel Fernández-Dols, Anthony SR Manstead, and Jane C Wellenkamp. 2013. *Everyday conceptions of emotion: An introduction to the psychology, anthropology and linguistics of emotion*. Vol. 81. Springer Science & Business Media.
- [97] Bruce Schneier. 2020. We're banning facial recognition. We're missing the point. *The New York Times* 20 (2020).
- [98] Evan Selinger and Woodrow Hartzog. 2020. The inconstitability of facial surveillance. *Loy. L. Rev.* 66 (2020), 33.
- [99] Alexis Shore. 2022. Talking about facial recognition technology: How framing and context influence privacy concerns and support for prohibitive policy. *Telematics and Informatics* 70 (2022), 101815.
- [100] Aaron Smith, Janna Anderson, and Lee Rainie. 2017. Automation in Everyday Life. <https://www.pewresearch.org/internet/2017/10/04/automation-in-everyday-life/> Accessed: 2024-10-05.
- [101] Daniel J Solove. 2023. Data Is What Data Does: Regulating Based on Harm and Risk Instead of Sensitive Data. *Nw. UL Rev.* 118 (2023), 1081.
- [102] Daniel J Solove. 2025. Artificial intelligence and privacy. *Florida Law Review* (2025).
- [103] Luke Stark. 2016. The emotional context of information privacy. *The Information Society* 32, 1 (2016), 14–27.
- [104] Luke Stark. 2018. Facial recognition, emotion and race in animated social media. *First Monday* (2018).
- [105] Luke Stark and Jesse Hoey. 2021. The ethics of emotion in artificial intelligence systems. In *Proceedings of the 2021 ACM conference on fairness, accountability, and transparency*, 782–793.
- [106] Anselm Strauss. 1997. *Grounded theory in practice*. Sage.
- [107] Nisha Talagala. 2024. California's New AI Laws—What You Should Know. *Forbes* (2024). <https://www.forbes.com/sites/nishatalagala/2024/10/14/californias-new-ai-laws-what-you-should-know/> Accessed: 2024-10-25.
- [108] Hannah Tyler. 2022. The increasing use of artificial intelligence in border zones prompts privacy questions. *Migration Policy Institute* (2022).
- [109] European Union. 2024. Artificial Intelligence Act. <https://artificialintelligenceact.eu/recital/44/> Accessed: 2024-09-05.
- [110] United Nations. 2021. Rights of Persons with Disabilities: Report of the Special Rapporteur on the Rights of Persons with Disabilities. <https://www.ohchr.org/en/issues/disability/sr-disability/pages/index.aspx> Accessed: 2024-09-24.
- [111] U.S. Department of Commerce, National Institute of Standards and Technology. 2023. *National Artificial Intelligence Advisory Committee Charter*. Technical Report. U.S. Department of Commerce, National Institute of Standards and Technology. <https://www.caipd.org/resources/naiaac/> Accessed: 2024-11-01.
- [112] U.S. Department of State. 2024. Risk Management Profile for Artificial Intelligence and Human Rights. <https://www.state.gov/risk-management-profile-for-ai-and-human-rights/> Accessed: 2024-10-25.
- [113] Konstantina Vemou, Anna Horvath, and T Zerdick. 2021. Facial emotion recognition. *TechDispatch* 1 (2021), 1–5.
- [114] Alicia von Schenk, Victor Klockmann, Jean-François Bonnefon, Iyad Rahwan, and Nils Köbis. 2024. Lie detection algorithms disrupt the social dynamics of accusation behavior. *Is Science* 27, 7 (2024).
- [115] Sandra Wachter. 2022. The theory of artificial immutability: Protecting algorithmic groups under anti-discrimination law. *Tul. L. Rev.* 97 (2022), 149.
- [116] Sandra Wachter. 2024. Limitations and loopholes in the EU AI Act and AI Liability Directives: what this means for the European Union, the United States, and beyond. *Yale Journal of Law and Technology* 26, 3 (2024).
- [117] Shangrui Wang and Zheng Liang. 2024. What does the public think about artificial intelligence? An investigation of technological frames in different technological context. *Government Information Quarterly* 41, 2 (2024), 101939.
- [118] White House Office of Science and Technology Policy. 2022. Blueprint for an AI Bill of Rights. <https://www.whitehouse.gov/ostp/ai-bill-of-rights/> Accessed: 2024-09-24.

