

Privacy Ethics Challenges and Response Mechanisms in Generative AI-Driven Shopping Recommendation Systems

Yingxuan Li

School of Artificial Intelligence, University of Science and Technology Beijing, Beijing, China
liyingxuan@xs.ustb.edu.cn

Abstract. A recommendation system is an important foundation for e-commerce platforms to deliver personalized services. The introduction of large language models (LLMs) has expanded recommendation systems from behavioral matching to semantic reasoning, user modeling, and personalized content generation. This paper argues that this transformation is restructuring the privacy risk landscape within recommendation systems: privacy risks have extended beyond data collection and storage security to encompass inferential privacy risks (where platforms infer sensitive user attributes from ordinary shopping behaviors) and behavioral influence risks (where personalized recommendation outputs reshape users' decision-making environments). Grounded in Nissenbaum's contextual integrity theory and Crawford and Schultz's predictive privacy harms framework, this paper proposes a three-stage privacy risk model of "data collection—attribute inference—behavioral influence," and examines the specific manifestations of this risk structure in shopping recommendation scenarios through conceptual analysis and case analysis. The paper further argues that while technical protection mechanisms such as federated learning and differential privacy can reduce data leakage risks, they struggle to address privacy issues at the inference and output stages. Therefore, a comprehensive governance framework needs to be established across three dimensions: technical protection, platform accountability, and user control rights.

Keywords: Large language models, Recommendation systems, Inferential privacy, User profiling

1. Introduction

Recommendation systems have been widely adopted by e-commerce platforms to personalize consumer experiences, yet their deep reliance on user data raises privacy concerns. A 2019 Pew Research Center survey showed that 81% of respondents believed the risks of corporate data collection outweigh its benefits [1]. In recent years, the introduction of large language models (LLMs) has further intensified this issue, transforming recommendation systems from behavioral matching tools into intelligent systems capable of semantic reasoning and personalized content generation. Privacy risks have consequently shifted from "data collection" toward "feature inference," yet existing research pays insufficient attention to this risk [2-5].

This paper argues that generative AI is restructuring the privacy risk architecture of recommendation systems. LLM-enhanced recommendation systems can infer user identities,

emotions, and even health conditions from unstructured text, with highly opaque inference pathways, and generate personalized content with behavioral steering effects, causing privacy risks to permeate the entire process of input, processing, and output. Accordingly, this paper addresses three research questions: (1) What privacy risks exist in current shopping recommendation systems at the stages of data collection, user profile construction, and recommendation generation? (2) How does the introduction of LLMs alter the privacy risk structure of traditional recommendation systems? (3) How can a privacy protection framework be constructed across the dimensions of technical design, platform governance, and users?

This paper adopts contextual integrity [6] as its core analytical perspective, combined with the predictive privacy harms framework [7] and the recommendation system ethics taxonomy [8], using Ctrip as the illustrative case object, and proposes a three-stage risk model of "data collection—attribute inference—behavioral influence." The research aims to reveal new forms of privacy risk in recommendation systems in the era of generative AI and provide theoretical references for related governance practices.

2. Related concepts and literature review

2.1. From traditional recommendation systems to LLM-enhanced recommendation systems

Traditional recommendation systems primarily employ two approaches—content-based recommendation and collaborative filtering—to predict user preferences based on historical behavior and item features [9]. The key transformation brought by LLM-enhanced recommendation systems lies in the expansion of recommendation mechanisms: LLMs can process unstructured information such as review texts, search queries, and conversation records, and incorporate them into user modeling and recommendation generation processes. In personalized RAG or Agent recommendation frameworks, user preferences may be continuously invoked as part of retrieval, memory, and generation, transforming user information from one-time recommendation inputs into ongoing conditions for model reasoning [4]. From a privacy perspective, this means LLMs can construct user profiles from more indirect and richer cues, laying the foundation for inferential privacy risks.

2.2. User data, user profiles, and inferred data

This paper categorizes the user information processed by recommendation systems into three levels: raw user data, user profile data, and inferred data.

