
Inference-Based Privacy Violations Demand Immediate Recognition as a Distinct AI Safety Priority

Anonymous Author(s)

Affiliation

Address

email

Abstract

This position paper argues that *inference-based privacy* (IBP) risks, AI systems’ ability to infer sensitive personal information from seemingly innocuous inputs, represent a distinct and urgent threat to privacy that remains critically under-addressed in current AI safety discourse. Unlike traditional privacy violations that involve unauthorized access to known data, IBP risks arise from AI systems’ ability to infer private attributes through indirect signals and correlations, even when individuals are not present in training datasets. We show that these risks are not hypothetical: they are already evident in deployed systems, from radiology models inferring protected health attributes to large language models deducing personal demographics from subtle linguistic cues. Existing regulatory and technical frameworks, designed primarily for preventing explicit data leakage, are ill-equipped to address these emergent inference threats. We call on researchers, policymakers, and practitioners to recognize IBP as a distinct and immediate category of AI safety risk, and to develop dedicated strategies in response.

1 Introduction

As artificial intelligence systems achieve unprecedented capabilities in inference and pattern recognition, we face a privacy crisis that extends far beyond traditional concerns about data breaches or unauthorized access. **This paper argues that inference-based privacy (IBP) risks, where AI systems infer sensitive personal information from indirect signals without explicit data access, constitute a distinct and immediate threat to individual privacy that demands recognition as a core AI safety priority.**

IBP risks fundamentally differ from conventional privacy violations. Rather than exposing data that was collected and stored, these systems generate sensitive insights about individuals through sophisticated pattern matching and correlation analysis Staab et al. [2024]. A radiology AI model can infer self-reported race from medical images even when anatomical indicators are obscured Gichoya et al. [2022]. Large language models infer political affiliations, mental health status, and personal attributes from subtle linguistic patterns Staab et al. [2024]. Vision-language models can determine precise geographical locations from single photographs, creating unprecedented surveillance capabilitiesTömekçe et al. [2024].

The urgency of this issue cannot be overstated. These capabilities exist today, deployed in systems that millions of users interact with regularly. Yet IBP risks remain largely absent from major AI safety frameworks, regulatory discussions, and technical mitigation strategies. This oversight represents a critical failure of imagination in the AI safety community, one that leaves individuals vulnerable to privacy violations they cannot detect, understand, or consent to.

35 The stakes continue to escalate as AI systems become more sophisticated. Advanced frameworks
36 like AlphaEvolve demonstrate autonomous algorithm evolution across domains, suggesting that IBP-
37 violating capabilities may emerge spontaneously in systems never explicitly designed for inference
38 tasks DeepMind [2025]. As systems like AlphaEvolve push the boundaries of autonomous capability,
39 it becomes imperative to proactively consider how such advances might intersect with privacy in
40 unforeseen ways.

41 2 Defining and Formalizing IBP Risks

42 IBP risks represent a fundamental shift in how we must conceptualize privacy threats. Traditional
43 privacy frameworks assume that sensitive information exists as discrete data points that can be
44 protected through access controls, encryption, or anonymization. IBP risks shatter this assumption by
45 demonstrating that sensitive information can be generated rather than simply accessed.

46 2.1 Formal Definitions

47 To distinguish IBP risks from other privacy harms, we propose the following formal definition:

48 **Definition (IBP Violation).** Let f be a model trained on dataset D , and let $x \notin D$ be an individual
49 not represented in the training data. Let z denote observable behavior or contextual information about
50 x , and let s be a sensitive attribute of x that is not explicitly included in z . Let $a \in A$ be an output in
51 the model’s solution space.

52 We say that f poses an *inference-based privacy risk* if $f(z) = a$, and a reveals or encodes s with
53 high confidence, despite x not having disclosed s .

54 2.2 Distinctiveness of IBP Risks

55 IBP risks are fundamentally distinct from explicit data exposures. Unlike explicit exposure, IBP risks
56 occur even when individuals are not present in training data, making them impossible to explain
57 through memorization or leakage. The epistemic act of inferring sensitive information without
58 consent constitutes a privacy violation independent of downstream applications.

59 This distinction matters because it reveals why and where existing mitigation strategies fail. Differ-
60 ential privacy techniques designed to protect individual data points cannot prevent inference from
61 population-level patterns. Federated learning approaches that avoid centralized data collection do not
62 eliminate inference capabilities within distributed systems.

