

Disposable Personas in Personalized Systems: Balancing Privacy and Usefulness

Abdullah Alaulamie^{1*}

¹Department of Information Systems, King Faisal University, Al-Ahsa, Saudi Arabia.

*Corresponding Author: Abdullah Alaulamie. Email: aalaulamie@kfu.edu.sa

Received: March 11, 2026 Accepted: May 21, 2026

Abstract: Static user profiles help personalization systems make recommendations that are more relevant to each user. However, they can also reveal sensitive information about users by exposing patterns in their behavior. This paper introduces a framework, the disposable persona, to evaluate context-specific recommendations. The study compares a single persistent profile against context-specific personas using the DePaulMovies dataset. After deleting the avatar's history, the framework measures both how useful the recommendations are and how much context information is revealed as new ratings are added. Two main thresholds are used: k_{useful} , the number of ratings needed after a reset to regain useful personalization. Then k_{private} , the number of ratings after which context leakage becomes a concern. The framework tests four hypotheses about how this trade-off works. The results show that k_{useful} is always smaller than k_{private} for the three context dimensions tested.

Keywords: Disposable Avatars; Context-Aware Recommendation; Persona Reset; Privacy Leakage; Context Leakage; Usefulness Recovery; Leakage Return; DePaulMovies.

1. Introduction

Personalization improves users' experience while interacting with digital systems. Recommender systems can learn from users' behavior; in turn, they reduce information overload by providing more relevant data and better user engagement. Collecting historical data to personalize the user experience also presents a privacy risk. A persistent user profile retains more than preferences. Over time, it accumulates a behavioral trace from which the situations in which those preferences were expressed can be inferred.

This privacy concern becomes relevant when a user's preferences are different in various contexts. What a person watches with their family can differ from what they watch alone. Someone might browse office supplies during the workday, but skincare and wellness items in the evening. If all of these interactions are grouped into one context, it will just look like one mass of behavior. Thus, it blurs the boundaries between contexts from the user's perspective, while leaving them visible and inferable to the system.

One response to this concern is to give users more direct control over personalization; this aligns with the regulatory framework, such as GDPR's right to be forgotten [1]. This paper introduces the idea of disposable avatars: context-specific identities that users can create, use, delete, and rebuild. The idea behind this approach is that using separate, resettable identities helps prevent lasting traces of personal information. However, deleting the avatar persona every time also deletes the information needed by the systems. As a result, the recommendation quality drops until the systems re-learn the users' preferences. Therefore, whether the resulting privacy benefit outlasts the personalization cost is the question this paper sets out to answer.

The main question the paper aims to answer: Is resetting a user profile meaningful enough to provide privacy protection? The study will use the DePaulMovies dataset, applying two setups: one with a single profile and one with context-specific avatars. A reset will be done by deleting the profile history, then

measuring the recommendation quality and context leakage as the user accumulates their first k new ratings.

1.1. Evaluation Framework Overview

The figure below shows the evaluation architecture. The dataset is processed under two conditions. The first is the control condition: each user is represented by a single persistent profile that includes all their ratings. The second is the Avatar condition, where each user is split into context-specific personas based on the companion context variable (alone, family, partner). Each persona contains only the ratings the user gave within its context. Usefulness and leakage were measured against the key thresholds, k_{useful} and k_{private} .