- [119] Baobao Zhang and Allan Dafoe. 2020. US public opinion on the governance of artificial intelligence. In *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*. 187–193.
- [120] Annuska Zolyomi and Jaime Snyder. 2024. An Emotion Translator: Speculative Design By Neurodiverse Dyads. In *Proceedings of the CHI Conference on Human Factors in Computing Systems*. 1–18.

A Appendix

A.1 Participant Demographics

Table 2: Demographics of participant sample

Variable	Identification	Count
Gender	Cisgender Woman	288
	Cisgender Man	212
	Transgender/Nonbinary	99
Race	White	302
	Black or African American	141
	Hispanic or Latino	37
	Mixed Race	76
	Southeast Asian	13
	East Asian	15
	South Asian	11
	American Indian/Alaska Native	3
	Middle Eastern	1
Disability	No Disability	286
	At least one Disability	313
Education	Less than High School	11
	High School or Equivalent	97
	Some College	200
	Bachelor's Degree	181
	Some Graduate School	18
	Master's/Professional Degree	79
	Doctoral Degree	13
Socio-economic status	Less than \$20,000	89
	\$20,000 to \$34,999	85
	\$35,000 to \$49,999	94
	\$50,000 to \$74,999	101
	\$75,000 to \$99,999	80
	\$100,000 to \$149,999	92
	\$150,000 to \$199,999	47
Age	\$200,000 or more	11
	18–24	98
	25–34	164
	35–44	127
	45–54	86
	55–82	124

A.2 Measurement Items

Responsibility Attribution

The following items were measured on a 7-point Likert scale ranging from "Strongly disagree" to "Strongly agree"

Government regulators: Emotion AI should be regulated by the government.

Platform developers: Emotion AI developers should be held responsible for ethical use of emotion AI.

Self-regulation (data subjects): The decision to make inferences about my emotions should be mine.

General Privacy Concerns: The following items were measured on a 7-point Likert scale ranging from "Strongly disagree" to "Strongly agree": All things considered, the internet causes serious privacy problems; Compared to others, I am more sensitive about the way my personal information is handled; To me, it is the most important thing to keep my privacy intact from companies and institutions; I believe other people are too much concerned with online privacy issues; Compared with other things, personal privacy is very important to me

A.3 Co-Occurrence Matrix

Options for conceptualizing emotion data are presented along both the x- and y-axis. Participants were able to select all options that they felt applied to the conceptualization of emotion data. Darker colors indicate a higher association between the two different categories (i.e., a larger number of participants chose *both* options presented on the x- and y-axis). Lighter colors indicate a lower association between two different categories (i.e., fewer people chose *both* the options presented on the x- and y- axis).

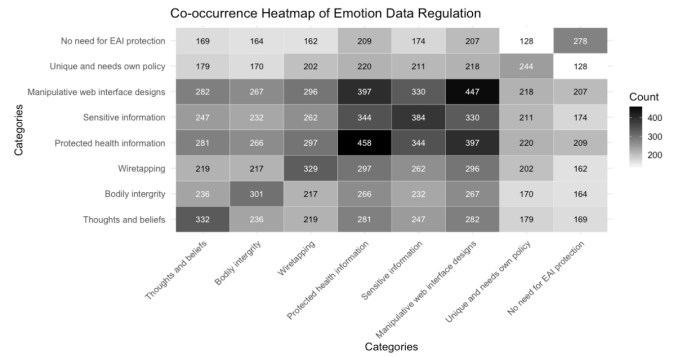


Figure 2: Heatmap of co-occurrence matrix

B Linear Regression Models by Context

Table 3: Support for Banning Emotion AI in Job Interviews

Variable	Estimate	Standard Error	t value
Intercept	5.03***	0.52	9.76
Cisgender Female	0.27*	0.11	2.38
Transgender/Nonbinary	-0.06	0.17	-0.37
Person of Color (POC)	-0.10	0.11	-0.91
Political Orientation	0.14**	0.05	2.96
Perceived Accuracy	-0.16***	0.04	-3.57
Perceived Bias	-0.19***	0.04	-4.63
Disability	-0.20.	0.11	-1.79
Privacy Concerns	0.08	0.06	1.44
Perceived Ownership	0.10.	0.06	1.75
Education Level	-0.07.	0.04	-1.73
Income	0.05.	0.03	1.87
Age	-0.01.	0.00	-1.85

Significance codes: *** $p < .001$, ** $p < .01$, * $p < .05$, . $p < .10$.
Reference groups: Male (Gender), White (POC), and People with a Disability (Disability).

