

Privacy Risks of Corporate Data Collection in the Age of Agentic Artificial Intelligence (AI): A Qualitative Document Analysis of Security and Individual Autonomy

Sharon L. Burton 

Embry-Riddle Aeronautical University, USA
burtons6@erau.edu

Dustin I. Bessette 

Mt. Hood Community College, USA
Dustin.Bessette@mhcc.edu

ABSTRACT: Agentic artificial intelligence (AI) changes corporate data collection because automated systems can retrieve, interpret, combine, and act on information across digital environments with reduced human intervention (National Institute of Standards and Technology [NIST], 2023, 2024). This qualitative document analysis examines how selected corporate privacy and AI governance documents frame data collection, security, autonomy, consent, retention, and user control (Bowen, 2009; Saldaña, 2021). Public privacy policies, AI principles, responsible AI standards, data processing addenda, and acceptable use documents from technology companies were reviewed as qualitative data because corporate documents can function as organizational evidence (Bowen, 2009; Microsoft, 2022; Salesforce, 2023). The analysis used purposive document sampling and descriptive coding to identify patterns in how organizations justify data collection while describing safeguards for privacy and security (Creswell & Creswell, 2023; Saldaña, 2021). Findings suggest that corporate documents often connect data collection to service delivery, safety, fraud prevention, system performance, and responsible AI, while public disclosures vary in detail about automated action and user remedies (Federal Trade Commission [FTC], 2024; NIST, 2024). The paper contributes a concise qualitative framework for reviewing public corporate documents as evidence of privacy risk governance in agentic AI environments (Bowen, 2009; NIST, 2020, 2023).

KEYWORDS: Agentic Artificial Intelligence, Corporate Data Collection, Individual Autonomy, Privacy Risk, Qualitative Document Analysis, Security Governance

Introduction

Corporate data collection is a routine feature of digital life because organizations collect user, employee, customer, device, behavioral, location, and interaction data for service delivery, security, fraud prevention, and system improvement (FTC, 2024; Schneier, 2015). AI creates additional privacy risk because automated systems can process large data collections, generate inferences, and support decisions that may affect individuals or groups (NIST, 2023; Zuboff, 2019). Agentic AI adds another layer of concern because agents may initiate tasks, call tools, access systems, retrieve records, and recommend actions with limited human awareness (NIST, 2024; OpenAI, 2026; Jaff et al., 2024). Agentic AI denotes systems capable of starting and performing multi-step tasks with reduced human oversight (NIST, 2024).

This paper examines the language that corporations use to explain data collection and AI governance because public documents shape notice, accountability, and user expectations (Bowen, 2009; Jobin et al., 2019). The analysis focuses on privacy statements, responsible AI principles, data processing addenda, and acceptable use policies because these documents contain official language about data use and governance (Google, n.d.-b; Microsoft, n.d.; Salesforce, 2023). The topic fits qualitative document analysis because the documents serve as the data source and can be coded for repeated meanings, gaps, and governance patterns (Bowen, 2009; Saldaña, 2021).

The purpose of this exploratory qualitative document analysis is to examine how corporate privacy and AI governance documents describe the balance between security and individual freedom in an era of agentic AI (NIST, 2020, 2023). The guiding research question asks how public corporate privacy and AI governance documents frame privacy risk, security justification, autonomy, and transparency when data collection supports AI-enabled services (FTC, 2024; OECD, 2019). The study addresses this question by reviewing selected public documents as qualitative evidence as opposed to conducting interviews, focus groups, or surveys (Bowen, 2009; Creswell & Creswell, 2023). This study of public corporate documents provides a qualitative framework that researchers and governance teams can employ to assess privacy risk in agentic AI environments (Bowen, 2009; NIST, 2024).