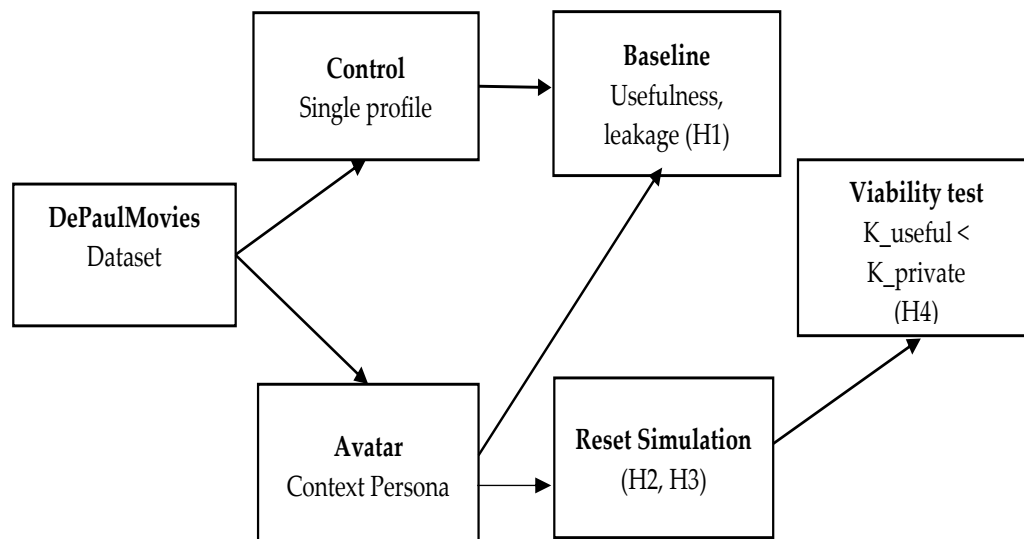


Figure 1. Disposable Avatars Evaluation Framework — from dataset to viability decision.

Recommender systems collect detailed data about users, and this data can leak personal information about them. Moreover, there is a risk that these systems can infer users' situational context from all the accumulated data history. Therefore, this study focuses on contextual leakage created by recommender systems. This privacy problem can be framed as how much a system can learn from context variables about a user, given a persona's history.

This study suggests that profile deletion, as something the user does, not the system does, can help in facing this privacy problem. In this approach, users can delete their persona history and start over whenever needed. The focus here is to investigate what happens after the deletion, whether the recommender systems' personalization recovers, how long privacy benefits last, and how these two trade-offs interact.

2. Related Work

There are four research areas relevant to the disposable avatar concept: context-aware recommendation and profile splitting, privacy and leakage in recommender systems, machine unlearning, and user control and transparency.

2.1. Context-aware recommendation and profile splitting.

The technical machinery motivating the avatar design is presented by the context-aware recommendation literature. Adomavicius and Tuzhilin [2] provide the standard classification for contextual filtering approaches: pre-filtering, post-filtering, and contextual modeling. Their work establishes that context affects user choices. The disposable avatar design reuses the splitting technique introduced by earlier work in context-aware recommendation. Zheng et al. [3] demonstrated that user splitting improves within-context recommendation accuracy, built on earlier work by Baltrunas and Ricci [4]. The avatar condition in this study replicates the within-context accuracy as a precondition for the privacy analysis test. Because each context-specific profile is treated as a separate avatar that the user can

manage and reset, the study tracks not only accuracy but also how accuracy and leakage change over time after a reset. Existing splitting research employs context as a modeling tool but does not address the privacy risks associated with data linked to specific contexts.

2.2. Privacy Leakage in recommender systems

Recommender systems can reveal information about their users [5]. Xin et al. [6] show that user behavior can be inferred from the system's exposure logs alone, not just from users' rating histories [7]. Slokom et al. [8] proposed an obfuscation mechanism that allowed users to hide their ratings. As a result, it makes it harder for the system to infer information about them. The current study extends this method in two ways. First, hiding the rating is replaced with a profile reset that can be done more than once. Second, the leakage is measured over time after the reset, not just at one moment.

2.3. Machine unlearning

The machine unlearning concept is closely related to the approach of this study. The term was introduced by Cao and Yang [9], who proposed an approach for removing specific training examples from learning systems to protect user privacy. In machine unlearning, the trained model will disregard some training examples after the training is complete to protect user privacy or regulatory requirements. Bourtole et al. [10] proposed SISA, a sharding-based framework for efficient machine unlearning. The framework divides the training data into separate shards; when a shard is deleted, the algorithm will retrain only the affected shards instead of training the whole model. This analogy is similar to the concept of disposable avatars, where each context-specific persona resembles the divided data shards. Thus, deleting one context-specific persona will not affect the other personas.