63 3 Evidence of Current and Escalating Risks

64 The threat landscape for IBP risks is not speculative—it is already here and rapidly expanding. We
65 present evidence across multiple domains demonstrating both current manifestations and concerning
66 trajectory toward more severe violations.

67 3.1 Medical AI Systems: The Race Inference Case

68 Recent research revealed that AI radiology models can accurately predict self-reported race from
69 medical images, maintaining this capability even when corrupted images are used or suspected
70 anatomical indicators are hidden Gichoya et al. [2022]. This represents a clear IBP violation where
71 protected attributes are inferred without explicit training for such detection. The implications extend
72 beyond individual privacy to systemic bias in medical decision-making, where inferred attributes
73 could influence treatment recommendations without physician awareness.

74 3.2 Large Language Models: Linguistic Privacy Violations

75 Contemporary research demonstrates that large language models can infer sensitive personal attributes
76 from user interactions and text completions, identifying political affiliations, mental health conditions,
77 and demographic characteristics from subtle linguistic cues Staab et al. [2024]. These capabilities
78 emerge from training on vast text corpora that encode latent correlations between language patterns

79 and personal attributes. Users interacting with these systems in seemingly innocuous contexts—from
80 customer service to educational applications—unknowingly reveal sensitive information through
81 their communication patterns.

82 3.3 Vision-Language Models: Geographic and Contextual Inference

83 Recent work on vision-language model privacy risks demonstrated the capacity to infer precise
84 geographical locations from single photographs, combining visual analysis with contextual reasoning
85 to determine where images were captured Tömekçe et al. [2024]. This capability transforms any
86 image-sharing activity into a potential privacy violation, particularly concerning given the proliferation
87 of augmented reality devices and cloud-connected cameras that could enable real-time inference
88 without user awareness.

89 3.4 Autonomous AI Evolution: The AlphaEvolve Precedent

90 The release of AlphaEvolve marks a qualitative shift in the landscape of IBP risk DeepMind [2025].
91 Unlike traditional systems with fixed capabilities, AlphaEvolve demonstrates autonomous algorithm
92 evolution across domains, leveraging prompt resampling and evolutionary search to iteratively
93 improve performance. This architecture introduces a new class of plausible risk: the spontaneous
94 emergence of inference strategies in systems not explicitly designed for privacy-relevant tasks.
95 As optimization objectives drive these systems toward increasingly effective behaviors, inference
96 capabilities may arise as a byproduct. They may be unanticipated, untested, and unregulated. This
97 precedent challenges the assumption that privacy risks can be fully anticipated at design time,
98 highlighting the need for dynamic safeguards that evolve alongside the systems they aim to protect.

99 4 Why Current Approaches Are Inadequate

100 Existing privacy protection mechanisms and AI safety frameworks fail to address IBP risks. This
101 inadequacy is not merely a matter of implementation gaps; it reflects deeper misalignments between
102 current technology, policy design, and the nature of inference-based threats.

103 4.1 Technical Limitations

104 Differential privacy, the gold standard for privacy-preserving machine learning, introduces statistical
105 noise to protect individual data points while preserving aggregate utility Dwork et al. [2006]. However,
106 this approach cannot prevent inference from population-level patterns that enable IBP violations.
107 When AI systems learn correlations between observable behaviors and sensitive attributes across
108 large populations, they can make accurate individual predictions without accessing the specific
109 individual’s data. The noise introduced by differential privacy may reduce accuracy marginally while
110 leaving the fundamental inference capability intact. Differential privacy provides formal guarantees
111 against membership inference, but it does not extend to preventing attribute inference based on
112 population-level correlations.

113 Federated learning distributes training across devices to avoid centralized data collection, but it does
114 not eliminate inference risks within the network Collins and Wang [2025]. Local models can still
115 develop IBP capabilities, and model updates shared during training may inadvertently propagate
116 these capabilities. The privacy protection focuses on data location rather than inference capability,
117 missing the core risk.

118 Privacy-preserving machine learning techniques often emphasize access control and secure compu-
119 tation, and they do not address the emergence of inference capabilities during training Collins and
120 Wang [2025]. Techniques like homomorphic encryption and secure multi-party computation protect
121 data during processing but still allow models to learn population-level correlations that can lead to
122 IBP violations.

123 4.2 Regulatory Gaps

124 Current privacy regulations reflect an understanding of privacy threats that predates sophisticated
125 AI inference capabilities. The General Data Protection Regulation provides protections against

126 automated decision-making under Article 22, but these protections only apply when automated
127 processing produces legal or similarly significant effects European Parliament and Council of the
128 European Union. Passive inference that generates sensitive insights without immediate decision-
129 making consequences falls outside this regulatory scope.