Background

Privacy and security frequently appear together in corporate governance language, although the concepts can create different operational demands (NIST, 2020; Schneier, 2015). Security programs may require monitoring, authentication logs, anomaly detection, fraud analytics, endpoint telemetry, and account activity review to reduce misuse and protect systems (NIST, 2020, 2023). Privacy programs emphasize collection limits, data minimization, dignity, notice, and user control because privacy risk includes the possibility that data processing can affect people in ways they cannot anticipate or manage (NIST, 2020; Zuboff, 2019). This analysis is based on the United States regulatory environment, where enforcement activity from the Federal Trade Commission serves as the foremost enforcement channel for addressing corporate data practices in the dearth of a comprehensive federal privacy statute (FTC, 2024).

Agentic AI intensifies the privacy and security relationship because agents may operate across applications, databases, collaboration tools, and external services (NIST, 2024; OpenAI, 2026). Traditional privacy notice may tell users what categories of information are collected, but agentic systems raise more questions about which tools the system can call, which approvals are necessary, and how automated activity is logged; these are concerns that empirical examination has revealed to be disconnectedly addressed in practice (NIST, 2024; OpenAI, 2025; Jaff et al., 2024). The Organisation for Economic Co-operation and Development (OECD) AI principles connect privacy, human-centered values, transparency, safety, and accountability, which makes them useful for assessing public-facing AI governance language (OECD, 2019).

Corporate documents also function as public governance artifacts because organizations use them to describe privacy, security, data processing, acceptable use, and responsible AI commitments (Bowen, 2009; Microsoft, 2022). Public documents may not disclose every internal control, but they provide evidence of how organizations communicate data practices to users, customers, regulators, and enterprise clients (Bowen, 2009; FTC, 2024). For that reason, privacy statements, responsible AI standards, data processing addenda, and acceptable use policies can be analyzed as qualitative evidence of

governance communication, a method developed in the Methodology and Design section below (Bowen, 2009; Saldaña, 2021).

Assumptions, Limitations, and Delimitations

Assumptions. The analysis assumes that public corporate documents are meaningful governance artifacts because they represent official language about privacy, security, data use, and AI governance (Bowen, 2009; Microsoft, 2022). The analysis also assumes that public-facing texts can reveal patterns in organizational framing, even when they do not disclose internal technical controls or operational details (Bowen, 2009; Saldaña, 2021). This assumption is appropriate for qualitative document analysis because documents can be treated as records that show how organizations present policies, values, and governance positions (Bowen, 2009). Additionally, the analysis assumes that uniform application of descriptive and pattern coding by a single researcher can produce dependable thematic findings across the document sample (Saldaña, 2021). As given by Gruwell and Ewing (2022), offering assumptions ascertain transparency and credibility. Further, by openly defining the foundational conditions required for research to be valid, this researcher can manage reader expectations, thwart misconception of findings, and offer a clear framework for future researchers to test or build upon the work (Gruwell & Ewing, 2022).

Limitations. The main limitation is that the document sample is small and relies on public materials, which means the findings address document language as opposed to verified internal practice (Bowen, 2009; Creswell & Creswell, 2023). An additional limitation is that one researcher conducted all coding without a second coder or interrater reliability check, which means the findings reflect a single interpretive perspective instead of a verified consensus (Saldaña, 2021; Creswell & Creswell, 2023). The paper does not test whether organizations follow each public statement in practice because the study design does not include audits, interviews, confidential records, or technical testing (Creswell & Creswell, 2023; Saldaña, 2021). The findings should consequently be interpreted as evidence of public governance communication as opposed to evidence of operational compliance (Bowen, 2009; NIST, 2020). This research acknowledges limitations based on guidance from Labaree (2026). These limitations define the constraints on how broadly the findings can be generalized, applied, or interpreted; they reflect the study's design choices, validity procedures, and any unforeseen challenges that occurred during the research process (Labaree, 2026).