Table 4: Support for Banning Emotion AI in the Workplace

Variable	Estimate	Standard Error	t value
Intercept	4.72***	0.51	9.31
Cisgender Female	0.28*	0.11	2.53
Transgender/Nonbinary	0.17	0.17	1.02
Person of Color (POC)	-0.21*	0.11	-2.04
Political Orientation	0.06	0.05	1.44
Perceived Accuracy	-0.11*	0.04	-2.53
Perceived Bias	-0.25***	0.04	-6.37
Disability	-0.17	0.11	-1.58
Privacy Concerns	0.11*	0.06	2.03
Perceived Ownership	0.15**	0.06	2.72
Education Level	-0.08	0.04	-1.91
Income	0.08**	0.03	2.75
Age	-0.01	0.00	-1.75

Significance codes: *** $p < .001$, ** $p < .01$, * $p < .05$, . $p < .10$.
Reference groups: Male (Gender), White (POC), and People with a Disability (Disability).

Table 5: Support for Banning Emotion AI on Social Media

Variable	Estimate	Standard Error	t value
Intercept	4.74***	0.56	8.53
Cisgender Female	0.25*	0.12	2.03
Transgender/Nonbinary	0.31	0.18	1.71
Person of Color (POC)	-0.31**	0.12	-2.70
Political Orientation	0.02	0.05	0.34
Perceived Accuracy	-0.10*	0.05	-2.03
Perceived Bias	-0.23***	0.04	-5.24
Disability	-0.07	0.12	-0.57
Privacy Concerns	0.10	0.06	1.67
Perceived Ownership	0.13*	0.06	2.04
Education Level	-0.05	0.05	-1.09
Income	0.01	0.03	0.32
Age	-0.00	0.00	-0.19

Significance codes: *** $p < .001$, ** $p < .01$, * $p < .05$, . $p < .10$.
Reference groups: Male (Gender), White (POC), and People with a Disability (Disability).

Table 6: Support for Banning Emotion AI in Children's Toys

Variable	Estimate	Standard Error	t value
Intercept	3.86***	0.57	6.72
Cisgender Female	0.01	0.13	0.06
Transgender/Nonbinary	0.02	0.19	0.12
Person of Color (POC)	-0.42***	0.12	-3.50
Political Orientation	0.01	0.05	0.25
Perceived Accuracy	-0.11*	0.05	-2.23
Perceived Bias	-0.23***	0.05	-5.19
Disability	0.24*	0.12	1.99
Privacy Concerns	0.11	0.06	1.70
Perceived Ownership	0.29***	0.06	4.49
Education Level	-0.01	0.05	-0.24
Income	0.02	0.03	0.67
Age	-0.01*	0.00	-2.38

Significance codes: *** $p < .001$, ** $p < .01$, * $p < .05$, . $p < .10$.
Reference groups: Male (Gender), White (POC), and People with a Disability (Disability).

Table 7: Support for Banning Emotion AI in Consumer Research

Variable	Estimate	Standard Error	t value
Intercept	4.60***	0.60	7.72
Cisgender Female	-0.04	0.13	-0.31
Transgender/Nonbinary	-0.05	0.19	-0.24
Person of Color (POC)	-0.19	0.12	-1.54
Political Orientation	0.00	0.05	0.09
Perceived Accuracy	-0.14**	0.05	-2.70
Perceived Bias	-0.20***	0.05	-4.17
Disability	-0.25*	0.13	-2.01
Privacy Concerns	0.16*	0.07	2.48
Perceived Ownership	0.13*	0.07	2.01
Education Level	-0.11*	0.05	-2.19
Income	0.04	0.03	1.27
Age	-0.00	0.00	-0.29

Significance codes: *** $p < .001$, ** $p < .01$, * $p < .05$, . $p < .10$.
Reference groups: Male (Gender), White (POC), and People with a Disability (Disability).