Raw data includes explicit data actively provided by users (such as age, gender, address, preference settings, product reviews, and purchase records) and implicit behavioral data (such as click paths, browsing duration, and scrolling speed). User profiles represent the platform's comprehensive modeling results of user preferences, spending capacity, interest categories, and behavioral patterns. In LLM-enhanced recommendation systems, user profiles may also exist in the form of user embeddings where systems encode users' historical behaviors, interest preferences, and interaction patterns into vector representations to facilitate similarity computation, ranking, and personalized generation [5].

Inferred data is the focus of this paper. It is not information directly provided by users, but rather user attributes reasoned by the system based on existing behavioral data and profile tags. For example, a platform might infer a user's price sensitivity from their repeated comparison of coupons or infer health anxiety from long-term purchases of health supplements. Compared with raw data,

inferred data involves sensitive attributes that users have not actively disclosed or may even be unwilling to disclose, posing higher privacy risks.

This issue is particularly pronounced in LLM-enhanced recommendation systems. LLMs can combine semantic information with long-term behavioral records for deep-level reasoning, and user information is no longer merely static data but is continuously transformed into input conditions for model reasoning and generation [4]. Therefore, privacy risks in recommendation systems cannot be understood solely as "what data the platform has collected" but must also ask "what can the platform infer from it."

2.3. Existing research and governance approaches for recommendation system privacy protection

Existing research on recommendation system privacy unfolds across three dimensions. At the legal and ethical level, China's Personal Information Protection Law and similar legislation emphasize informed consent, data minimization, and automated decision-making control rights [10]. Milano et al. [8] point out that recommendation system ethics involves multiple dimensions, including transparency, autonomy, and fairness; Crawford and Schultz [7] argue that predictive privacy harms occur at the data use stage, which traditional informed consent mechanisms struggle to cover. At the technical level, methods such as federated learning, differential privacy, anonymization, and encryption primarily reduce data leakage and individual identification risks but cannot prevent models from making sensitive inferences from lawfully collected data. At the platform governance level, Friedman and Knijnenburg [11] note that recommendation system privacy protection involves comprehensive issues of system architecture, algorithms, and institutional arrangements that cannot be solved by technology alone.

3. Privacy infringement risk analysis of shopping recommendation systems

3.1. Boundary-crossing risks at the data collection stage: From necessary data to excessive collection

Recommendation systems inevitably require a certain degree of user data input to deliver personalized services, such as which products users have purchased and which content they have spent more time viewing, to predict user preferences and generate recommendation results. However, the performance optimization of recommendation systems typically depends on richer and longer-term behavioral data, giving platforms a commercial incentive to expand the scope of data collection. Data originally necessary for product recommendations may gradually extend to continuous collection of user location, cross-platform behavior, and environmental information.

Data collection risks can be understood from three dimensions. On the user side, e-commerce applications may request device permissions for geolocation, photo albums, contacts, and clipboards, while users find it difficult to determine which permissions are genuinely necessary for recommendations [12]. On the service provider side, user data from shopping platforms may flow into advertising systems, data analytics tools, and marketing networks through third-party SDKs and advertising identifiers, creating cross-context data flows that exceed user expectations [13]. On the algorithm side, recommendation systems track micro-behavioral signals such as video dwell time, scrolling speed, comment interactions, and clicks without purchases to assess user interest intensity and purchase likelihood. LLM-enhanced recommendation systems further intensify this data

demand, as semantic reasoning and personalized generation typically require more complete user background information.

In this process, traditional informed consent mechanisms struggle to play a substantive role. Solove [14] points out that individuals typically cannot foresee the future uses of their data. Obar and Oeldorf-Hirsch [15] further demonstrate that users often click "agree" without actually reading privacy policies. Therefore, the risks at the data collection stage are not only about "unauthorized collection" but also include substantive loss of control under formal authorization. These continuously collected behavioral data do not remain at the input stage but enter the user profiling and personalized generation process, forming the data foundation for inferential privacy risks in the next stage.

3.2. Inferential privacy risks at the user profiling stage: From behavioral cues to sensitive attributes

The core question after data collection is "what can the platform infer from this data?" Inferential privacy risk refers to the system's inference of sensitive attributes that users have not actively disclosed, based on ordinary behavioral data. The key issue is not whether users have provided sensitive information, but whether the system has generated sensitive judgments from non-sensitive inputs. This risk did not originate in the LLM era—as early as 2012, Target stores inferred customers' pregnancy likelihood from purchases of folic acid and loose-fitting clothing and sent maternity product coupons, revealing the basic structure of inferential privacy: users provide consumption data, and the platform generates sensitive attribute judgments.