Delimitations. The study is delimited to selected technology-sector documents that are publicly available and connected to privacy, responsible AI, data processing, or acceptable use (Google, n.d.-b; Microsoft, n.d.; Salesforce, 2023). The five organizations examined were selected for public prominence in enterprise or agentic AI deployment during the study period, which supports a manageable and relevant document sample as opposed to attempting an exhaustive industry coverage. The study does not compare all industries, all global privacy laws, or all AI products because the purpose is to present a manageable qualitative model (Creswell & Creswell, 2023; Saldaña, 2021). The delimitations support a focused review of public corporate language in relation to agentic AI privacy risk (Bowen, 2009; NIST, 2024). Delimitations were included to explain the boundaries of this study by stating what was included and excluded (Didelot, 2024).

Research Gap

Prior scholarship and policy guidance have examined privacy, AI ethics, responsible AI principles, and enterprise risk management (Floridi et al., 2018; Jobin et al., 2019; NIST, 2023). Jobin et al. (2019) identified recurring AI ethics themes such as transparency, justice, responsibility, and privacy across many guidelines. Floridi et al. (2018) presented ethical

principles for a good AI society, and the NIST AI Risk Management Framework offers a structured approach for identifying and managing AI risk (NIST, 2023).

Amongst the works reviewed above, none examine how public corporate documents themselves frame privacy risk when agentic AI may act across systems and contexts, which indicates a gap in qualitative attention to this specific evidence source (Floridi et al., 2018; Jobin et al., 2019; NIST, 2023, 2024). The gap is not limited to whether a privacy policy mentions collection, consent, or sharing because those topics already appear in many corporate privacy documents (Google, n.d.-b; Microsoft, n.d.; OpenAI, 2026). The concern is whether public language explains the relationship among security justification, user autonomy, automated action, retention, and governance in a way that supports meaningful comprehension and accountability (NIST, 2020, 2024; OECD, 2019).

Methodology and Design

This study used an exploratory qualitative document analysis design because documents, records, and organizational texts can serve as primary qualitative data (Bowen, 2009). Document analysis fit the study because the research did not require human subjects, interviews, focus groups, or survey responses (Creswell & Creswell, 2023; Saldaña, 2021). The design examined public corporate texts as evidence of how organizations describe privacy, security, AI governance, and user control (Bowen, 2009; NIST, 2020).

The document sample was purposive because documents were selected for relevance to privacy, data processing, AI governance, and agentic AI risk (Creswell & Creswell, 2023; Saldaña, 2021). Documents were retrieved between June 12, 2026, and July 2, 2026, which grounds the document sample in a specific, verifiable time frame instead of leaving retrieval dates unspecified (Bowen, 2009).

Eligible documents included privacy statements, responsible AI principles, data processing addenda, acceptable use policies, and public trust-center materials (Google, n.d.-b; Microsoft, 2022; OpenAI, 2025; Salesforce, 2023). Documents were excluded when they were marketing-only pages, news releases, duplicate pages, or materials that did not address data collection, privacy, security, or AI-related governance (Bowen, 2009; Creswell & Creswell, 2023).

Table 1. Corporate document sample for qualitative review

Organization	Public document type reviewed	Reason for inclusion
Google	Privacy policy and AI principles	Provides public language on data collection, user controls, and AI principles (Google, n.d.-b.).
Microsoft	Privacy statement and Responsible AI Standard	Provides privacy notice and AI governance requirements (Microsoft, 2022, n.d.-b).
OpenAI	Enterprise privacy and data processing addendum	Provides public language on enterprise data handling and processing terms (OpenAI, 2025, 2026).
Salesforce	AI acceptable use policy and trusted AI principles	Provides public language on acceptable AI use and trust principles (Salesforce, 2023, n.d.-b).
IBM	Responsible technology statement	Provides public language on responsible technology and governance commitments (IBM, n.d.-b).

The search strategy used targeted searches of company websites, trust centers, legal pages, privacy pages, and responsible AI pages because those locations commonly host public governance documents (Bowen, 2009; Google, n.d.-b; Microsoft, n.d.). Search phrases included privacy policy, responsible AI, AI governance, data processing addendum, data retention, user consent, automated systems, trust center, and acceptable use policy

(Bowen, 2009; Saldaña, 2021). Documents were included when they addressed privacy, data collection, data processing, AI governance, automated systems, security justification, user choice, retention, or customer control (Creswell & Creswell, 2023; NIST, 2020, 2023).