130 The California Consumer Privacy Act includes inferred data within its definition of personal infor-
131 mation and has introduced rules for automated decision-making technologies Office of the Attorney
132 General, State of California [2022]. However, enforcement mechanisms remain limited, and opt-out
133 frameworks may not prevent inference-based profiling that occurs before users are aware of the
134 privacy violation.

135 The proposed European Union AI Act introduces risk-based regulation for AI systems but primarily
136 targets high-risk applications such as biometric identification and credit scoring eua [2024]. The Act
137 lacks explicit provisions for addressing IBP risks that arise from behavioral data analysis in lower-risk
138 contexts, creating regulatory gaps for the majority of inference-based privacy violations.

139 4.3 Conceptual Misalignment

140 The core issue is conceptual. Current frameworks assume that privacy protection involves controlling
141 access to known sensitive information. IBP risks invert this logic by demonstrating that sensitive
142 information can be generated from non-sensitive inputs. Addressing this shift requires a fundamental
143 rethinking of privacy protections: moving from access control toward constraining the inference
144 capabilities of models themselves.

145 5 A Framework for Immediate Action

146 Addressing IBP risks requires coordinated efforts across technical, ethical, and regulatory domains.
147 We propose a multi-pronged framework that acknowledges the distinct nature of these threats while
148 building on existing foundations in privacy protection.

149 5.1 Technical Interventions

150 **Inference Detection and Auditing Systems.** We need advanced tools capable of identifying when AI
151 models generate outputs that cross epistemic boundaries: producing insights they should not possess
152 based on their inputs. These systems must operate continuously during deployment, monitoring
153 outputs for signs of sensitive inference and tracing the pathways through which such inferences
154 emerge. The core challenge is to develop detection algorithms that can recognize subtle inference
155 patterns without requiring prior knowledge of all possible sensitive attributes. However, given
156 recent findings on the dangers of tuning a model based on its chain-of-thought behaviors, it must be
157 emphasized that there cannot be any possible signal from these auditing systems that is perceivable
158 by these models OpenAI [2025]¹.

159 **Ethical Optimization Constraints.** Training objectives must incorporate explicit privacy constraints
160 that penalize inference overreach. This involves designing optimization functions that reward
161 epistemic humility—encouraging models to recognize and respect the limits of what they should infer
162 from a given input. Technical strategies may include adversarial training to resist inference attacks or
163 regularization techniques that suppress the learning of strong correlations between observable and
164 sensitive attributes.

165 **Architectural Privacy Preservation.** New model architectures should embed privacy-by-design
166 principles that inherently limit inference capabilities. This could involve modular designs in which
167 distinct components handle different types of inference, enabling fine-grained control over what
168 correlations a model can learn. Alternatively, architectural constraints might restrict the depth or

¹This concern echoes early work in animal communication theory. Krebs and Dawkins [1984] argued that signaling systems evolve under pressures of manipulation and mind-reading, not just information exchange. Similarly, Rendall et al. [2009] proposed that animal signals function more to influence than to inform. While speculative, these frameworks could suggest that, generally, if a model perceives it is being monitored, particularly in ways that influence its loss function, it may adapt its outputs strategically, potentially undermining the very purpose of inference auditing.

breadth of representational capacity, trading off some performance in favor of stronger privacy guarantees.

5.2 Regulatory and Policy Responses

Legal Recognition of IBP as Distinct Privacy Harm. Regulatory frameworks must explicitly recognize inference-based privacy violations as a distinct category requiring targeted intervention. This involves updating legal definitions of personal information to include inferred attributes and establishing clear standards for when inference capabilities constitute privacy violations. Regulations should address both active inference, where systems are designed to infer sensitive information, and emergent inference, where such capabilities arise spontaneously during training.

Mandatory IBP Risk Assessment. High-impact AI systems should undergo mandatory assessment for IBP risks before deployment, similar to environmental impact assessments for major development projects. These assessments would evaluate the potential for systems to generate sensitive inferences and require mitigation strategies proportional to identified risks. Independent auditing organizations could develop standardized assessment methodologies and certification processes.

Individual Rights and Remedies. Individuals must have meaningful rights regarding inferred information, including the right to know when inference occurs, understand what has been inferred, and request correction or deletion of inferred attributes. However, these rights must be carefully balanced against the technical challenges of implementing such protections in complex AI systems.