LLM-enhanced recommendation systems exacerbate this risk on three levels.

First, inference inputs are richer—LLMs can process unstructured information, such as search phrasing, review texts, and customer service conversations [2]. Second, inference pathways are more opaque—the logic of traditional collaborative filtering is relatively traceable, whereas LLMs' semantic reasoning processes are difficult for users or auditors to reconstruct. Third, inference results are more persistent—in RAG or agent architectures, user preferences may be continuously invoked and updated as long-term memory [16], transforming user profiles from static tags into dynamic behavioral representations.

According to Nissenbaum's contextual integrity theory [6], such cross-context associations may cause data from shopping contexts to be reinterpreted as health, psychological, or economic status information. Viewed individually, behaviors such as purchasing sleep aids or comparing low-priced products do not constitute sensitive information. However, when these behaviors are recorded over time, correlated, and semantically interpreted, they may be used to infer users' sleep problems, emotional states, or economic conditions. As Barocas and Nissenbaum [17] point out, privacy risks in modern data processing often arise from the combination and reinterpretation of data, rather than the disclosure of any single data point.

User embeddings present similar risks. When user behavior is transformed into vector representations—a common practice in machine learning systems—these embeddings may inadvertently encode sensitive personal attributes. Song and Raghunathan [18] demonstrate that embedding representations may leak sensitive attribute information; M. Fredrikson, S. Jha, and T. Ristenpart [19] show that model intermediate representations may be used to reverse-engineer individuals' sensitive characteristics.

This stage reveals an important characteristic of inferential privacy risk: its invisibility. Users cannot see what profile tags or sensitive inferences the platform has generated. The system may

have already formed probabilistic judgments and used them for product ranking, ad targeting, or recommendation explanations, while users lack awareness of and control over this process.

Therefore, the key to how LLM-enhanced recommendation systems exacerbate inferential privacy risk lies not in their collecting more sensitive information, but in their ability to generate sensitive conclusions from non-sensitive inputs.

3.3. Behavioral influence risks at the recommendation output stage: From personalized output to behavioral intervention

This paper does not equate all behavioral influence with privacy infringement—personalized recommendations can reduce search costs and improve user experience in many scenarios. This section focuses on a more specific risk: when platforms restructure users' choice environments based on user profiles and inferred data, the privacy concern extends from "what the platform knows" to "how the platform uses this information to influence user decisions."

Differential pricing and differential treatment. Platforms can infer users' willingness to pay based on information such as historical spending amounts, brand preferences, return frequency, and coupon usage records, and, accordingly, display different prices, promotions, or ranking intensities. At this point, user profiles are no longer merely input conditions for recommendation relevance but become the basis for differentiated transaction terms. In China, the "big data price discrimination" controversies surrounding platforms such as Ctrip reflect this concern. Article 24 of China's Personal Information Protection Law requires that automated decision-making should ensure transparency and outcome fairness and must not impose unreasonable differential treatment in transaction prices and other conditions, indicating that once personalized recommendations affect prices and transaction opportunities, issues of fair dealing and autonomous choice are implicated [10].

Choice environment shaping and preference reinforcement. Recommendation systems continuously adjust recommendation results through ongoing learning from user behavior, potentially exposing users to highly similar products and consumption scenarios over extended periods. Yeung's [20] concept of "hypernudge" helps explain this process: platforms exert fine-grained, persistent, and environmentally embedded guidance on user choices through real-time data analysis, interface design, and dynamic ranking. The issue lies not in personalized ranking itself, but in the fact that when platforms primarily optimize recommendations around click-through rates and conversion rates, they may tend to reinforce content with high commercial returns rather than prioritize users' choice diversity and decision-making autonomy.