Coding followed a descriptive process because first-cycle coding can capture repeated words, meanings, and policy claims in qualitative data (Saldaña, 2021). Coding was conducted using a structured matrix organized by document and theme, which created a traceable record of coding decisions (Saldaña, 2021). The first coding stage captured terms and meanings connected to consent, security justification, retention, transparency, human oversight, automated action, user control, and downstream data sharing (Saldaña, 2021; NIST, 2024). Pattern coding then grouped descriptive codes into broader governance themes so that documents could be compared across privacy, security, and AI risk dimensions (Bowen, 2009; Saldaña, 2021).

Originality and Contribution of the Text

The contribution of this paper is a concise qualitative framework for examining public corporate documents as evidence of privacy risk governance in agentic AI environments (Bowen, 2009; NIST, 2024). Many discussions of AI privacy focus on laws, technical controls, or ethical principles, while this paper focuses on the documents that users and enterprise customers may read when trying to recognize data practices (Floridi et al., 2018; Jobin et al., 2019). This focus matters because public documents can show how organizations translate privacy and responsible AI commitments into user-facing and customer-facing language (FTC, 2024; Microsoft, 2022).

The paper also contributes an agentic AI privacy lens because a privacy document should be reviewed for more than collection, use, sharing, and retention statements (NIST, 2020, 2024). Under this lens, reviewers also ask whether the document explains what automated or agentic systems can do, how system actions are constrained, when human review occurs, how data movement is limited, and how individuals can exercise control (NIST, 2024; OECD, 2019). This approach extends qualitative document analysis by connecting document language to privacy, autonomy, security, and contestability in AI-enabled workflows (Bowen, 2009; Saldaña, 2021).

Research Insights

Five themes emerged from the document review: security as a justification for data collection, privacy as governance promise, transparency as notice, autonomy through controls, and agentic risk as an underdeveloped disclosure area (Bowen, 2009; Saldaña, 2021). These themes reflect how public documents describe why data are collected and how privacy safeguards are presented to users and customers (FTC, 2024; NIST, 2020). The themes also show that agentic AI risk requires language about tool access, cross-system action, logging, human review, and user remedies (NIST, 2024; OpenAI, 2026).

Security as a justification for data collection appeared across documents because privacy language often ties collection to fraud prevention, abuse detection, account protection, service reliability, and system safety (FTC, 2024; Google, n.d.-b; Microsoft, n.d.). This framing is reasonable from an enterprise risk management perspective because organizations use data to detect misuse, protect systems, and reduce harm (Jones, 2024; NIST, 2020, 2023). The privacy concern is that broad security language can normalize broad data collection unless collection limits, retention rules, and user controls are explained with enough operational detail (NIST, 2020; Schneier, 2015).

Privacy as governance promise appeared through commitments to responsible AI, data protection, privacy design, access controls, and data processing terms (Microsoft, 2022; OpenAI, 2025; Salesforce, n.d.). Microsoft, Google, IBM, OpenAI, and Salesforce

provide public language that connects AI to privacy or responsible use (Google, n.d.-a; IBM, n.d.; Microsoft, 2022; OpenAI, 2026; Salesforce, n.d.). The level of detail varied by document type because privacy statements tend to speak to users in general terms, while data processing addenda and standards tend to use enterprise or contractual language (Microsoft, 2022; OpenAI, 2025).

Transparency often appeared as notice instead of operational explanation because documents describe categories of data, purposes of use, and general safeguards more often than they explain agentic workflows (Google, n.d.-b; Microsoft, n.d.; NIST, 2024). Public documents were less consistent in explaining how automated systems combine data, what data an agent can access during a task, or how a user can review an automated recommendation or action (NIST, 2024; OpenAI, 2026). Transparency without operational meaning may give users awareness without usable control, which conflicts with the accountability and human-centered values emphasized in AI governance guidance (Jobin et al., 2019; OECD, 2019).