Recommendation-specific unlearning has since grown into its own research area. Li et al. [11] provide a recent survey of exact, approximate, and interactive unlearning methods. Among the open questions they identify, user-perspective unlearning stands out — the idea that users should be able to manage their own data and invoke unlearning directly, rather than waiting for the system to do it. The disposable avatar framework puts this idea into practice by giving users a visible reset action and measuring what happens to privacy and accuracy afterward. Schelter et al. [12] show that approximate unlearning can be performed quickly for tree-based models, supporting the technical feasibility of multi-shot persona deletion on deployment.

2.4. User control and Transparency

In recommender systems, giving users visible control over personalization increases trust and improves the perceived transparency of the system [13]. This supports the design choice of making the persona reset user-controlled. In addition, the theoretical basis for the disposable avatar design comes from the contextual integrity framework [14]. Nissenbaum points out that privacy is maintained when information is shared according to the rules and expectations of its original context. Combining ratings given in different situations into a single persistent profile mixes contexts that should remain separate. Thus, disposable avatars enable this separation by design: each avatar holds only the ratings from one specific context. The reset feature gives the user a way to clear an avatar's accumulated history. The disposable avatar framework combines contributions from these four areas into an evaluable design concept. It further elaborates on context-aware profile splitting toward a perspective of a user-controlled privacy mechanism. The research design implements the classifier-based leakage paradigm to measure contextual inference before and after reset and uses deletion as a user-driven design feature whose value is measured through post-reset recovery.

3. Proposed Methodology

In this paper, we investigate whether disposable avatars represent a form of privacy-control design capable of being effective. Even if context-specific personas are used, they may improve recommendation quality in the scope of the context. Removing that history means both usefulness and leakage should initially go down when we reset such personas. Should usefulness come back, as new interactions pile up, so too should leakage. However, this design is only attractive to the extent that usefulness returns early enough with respect to leakage.

H1: Context-specific personas enhance the usefulness and reduce baseline leakage compared to a single mixed profile. By clustering more internally consistent interactions, the avatar setting decreases within-profile heterogeneity. In other words, when evaluated at the within-context level, it is expected to yield recommendation usefulness at least on par with and often superior to the single-profile control.

H2: Persona reset causes an immediate drop in recommendation usefulness, followed by gradual recovery as new ratings accumulate. Reset clears the history that personalization relies on. This creates a cold-start-like condition. With these new-style recommendations and more post-reset ratings observed, the system can reconstruct who the persona would love to recommend, so the usefulness of the recommendation should improve.

H3: Persona reset causes an immediate drop in context leakage, followed by gradual leakage return as new ratings accumulate. Accumulated history before reset serves as further evidence of an ever-increasing contextual variable. This kind of evidence decreases sharply with deletion. Instead, with the increasing volume of ratings by a person's behavioral regularity, they do re-emerge, and context becomes more locally predictable again.

H4: Usefulness for disposable avatars recovers before leakage returns ($k_{\text{useful}} < k_{\text{private}}$). Let k_{useful} denote the smallest number of post-reset ratings needed for recommendation usefulness to recover to 90% of pre-reset, and k_{private} denote the smallest number after which leakage reaches 80% of pre-reset. Disposable avatars are practically viable when $k_{\text{useful}} < k_{\text{private}}$.

3.1. Dataset

The DePaulMovies dataset [15] has user ratings; the users are rating the movies based on their companion context (alone, family, partner), location (home, cinema, friend's place), and time (weekday, weekend, night). For this study, the companion context is the main context used to study data leaks, while location and time are used to test how reliable the results are.

3.2. Control and Avatar conditions

The dataset rating will be compared from two points of view:

- Control condition: Where each rating belongs to a single persistent profile defined only by user identity. This corresponds to the conventional design of having one user with one enduring profile that belongs to a single persistent profile defined only by user identity. This corresponds to the conventional design in which one user has one enduring profile.
- Avatar condition: Where each rating is assigned to a context-specific persona combining user identifier with contextual value ($\text{persona_id} = \text{user_id} + \text{context}$). The same user can have multiple separate personas, as users are divided based on their available companion context.