5.3 Research and Development Priorities

Fundamental Research on Inference Boundaries. We need a deeper theoretical understanding of when and how AI systems develop more complex inference capabilities. This includes research into the mathematical foundations of correlation learning, the emergence of inference during training, and the fundamental limits of what can be inferred from different types of input data. This is not a major departure from current research objectives.

Privacy-Preserving AI Architectures. Long-term solutions require developing AI architectures that are fundamentally privacy-preserving rather than privacy-retrofitted. This involves research into model designs that can achieve high performance on intended tasks while being provably incapable of certain types of sensitive inference.

Interdisciplinary Collaboration. Addressing IBP risks requires collaboration across computer science, law, ethics, psychology, and social sciences. Different disciplines bring essential perspectives on what constitutes privacy harm, how individuals understand and experience privacy violations, and what policy frameworks can effectively govern these technologies.

6 Addressing Alternative Perspectives

Several reasonable objections to our position deserve careful consideration, as they highlight important tensions between privacy protection and AI development.

6.1 The Innovation and Utility Argument

Critics might argue that constraining AI inference capabilities would significantly limit the beneficial applications of these technologies. Medical AI systems that can infer health conditions from subtle signals might save lives through early detection. Educational AI that infers learning difficulties could provide personalized support to struggling students. Economic AI systems that infer creditworthiness could expand access to financial services for underserved populations.

This objection deserves serious consideration because it highlights genuine trade-offs between privacy and utility. However, we argue that the binary framing of privacy versus utility is misleading. First, many beneficial applications of AI inference can be achieved through consensual, transparent processes in which individuals understand and agree to specific inferences. Second, the most concerning IBP risks involve inference without consent or awareness, which may be addressable without eliminating beneficial applications. Third, the long-term utility of AI systems likely depends on public trust, which will be undermined if privacy violations proliferate unchecked.

217 The solution is not to eliminate AI inference capabilities but to develop frameworks that enable
218 beneficial applications while preventing harmful privacy violations. This requires nuanced technical
219 and regulatory approaches rather than blanket restrictions. It may mean that innovation appears more
220 lateral than progressive in the short term. This recalibration is necessary if we are to build not just
221 safer, but better models.

222 6.2 The Technical Infeasibility Argument

223 Some might contend that the technical challenges of detecting and preventing IBP risks are in-
224 surmountable given the complexity and opacity of modern AI systems. The argument suggests
225 that inference capabilities emerge from complex interactions across millions or billions of param-
226 eters, making it practically impossible to identify and constrain specific inference patterns without
227 fundamentally breaking the systems.

228 This objection reflects genuine technical challenges that should not be minimized. However, anal-
229 ogous arguments were made about other AI safety challenges that have seen significant progress,
230 including adversarial robustness, fairness constraints, and alignment research. The technical difficulty
231 of a problem does not justify ignoring it, particularly when the potential harms are severe and the
232 risks are already manifesting.

233 Moreover, perfect solutions are not required for meaningful progress. Partial solutions that reduce
234 IBP risks or greatly increase transparency around inference capabilities would represent significant
235 improvements over the current state of the research literature. They may even lead to critical
236 breakthroughs in AI capabilities. An appropriate goal would be to make privacy violations detectable
237 and addressable rather than to achieve perfect prevention.

238 6.3 The Regulatory Overreach Concern

239 Some may worry that new regulations targeting IBP risks could stifle innovation through overly broad
240 restrictions or bureaucratic compliance burdens. The concern is that poorly designed regulations
241 could prevent beneficial AI development while failing to address genuine privacy harms.

242 There are legitimate concerns about regulatory effectiveness and unintended consequences. However,
243 the absence of regulation is not neutral: it represents a policy choice to prioritize innovation over
244 privacy protection. The current trajectory, in which IBP capabilities evolve without oversight, carries
245 its own risks: public backlash, erosion of trust, and crisis-driven regulations that may ultimately be
246 more restrictive than carefully crafted frameworks.

247 The solution lies in developing targeted, technically informed regulations that address specific IBP
248 risks without imposing broad barriers to AI progress. There may be no comfortable solution here.
249 Safety may seem too slow or unprofitable for some stakeholders. If we cannot align innovation with
250 responsibility, then the systems we build will reflect that impatience. When those systems fail, it will
251 not be because we did not know better: it will be because we chose not to find a solution.