Autonomy appeared through privacy settings, account tools, deletion processes, customer controls, disclosure obligations, and user-facing choices (Google, n.d.-b; Microsoft, n.d.; OpenAI, 2026). These mechanisms support user agency because they give individuals and customers some ability to manage data or request changes (NIST, 2020; OECD, 2019). The mechanisms may not fully address agentic AI workflows because an individual privacy choice in one service may interact with another platform, plug-in, customer database, or third-party service during an automated task (NIST, 2024; OpenAI, 2025).

Agentic risk was the least developed disclosure area across the reviewed public documents because public language often does not explain the full chain of agentic action (NIST, 2024; OpenAI, 2026). Some documents refer to AI systems, automated systems, bots, or agents, but the public text may not specify tool-access limits, cross-system retrieval limits, human approval points, or agent activity log retention (NIST, 2024; Salesforce, 2023). Missing details may leave users unable to challenge or comprehend an action taken through an AI-enabled workflow, which raises accountability concerns in relation to responsible AI principles (Jobin et al., 2019; OECD, 2019).

Table 2. Coding themes and privacy-risk meaning

Theme	Evidence in document language	Privacy-risk meaning
Security justification	Fraud prevention, abuse detection, service safety, account protection, reliability.	Security language can justify broad data collection unless collection limits are stated (NIST, 2020; Schneier, 2015).
Governance promise	Responsible AI, privacy, security, accountability, data processing terms.	Governance claims need operational detail to support trust (Jobin et al., 2019; NIST, 2023).
Transparency as notice	Categories of data, purpose statements, user-facing policy language.	Notice may not explain agent decisions, tool access, or downstream movement (NIST, 2024).
Autonomy through controls	Settings, deletion tools, customer controls, disclosure duties, human contact options.	Controls may be fragmented across services and agents (NIST, 2020; OECD, 2019).
Agentic-risk gap	Limited detail on autonomous action, human approval, logging, and cross-system retrieval.	Users may not know how data can be acted upon by agents (NIST, 2024; OpenAI, 2026).

Discussion

The five themes identified above touch on a single interpretive pattern (Bowen, 2009; Saldaña, 2021). The findings suggest that corporate privacy and AI governance documents are useful but incomplete sources for interpreting privacy risk in agentic AI environments (Bowen, 2009; NIST, 2024). Public documents tend to describe why data are collected and how privacy is valued, but they provide less detail about how agentic systems may operate after data are collected (FTC, 2024; NIST, 2024). This pattern creates a governance communication problem because people are asked to trust data practices without always receiving workflow-level explanations of automated action, tool access, retention, and user remedies (NIST, 2020, 2024).

The security and autonomy tension is visible because security programs need information to detect threats, prevent fraud, authenticate users, protect systems, and investigate abuse (NIST, 2020, 2023). Individual autonomy requires limits on surveillance, proportionate data use, meaningful choice, and protection against unnecessary profiling or automated pressure (OECD, 2019; Zuboff, 2019). When agentic AI is added, the same data that support security can become input for automated interpretation, prediction, recommendation, or action (NIST, 2024; OpenAI, 2026).

A privacy policy may satisfy notice expectations while still being difficult for users to apply to agentic AI workflows (FTC, 2024; NIST, 2024). A user may comprehend that a service collects interaction data, but the user may not know whether an AI agent can use the information to complete a task, connect to another application, trigger a workflow, or share outputs with a third party (NIST, 2024; OpenAI, 2025). This gap supports the need for agent-specific transparency statements that explain data access, tool use, human review, retention, logging, and user remedies (NIST, 2024; OECD, 2019).

The analysis also suggests that corporate documents should be read as a document system instead of as isolated pages (Bowen, 2009; Saldaña, 2021). A privacy statement, responsible AI standard, trust center, data processing addendum, and acceptable use policy may each address a different part of governance (Microsoft, 2022, n.d.; OpenAI, 2025; Salesforce, 2023). Fragmented disclosure may weaken user perception because agentic AI risk exists across the workflow instead of within one document or platform (McDonald & Cranor, 2008; NIST, 2024; OECD, 2019).