Personalized persuasion: from "what to recommend" to "how to recommend." LLM-enhanced recommendation systems further change "how to recommend." Systems can automatically generate recommendation rationales and personalized copy through generative models, influencing users' understanding and judgment at the expression level. For instance, they may emphasize discounts and limited-time offers for price-sensitive users, while highlighting brand reputation and service guarantees for quality-oriented users. Research by Fan et al. and Wu et al. demonstrates that generative models are expanding the functional boundaries of recommendation systems, enabling them to participate in recommendation explanation, conversational interaction, and content generation, meaning platforms can continuously adjust persuasion strategies based on user profiles and real-time feedback, making commercial recommendations more dynamic, interactive, and difficult to detect [2]. Susser, Roessler and Nissenbaum [21] argue that digital platforms' optimization mechanisms may influence user decision-making processes through personalized ranking and choice environment design, and, to some extent, undermine user autonomy. When such influence is based on sensitive inferences unknown to users, its legitimacy warrants further scrutiny.

3.4. Summary of the three-stage risk model

Table 1. Three-stage privacy risk model for LLM-enhanced shopping recommendation systems

Stage	System Component	Key Mechanisms	Privacy Risks	Typical Manifestations
Data Collection	Input	App permissions, browsing history, click behavior, third-party SDKs, location data	Excessive collection, informed consent failure, cross-context data flows	Clipboard reading controversies, permission abuse, third-party sharing
Attribute Inference	Processing	User profiling, semantic reasoning, behavioral correlation, user embedding, RAG memory	Sensitive attribute inference, profile invisibility, inferential privacy infringement	Pregnancy inference, health inference, economic status inference
Behavioral Influence	Output	Personalized ranking, dynamic pricing, recommendation copy, conversational shopping guidance	Differential treatment, choice environment contraction, personalized persuasion	Differential pricing, preference reinforcement, manipulative copy

The three stages are interconnected: boundary-crossing at the data collection stage provides material for inference, inference results are used for personalized output, and behavioral influence at the output stage further depends on user profiles formed in the preceding two stages (Table 1). This risk structure also exposes the inadequacy of existing privacy protection frameworks—for instance, informed consent struggles to cover downstream inferences, technical protections struggle to restrict models from forming sensitive inferences based on lawfully collected data, and platform governance rarely provides substantive user control over profiles and inference tags.

4. Response strategies for privacy protection

Regarding privacy protection responses, this paper does not advocate abolishing personalized recommendations, but emphasizes that when recommendation systems begin conducting highly personalized recommendation generation based on sensitive inferences, the traditional privacy protection framework centered on data collection and storage security is no longer sufficient. The following discussion proceeds across three dimensions: the boundaries of technical protection, the responsibilities of platform governance, and the safeguarding of user control rights.

4.1. Boundaries of technical protection

Technologies such as federated learning, differential privacy, anonymization, and encryption remain significant for reducing data leakage risks, but face structural limitations in LLM-enhanced recommendation systems. (1) While these technologies effectively safeguard raw data from unauthorized disclosure, they fail to mitigate the risk of sensitive attribute inference by models trained on lawfully acquired datasets. That is, federated learning can reduce centralized data uploads and differential privacy can reduce individual identification risks, but neither can restrict models

from forming sensitive conclusions based on lawful inputs. (2) Existing technical protections primarily operate at the input stage and offer limited coverage at the output stage—even if input data has undergone protective processing, output content may still leverage the model's internal user profiles to exert highly personalized behavioral influence. Therefore, technical protection is a necessary foundation but not a sufficient condition, requiring further responses through platform governance and user control rights.

4.2. Platform governance: Inference boundaries, persuasion constraints, and algorithmic accountability

The core question that platform governance must address is "what is the platform permitted to infer, and how is it permitted to use inference results?"

Sensitive inference boundaries and inferred data management. Inferring that a user prefers running shoes or a particular price range constitutes normal inference within the consumption context, but inferring that a user may be pregnant, experiencing psychological stress, or facing economic hardship exceeds the reasonable boundaries of the consumption context, constituting what Nissenbaum [6] describes as inappropriate cross-context information flow.

Platforms should bear governance obligations for inferred data: clearly defining the generation purposes and usage scope of inferred data, setting retention periods, and not using sensitive inferences for differential pricing or targeted persuasion. As Crawford and Schultz [7] argue, predictive privacy harms require an independent governance framework because they occur at the data use stage rather than the collection stage.