Table 8: Support for Banning Emotion AI in Public Spaces

Variable	Estimate	Standard Error	t value
Intercept	4.60***	0.67	6.89
Cisgender Female	-0.10	0.15	-0.66
Transgender/Nonbinary	0.31	0.22	1.41
Person of Color (POC)	-0.00	0.14	-0.00
Political Orientation	0.07	0.06	1.10
Perceived Accuracy	-0.32***	0.06	-5.67
Perceived Bias	-0.26***	0.05	-4.86
Disability	-0.07	0.14	-0.50
Privacy Concerns	0.12	0.07	1.68
Perceived Ownership	0.22**	0.07	2.94
Education Level	0.00	0.05	0.08
Income	0.02	0.04	0.57
Age	-0.01*	0.00	-2.34

Significance codes: *** $p < .001$, ** $p < .01$, * $p < .05$, . $p < .10$.
Reference groups: Male (Gender), White (POC), and People with a Disability (Disability).

Table 9: Support for Banning Emotion AI in Education

Variable	Estimate	Standard Error	t value
Intercept	4.69***	0.58	8.08
Cisgender Female	-0.04	0.13	-0.32
Transgender/Nonbinary	0.06	0.19	0.34
Person of Color (POC)	-0.21	0.12	-1.71
Political Orientation	0.01	0.05	0.14
Perceived Accuracy	-0.21***	0.05	-4.28
Perceived Bias	-0.23***	0.05	-4.99
Disability	-0.15	0.12	-1.18
Privacy Concerns	0.04	0.06	0.62
Perceived Ownership	0.20**	0.06	3.17
Education Level	-0.03	0.05	-0.68
Income	0.04	0.03	1.29
Age	-0.00	0.00	-0.84

Significance codes: *** $p < .001$, ** $p < .01$, * $p < .05$, . $p < .10$.
Reference groups: Male (Gender), White (POC), and People with a Disability (Disability).

Table 10: Support for Banning Emotion AI in Border Control

Variable	Estimate	Standard Error	t value
Intercept	5.11***	0.60	8.46
Cisgender Female	0.04	0.13	0.27
Transgender/Nonbinary	0.21	0.20	1.04
Person of Color (POC)	-0.02	0.13	-0.14
Political Orientation	0.29***	0.05	5.31
Perceived Accuracy	-0.23***	0.05	-4.50
Perceived Bias	-0.26***	0.05	-5.48
Disability	-0.21	0.13	-1.61
Privacy Concerns	0.01	0.07	0.08
Perceived Ownership	0.10	0.07	1.49
Education Level	0.00	0.05	-0.01
Income	-0.02	0.03	-0.62
Age	-0.02***	0.00	-4.13

Significance codes: *** $p < .001$, ** $p < .01$, * $p < .05$, . $p < .10$.
Reference groups: Male (Gender), White (POC), and People with a Disability (Disability).

Table 11: Support for Banning Emotion AI in Healthcare

Variable	Estimate	Standard Error	t value
Intercept	4.54***	0.62	7.28
Cisgender Female	0.36**	0.14	2.63
Transgender/Nonbinary	0.34.	0.20	1.68
Person of Color (POC)	-0.15	0.13	-1.13
Political Orientation	-0.02	0.06	-0.36
Perceived Accuracy	-0.33***	0.05	-6.22
Perceived Bias	-0.14**	0.05	-2.81
Disability	-0.10	0.13	-0.78
Privacy Concerns	0.04	0.07	0.61
Perceived Ownership	0.11	0.07	1.56
Education Level	0.00	0.05	-0.01
Income	0.03	0.03	0.99
Age	-0.00	0.01	-0.52

Significance codes: *** $p < .001$, ** $p < .01$, * $p < .05$, . $p < .10$.
Reference groups: Male (Gender), White (POC), and People with a Disability (Disability).

Table 12: Support for Banning Emotion AI in Cars

Variable	Estimate	Standard Error	t Value
Intercept	3.96***	0.66	5.98
Cisgender Female	0.10	0.15	0.66
Transgender/Nonbinary	0.20	0.22	0.91
Person of Color (POC)	-0.10	0.14	-0.75
Political Orientation	0.06	0.06	0.94
Perceived Accuracy	-0.34***	0.06	-6.14
Perceived Bias	-0.22***	0.05	-4.22
Disability	-0.12	0.14	-0.85
Privacy Concerns	0.18*	0.07	2.44
Perceived Ownership	0.24**	0.07	3.27
Education Level	-0.04	0.05	-0.72
Income	0.05	0.04	1.43
Age	-0.00	0.01	-0.96

Significance codes: *** $p < .001$, ** $p < .01$, * $p < .05$, . $p < .10$.
Reference groups: Male (Gender), White (POC), and People with a Disability (Disability).