Practical Implication

Organizations can strengthen public privacy and AI governance documents by adding agentic AI disclosure elements that describe data access, tool use, automated action, logging, retention, and user remedies (NIST, 2024; OECD, 2019). At least one organization in this sample already necessitates disclosure when users interact with a named automated agent, which offers a starting model other organizations could extend to cover tool access and logging as well (Salesforce, 2023). These elements should identify what data an agent may access, what actions it may initiate, whether it can use external tools, when human review is required, and how people can contest or stop an action (NIST, 2024; Nobles & Burrell, 2024; OpenAI, 2026). These additions would make privacy communication more useful for individuals and enterprise customers because they would connect privacy notice to agentic system behavior (FTC, 2024; NIST, 2020).

Organizations can also use document analysis as an internal governance audit because public and internal documents should tell a consistent story about privacy, security, and AI governance (Bowen, 2009; NIST, 2023). A cross-functional team can compare privacy statements, responsible AI principles, acceptable use policies, data processing terms, and security documentation to identify gaps in language and accountability (Burton et al., 2023; Microsoft, 2022; OpenAI, 2025; Salesforce, 2023). When documents describe

security, privacy, and AI in disconnected ways, the organization may need stronger governance alignment before deploying agents in sensitive workflows (NIST, 2023, 2024).

For researchers, the paper offers a manageable qualitative method because a scholar can select a document sample, define inclusion criteria, code recurring themes, and compare governance framing (Bowen, 2009; Creswell & Creswell, 2023). This approach is publishable as empirical qualitative work because it analyzes real documents as data instead of summarizing literature alone (Bowen, 2009; Saldaña, 2021). Future qualitative studies can expand this model by comparing sectors, document types, employee-facing policies, and consumer-facing disclosures (Creswell & Creswell, 2023; NIST, 2024).

Conclusions

Agentic AI expands the privacy-risk landscape because data collection can lead to automated retrieval, interpretation, linking, recommendation, and action (NIST, 2024; OpenAI, 2026). Corporate privacy and AI governance documents are valuable public evidence of how organizations explain these practices, even though public documents do not prove internal compliance (Bowen, 2009; FTC, 2024). This qualitative document analysis found that reviewed documents frequently connect data collection to security, safety, service delivery, and responsible AI, but agentic AI risks are not always described with workflow-level detail (NIST, 2020, 2023, 2024).

The central conclusion is that privacy governance for agentic AI requires more than general privacy notice and broad responsible AI principles (Jobin et al., 2019; NIST, 2024). Public documents should connect security needs with individual autonomy, contestability, tool access limits, logging, retention, and human oversight (NIST, 2020, 2024; OECD, 2019). Future research should expand the document sample across industries, compare consumer-facing and employee-facing documents, and examine whether public disclosures align with internal governance controls (Bowen, 2009; Creswell & Creswell, 2023).