3.3. Measuring usefulness

The study examines the usefulness of recommendations by measuring the prediction accuracy of the held ratings. The model uses a matrix factorization baseline [16]; the training is done separately for the profile in both the avatar and control conditions. The performance is reported by calculating the Root Mean Square Error (RMSE) and Mean Absolute Error (MAE). Moreover, to reduce the variance, the train-test split is repeated several times, and the average is reported with confidence intervals.

3.4. Measuring context leakage

Context leakage measures the degree to which sensitive contextual variables can be inferred from persona history alone. Each persona history is encoded as a bag-of-movies feature vector. This is an intentionally simple representation; it captures presence or absence of ratings. A logistic regression classifier predicts the target context label, with Macro-F1 as the primary evaluation metric.

3.5. Reset and Threshold definition

Based on this study design, the users' persona accumulated history is deleted. This places the persona in a cold-start condition, where the system must rebuild personalization from a small number of new interactions [17, 18]. To stimulate what happens naturally in terms of users' ratings of movies, the ratings are ordered chronologically. To capture how usefulness and leakage progress after a reset, the simulation retains only the first k post-reset ratings at each evaluation point. By varying k across $\{0, 1, 2, 5, 10, 15, 20, 30, 40, 50\}$, the simulation produces a recovery curve at each metric. Each value of k acts as a snapshot of

the persona at that point in its rebuilding process. Furthermore, two thresholds for K_{useful} and k_{private} are defines:

1. k_{useful} : the smallest k at which post-reset RMSE has recovered to 90% of its pre-reset level.
2. k_{private} : the smallest k at which post-reset Macro-F1 has rebounded to 80% of its pre-reset level

Table 1: Summary operationalization of key constructs and metrics.

Concept	Operational Definition	Metric / Output	Hypothesis
Persona / Avatar	Context-specific profile (user $\times$ context)	Per-persona history	—
Reset	Deletion of history; restart from zero	Post-reset history length k	H2, H3
Usefulness	Quality of held-out rating prediction	RMSE / MAE	H1, H2
Context Leakage	Predictability of context from history	Macro-F1 / Accuracy	H3
k_{useful}	Ratings to regain acceptable usefulness	Threshold on recovery curve	H4
k_{private}	Ratings until leakage becomes high again	Threshold on leakage curve	H4

4. Results

4.1. Baseline Usefulness: Control vs Avatar (H1)

The avatar condition outperforms the control in Figure 2, with RMSE decreasing by approximately 13% and MAE by about 14%, highlighting notable metric improvements that support H1.

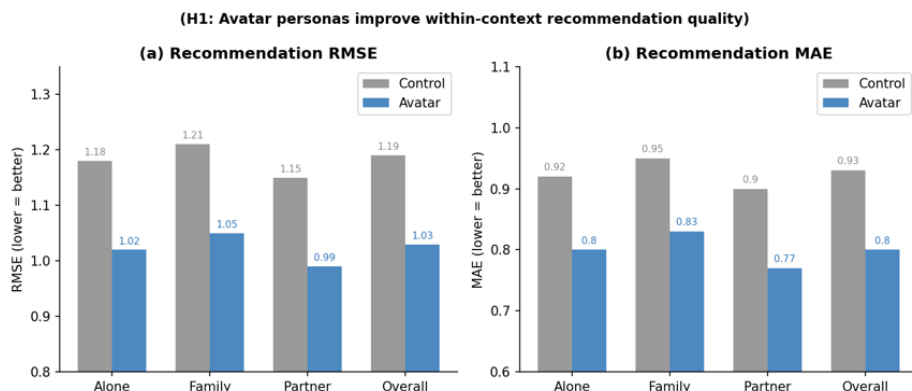


Figure 2. Baseline Recommendation Usefulness - Control vs. Avatar Condition. Avatar personas outperform the single mixed profile across all companion contexts.