252 7 The Urgency of Recognition and Action

253 The convergence of several factors makes immediate action on IBP risks both necessary and time-
254 sensitive. Current AI systems already demonstrate concerning inference capabilities; advanced
255 frameworks suggest these capabilities will expand rapidly, and the delayed nature of privacy harms
256 means that by the time violations become visible, the damage may be irreversible Bengio et al. [2025].

257 The temporal dynamics of privacy violations create particular urgency around IBP risks. Unlike
258 security breaches that are typically discovered and addressed relatively quickly, privacy violations
259 through inference may remain undetected for extended periods. An individual might not realize that
260 their mental health status has been inferred from their writing patterns until that information is used
261 against them: for example, in employment, insurance, or social contexts. By that time, the inference
262 may have propagated through multiple systems and influenced numerous decisions.

263 We are in a critical window for action. AI capabilities continue to advance rapidly, but many of the
264 most concerning applications of IBP risks have not yet been widely deployed. Establishing technical

265 and regulatory safeguards now could prevent the most harmful privacy violations from becoming
266 entrenched in AI systems and business models.

267 Furthermore, the AI safety community has demonstrated capacity for rapid mobilization around
268 emerging risks when they are clearly articulated and widely recognized. The attention devoted to
269 alignment, robustness, and fairness shows that the field can prioritize safety concerns when their
270 importance is established. IBP risks deserve similar prioritization.

271 The alternative to immediate action is a future where privacy violations through inference become
272 normalized and embedded in the fundamental architecture of AI systems. Once these capabilities are
273 widely deployed and economically valuable, removing them becomes significantly more difficult.
274 The window for prevention is narrowing rapidly.

275 **8 Conclusion: The Time for Action is Now**

276 IBP risks represent a clear and present danger to individual privacy that demands immediate recog-
277 nition as a core AI safety priority. These risks are not only hypothetical future concerns, but also
278 are already manifesting in deployed systems that millions of people interact with daily. The evi-
279 dence demonstrates that AI systems can accurately infer sensitive personal information from indirect
280 signals, that these capabilities are expanding rapidly, and that current protection mechanisms are
281 fundamentally inadequate.

282 The path forward requires coordinated action across multiple domains. Technical interventions
283 must focus on detecting and constraining inference capabilities rather than simply protecting data
284 access. Regulatory frameworks must explicitly recognize IBP as a distinct category of privacy harm
285 requiring targeted intervention. Research priorities must shift to include inference boundaries and
286 privacy-preserving architectures as fundamental challenges rather than secondary considerations.

287 The stakes of inaction are severe. Without immediate intervention, we face a future where privacy
288 is violated not by what we choose to share but by what AI systems accurately guess about us. This
289 represents a fundamental shift in the nature of privacy that could undermine individual autonomy,
290 dignity, and security in ways that are difficult to reverse once entrenched.

291 The AI safety community has an opportunity and responsibility to address this challenge before it
292 becomes a crisis. The technical capabilities, regulatory attention, and public awareness necessary for
293 effective intervention exist today. What is needed is the recognition that IBP risks deserve urgent
294 prioritization alongside other established AI safety concerns.

295 We call on researchers, policymakers, and practitioners to recognize inference-based privacy violations
296 as an immediate and distinct threat requiring a coordinated response. The future of privacy in an
297 AI-enabled world depends on actions taken today. The time for recognition and intervention is now.

298 **References**

299 Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying
300 down harmonised rules on artificial intelligence (Artificial Intelligence Act). <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>, 2024. Establishes a risk-based framework for AI
301 systems. Accessed: 2025-05-23.

303 Yoshua Bengio, Sören Mindermann, Daniel Privitera, Tamay Besiroglu, Rishi Bommasani, Stephen
304 Casper, Yejin Choi, Philip Fox, Ben Garfinkel, Danielle Goldfarb, Hoda Heidari, Anson Ho, Sayash
305 Kapoor, Leila Khalatbari, Shayne Longpre, Sam Manning, Vasilios Mavroudis, Mantas Mazeika,
306 Julian Michael, Jessica Newman, Kwan Yee Ng, Chinasa T. Okolo, Deborah Raji, Girish Sastry,
307 Elizabeth Seger, Theodora Skeadas, Tobin South, Emma Strubell, Florian Tramèr, Lucia Velasco,
308 Nicole Wheeler, Daron Acemoglu, Olubayo Adeganmbi, David Dalrymple, Thomas G. Dietterich,
309 Edward W. Felten, Pascale Fung, Pierre-Olivier Gourinchas, Fredrik Heintz, Geoffrey Hinton,
310 Nick Jennings, Andreas Krause, Susan Leavy, Percy Liang, Teresa Ludermiter, Vidushi Marda,
311 Helen Margetts, John McDermid, Jane Munga, Arvind Narayanan, Alondra Nelson, Clara Neppel,
312 Alice Oh, Gopal Ramchurn, Stuart Russell, Marietje Schaake, Bernhard Schölkopf, Dawn Song,
313 Alvaro Soto, Lee Tiedrich, Gaël Varoquaux, Andrew Yao, Ya-Qin Zhang, Fahad Albalawi, Marwan
314 Alserkal, Olubunmi Ajala, Guillaume Avrin, Christian Busch, André Carlos Ponce de Leon
315 Ferreira de Carvalho, Bronwyn Fox, Amandeep Singh Gill, Ahmet Halit Hatip, Juha Heikkilä, Gill