References

- Bowen, G. A. (2009). Document analysis as a qualitative research method. *Qualitative Research Journal*, 9(2), 27-40. <https://doi.org/10.3316/QRJ0902027>
- Burton, S. L., Burrell, D. N., Nobles, C., & Jones, L. A. (2023). Exploring the nexus of cybersecurity leadership, human factors, emotional intelligence, innovative work behavior, and critical leadership traits. *Scientific Bulletin*, 28(2), 162-175. <https://doi.org/10.2478/bsaft-2023-0016>
- Creswell, J. W., & Creswell, J. D. (2023). *Research design: Qualitative, quantitative, and mixed methods approaches* (6th ed.). SAGE Publications.
- Didelot, T. R. (2024). *A little bit of everything: Persistence though the pre-medical pipeline for first and continuing generation students* [Doctoral dissertation, University of Louisville]. Electronic Theses and Dissertations. <https://ir.library.louisville.edu/etd/4477>
- Federal Trade Commission. (2024). *A look behind the screens: Examining the data practices of social media and video streaming services*. <https://www.ftc.gov/reports/look-behind-screens-examining-data-practices-social-media-video-streaming-services>
- Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People, an ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689-707. <https://doi.org/10.1007/s11023-018-9482-5>
- Google. (n.d.-a). *AI principles*. <https://ai.google/principles/>
- Google. (n.d.-b). *Google privacy policy*. <https://policies.google.com/privacy>
- Gruwell, C., & Ewing, R. (2022). What about assumptions? In *Critical thinking in academic research* (2nd ed.). Minnesota State Pressbooks. <https://minnstate.pressbooks.pub/ctar2/chapter/what-about-assumptions/>
- IBM. (n.d.). *IBM's approach to responsible technology*. <https://www.ibm.com/trust/responsible-technology>
- Jaff, E., Wu, Y., Zhang, N., & Iqbal, U. (2024). Data exposure from LLM apps: An in-depth investigation of OpenAI's GPTs. *arXiv*. <https://arxiv.org/abs/2408.13247>

- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389-399. <https://doi.org/10.1038/s42256-019-0088-2>
- Jones, L. A. (2024). Unveiling human factors: Aligning facets of cybersecurity leadership, insider threats, and arsonist attributes to reduce cyber risk. *SocioEconomic Challenges*, 8(2), 44-63. [https://doi.org/10.61093/sec.8\(2\).44-63.2024](https://doi.org/10.61093/sec.8(2).44-63.2024)
- Labaree, R. V. (2026). Limitations of the study. In *Organizing your social sciences research paper*. USC Libraries. <https://libguides.usc.edu/writingguide/limitations>
- McDonald, A. M., & Cranor, L. F. (2008). The cost of reading privacy policies. *I/S: A Journal of Law and Policy for the Information Society*, 4(3), 543-568. [PDF] The Cost of Reading Privacy Policies | Semantic Scholar
- Microsoft. (2022). *Microsoft Responsible AI Standard, v2: General requirements*. <https://cdn-dynmedia-1.microsoft.com/is/content/microsoftcorp/microsoft/final/en-us/microsoft-brand/documents/Microsoft-Responsible-AI-Standard-General-Requirements.pdf>
- Microsoft. (n.d.). *Microsoft privacy statement*. <https://www.microsoft.com/en-us/privacy/privacystatement>
- National Institute of Standards and Technology. (2020). *NIST Privacy Framework: A tool for improving privacy through enterprise risk management*, Version 1.0. <https://doi.org/10.6028/NIST.CSWP.01162020>
- National Institute of Standards and Technology. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. <https://doi.org/10.6028/NIST.AI.100-1>
- National Institute of Standards and Technology. (2024). *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile*. <https://doi.org/10.6028/NIST.AI.600-1>
- Nobles, C., & Burrell, D. N. (2024). Exploring the variability of human factors definitions in cybersecurity literature. *MWAIS 2024 Proceedings*, Paper 28. <https://aisel.aisnet.org/mwais2024/28>
- OpenAI. (2025). *Data processing addendum*. <https://openai.com/policies/data-processing-addendum/>
- OpenAI. (2026). *Enterprise privacy at OpenAI*. <https://openai.com/enterprise-privacy/>
- Organisation for Economic Co-operation and Development. (2019). *Recommendation of the Council on Artificial Intelligence*, OECD/LEGAL/0449. <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>
- Saldaña, J. (2021). *The coding manual for qualitative researchers* (4th ed.). SAGE Publications.
- Salesforce. (2023). *AI acceptable use policy*. <https://www.salesforce.com/en-us/wp-content/uploads/sites/4/documents/legal/Agreements/policies/ai-acceptable-use-policy.pdf>
- Salesforce. (n.d.). *Trusted AI: Key principles*. <https://www.salesforce.com/artificial-intelligence/trusted-ai/>
- Schneier, B. (2015). *Data and Goliath: The hidden battles to collect your data and control your world*. W. W. Norton.
- Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. PublicAffairs.