Constraints on personalized persuasion intensity. This paper argues that the risk of personalized persuasion escalates under the following conditions: the inference is based on sensitive attributes (such as emotional state, economic pressure, or health issues); the recommendation content employs manipulative expression strategies (such as creating urgency, exploiting anxiety, or implying scarcity); and the user is in a relatively vulnerable consumption state (such as late-night shopping, impulse buying tendencies, or periods of financial stress).

Therefore, platforms need to distinguish between the boundary of "helping users make better choices" and "exploiting user vulnerabilities to facilitate transactions." Stricter content review protocols should be mandated for high-risk sectors, including health products, financial services, and services targeting minors and elderly consumers.

Algorithmic audit mechanisms. Platforms should establish regular algorithmic audit mechanisms to examine whether recommendation systems form sensitive inferences from ordinary consumption behaviors and use them for non-recommendation purposes; whether they implement non-transparent differential pricing for different users; and whether systematic emotional manipulation or scarcity interventions exist. Audit results should be disclosed in appropriate forms to the public or regulatory authorities.

4.3. User control rights: Profile visibility, recommendation comprehensibility, and autonomous choice

Solove notes that expecting users to independently understand and control complex data practices is often unrealistic and user protection should confer substantive control rights through institutional design.

Profile visibility and editing rights. Platforms should display to users the main interest tags and recommendation bases, and allow deletion or modification of inaccurate or unwanted tags. For

profile dimensions that may involve sensitive inferences, they should not be generated by default or used without the user's knowledge.

Recommendation explainability. Platforms should explain recommendation bases (e.g., "based on outdoor products you have browsed"), but should not expose or normalize sensitive inferences—they should not display "because you may be pregnant" or "because you may be experiencing financial difficulties recently."

Data deletion and profile reset. Platforms should distinguish between "deleting raw data" and "deleting inference results," providing synchronized clearing options. Users should be able to reset recommendation profiles or turn off personalized recommendations without suffering punitive degradation in service quality.

Contextualized authorization. Based on contextual integrity theory, layered authorization mechanisms can help users set privacy boundaries across different dimensions (e.g., allowing the use of browsing history for product recommendations while refusing health status inferences based on it), rather than being limited to "agree to all" or "reject entirely."

Regulatory safeguards. Regulatory authorities should require platforms to disclose the basic operational logic of recommendation systems; establish higher compliance requirements for high-risk scenarios; and promote the establishment of independent governance rules for inferred data.

The privacy governance framework for LLM-enhanced recommendation systems should center on inference boundaries, persuasion constraints, and user autonomy. Technical protection is a necessary foundation but cannot cover inference and behavioral influence risks. Platforms should bear independent governance obligations for inferred data, restrict personalized persuasion based on sensitive inferences, and establish algorithmic audit mechanisms. Users should be able to view, modify, and delete profiles, understand recommendation bases, and control the boundaries of data use through contextualized authorization.

5. Conclusion

This paper proposes a three-stage risk model of "data collection—attribute inference—behavioral influence," analyzing the generation and escalation mechanisms of privacy risks in LLM-enhanced shopping recommendation systems. The paper further argues that technical protections such as federated learning and differential privacy cover input-stage risks but struggle to govern issues at the inference and output stages. Therefore, a comprehensive governance framework needs to be established around sensitive inference boundaries, personalized persuasion constraints, platform accountability, and user profile control rights.

The core contribution of this paper lies in shifting the analytical focus from "whether data has been leaked" to "what the platform can infer" and "how the platform uses these inferences." Traditional recommendation systems predict preferences based on behavioral co-occurrence, with relatively traceable inference pathways. LLM-enhanced recommendation systems can process natural language expressions and long-term behavioral records, forming more complex and less transparent user profiles, and reshaping users' choice environments through personalized content generation. The resulting inferential privacy risks do not depend on data leakage but manifest as platforms generating sensitive attribute judgments from ordinary consumption behaviors that users have not actively disclosed.

However, this paper has limitations. It adopts conceptual analysis and case analysis methods without conducting systematic empirical audits of specific platforms; some analyses reflect more of a normative deduction of risk mechanisms. Future research could examine the actual occurrence mechanisms of inferential privacy risks and personalized behavioral influence through user

interviews, algorithmic audits, or experimental methods. As generative AI applications deepen in e-commerce recommendations, inferential privacy governance and recommendation output accountability will become important topics in this field.