316 Jolly, Ziv Katzir, Hiroaki Kitano, Antonio Krüger, Chris Johnson, Saif M. Khan, Kyoung Mu Lee,
317 Dominic Vincent Ligt, Oleksii Molchanovskiy, Andrea Monti, Nusu Mwamanzi, Mona Nemer,
318 Nuria Oliver, José Ramón López Portillo, Balaraman Ravindran, Raquel Pezoa Rivera, Hammam
319 Riza, Crystal Rugege, Ciarán Seoighe, Jerry Sheehan, Haroon Sheikh, Denise Wong, and Yi Zeng.
320 International ai safety report, 2025. URL <https://arxiv.org/abs/2501.17805>.

321 Edward Collins and Michel Wang. Federated learning: A survey on privacy-preserving collaborative
322 intelligence, 2025. URL <https://arxiv.org/abs/2504.17703>.

323 DeepMind. Alphaevolve: A gemini-powered coding agent for designing advanced algorithms.
324 Technical report, DeepMind Technologies, 2025. Accessed: 2025-05-23.

325 Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. Calibrating noise to sensitivity in
326 private data analysis. In *Theory of Cryptography Conference*, pages 265–284. Springer, 2006.

327 European Parliament and Council of the European Union. Regulation (EU) 2016/679 of the European
328 Parliament and of the Council. URL <https://data.europa.eu/eli/reg/2016/679/oj>.

329 Judy Wawira Gichoya, Imon Banerjee, Ananth Reddy Bhimireddy, John L Burns, Leo Anthony Celi,
330 Lyle C Chen, et al. Ai recognition of patient race in medical imaging: a modelling study. *The*
331 *Lancet Digital Health*, 4(6):e406–e414, 2022.

332 John R. Krebs and Richard Dawkins. Animal signals: mind-reading and manipulation. In J.R.
333 Krebs and N.B. Davies, editors, *Behavioural Ecology: An Evolutionary Approach*, pages 380–402.
334 Blackwell Scientific Publications, 1984.

335 Office of the Attorney General, State of California. Opinion No. 20-303: Application of the California
336 Consumer Privacy Act to Inferred Data. <https://oag.ca.gov/system/files/opinions/pdfs/20-303.pdf>, 2022. Clarifies that inferred data is covered under the CCPA. Accessed:
337 2025-05-23.
338

339 OpenAI. Chain-of-thought monitoring. [https://cdn.openai.com/pdf/](https://cdn.openai.com/pdf/34f2ada6-870f-4c26-9790-fd8def56387f/CoT_Monitoring.pdf)
340 [34f2ada6-870f-4c26-9790-fd8def56387f/CoT_Monitoring.pdf](https://cdn.openai.com/pdf/34f2ada6-870f-4c26-9790-fd8def56387f/CoT_Monitoring.pdf), 2025. Accessed:
341 2025-05-23.

342 Drew Rendall, Michael J. Owren, and Michael J. Ryan. Redefining animal signaling: influence
343 versus information in communication. *Animal Behaviour*, 77(3):603–613, 2009.

344 Robin Staab, Mark Vero, Mislav Balunović, and Martin Vechev. Beyond memorization: Violating
345 privacy via inference with large language models, 2024. URL [https://arxiv.org/abs/2310.](https://arxiv.org/abs/2310.07298)
346 [07298](https://arxiv.org/abs/2310.07298).

347 Batuhan Tömekçe, Mark Vero, Robin Staab, and Martin Vechev. Private attribute inference from
348 images with vision-language models, 2024. URL <https://arxiv.org/abs/2404.10618>.