References

- [1] B. Auxier, L. Rainie, M. Anderson, et al. (2019) Americans and privacy: Concerned, confused and feeling lack of control over their personal information. Pew Research Center, Washington, DC, Tech. Rep..
- [2] W. Fan et al. (2023) Recommender systems in the era of large language models. arXiv preprint arXiv: 2307.02046.
- [3] L. Li, Y. Zhang, and L. Chen. (2023) A survey on Large language models for recommendation: Progresses and future directions. arXiv preprint arXiv: 2305.19860.
- [4] X. Li, P. Jia, et al. (2025) From RAG to agent: A survey on personalization with large language models. arXiv: 2504.10147.
- [5] A. Salemi, H. Zamani, et al. (2024) Personalization of large language models: A survey. Transactions on Machine Learning Research.
- [6] H. Nissenbaum. (2004) Privacy as contextual integrity. Washington Law Review, vol. 79, pp. 119-158.
- [7] K. Crawford and J. Schultz. (2014) Big data and due process: Toward a framework to redress predictive privacy harms. Boston College Law Review, vol. 55, no. 1, pp. 93-128.
- [8] S. Milano, M. Taddeo, and L. Floridi. (2020) Recommender systems and their ethical challenges. AI & Society, vol. 35, pp. 957-967.
- [9] G. Adomavicius and A. Tuzhilin. (2005) Toward the next generation of recommender systems: A survey of the state-of-the-art and possible extensions. IEEE Transactions on Knowledge and Data Engineering, vol. 17, no. 6, pp. 734-749.
- [10] PRC. (2021) Personal Information Protection Law of the People's Republic of China, August 20, effective November 1, Art. 24.
- [11] A. Friedman, B. P. Knijnenburg, K. Vanhecke, et al. (2015) Privacy and recommender systems. Recommender Systems Handbook. Boston: Springer, pp. 649-688.
- [12] J. Reardon, Á. Feal, P. Wijesekera, et al. (2019) 50 ways to leak your data: An exploration of apps' circumvention of the android permissions system. Proc. 28th USENIX Security Symposium. Santa Clara: USENIX Association, pp. 603-620.
- [13] S. Englehardt and A. Narayanan. (2016) Online tracking: A 1-million-site measurement and analysis. Proc. 2016 ACM SIGSAC Conf. Computer and Communications Security. New York: ACM, pp. 1388-1401.
- [14] D. J. Solove. (2013) Privacy self-management and the consent dilemma. Harvard Law Review, vol. 126, pp. 1880-1903.
- [15] J. A. Obar and A. Oeldorf-Hirsch. (2020) The biggest lie on the Internet: Ignoring the privacy policies and terms of service policies of social networking services. Information, Communication & Society, vol. 23, no. 1, pp. 128-147.
- [16] J. S. Park, J. C. O'Brien, C. J. Cai, et al. (2023) Generative agents: Interactive simulacra of human behavior. Proc. 36th ACM Symp. User Interface Software and Technology. New York: ACM, Art. 2.
- [17] S. Barocas and H. Nissenbaum. (2014) Big data's end run around anonymity and consent, Privacy, Big Data, and the Public Good: Frameworks for Engagement. Cambridge: Cambridge University Press, pp. 44-75.
- [18] C. Song and A. Raghunathan. (2020) Information leakage in embedding models. Proc. 2020 ACM SIGSAC Conf. Computer and Communications Security. New York: ACM, pp. 377-390.
- [19] M. Fredrikson, S. Jha, and T. Ristenpart. (2015) Model inversion attacks that exploit confidence information and basic countermeasures. 22nd ACM SIGSAC Conference on Computer and Communications Security, Denver, CO, USA: Association for Computing Machinery, pp. 1322-1333.
- [20] K. Yeung. (2017) "Hypernudge": Big data as a mode of regulation by design. Information, Communication & Society, vol. 20, no. 1, pp. 118-136.
- [21] D. Susser, B. Roessler, and H. Nissenbaum. (2019) Online manipulation: Hidden influences in a digital world. Georgetown Law Technology Review, vol. 4, no. 1, pp. 1-45.