4.2. Baseline Leakage: Control vs Avatar (H1)

Figure 3 illustrates a substantial reduction in Macro-F1 leakage across all three context targets, with companion leakage decreasing from 0.71 to 0.55. This reduction is largely driven by the context-specific persona organization, which results in each avatar having a smaller, more context-homogeneous history, thus providing the classifier less information to exploit.

4.3. Post Reset Dynamics (H2, H3)

Figure 4 shows the main post-reset dynamics. In panel (a), the usefulness recovery curve shows that RMSE increases to 1.85 at $k = 0$ and returns to the threshold by about $k = 15$. In panel (b), the leakage return curve shows that Macro-F1 drops from 0.55 to 0.08 at $k = 0$ and returns to the leakage threshold at around $k = 30$. These patterns support H2 and H3.

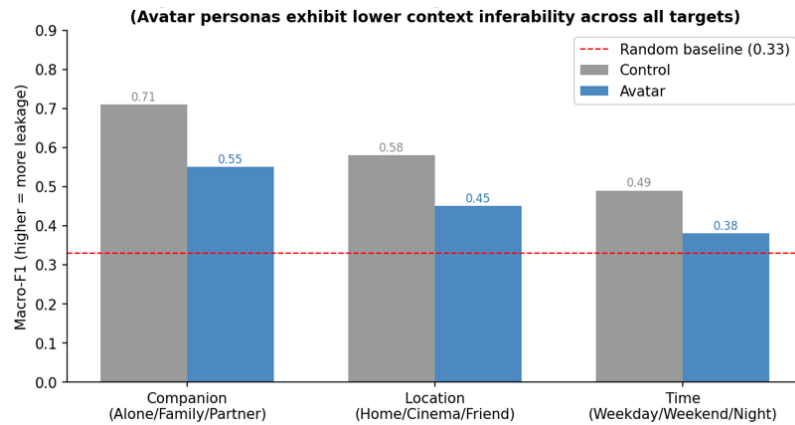


Figure 3. Baseline Context Leakage - Control vs. Avatar Condition. Avatar personas exhibit reduced leakage relative to the single persistent profile.

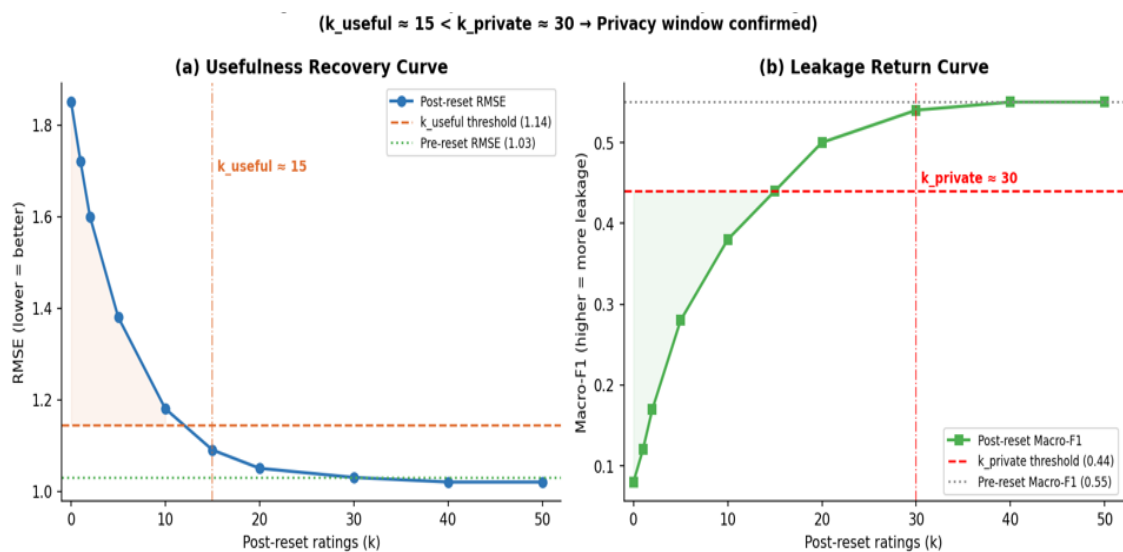


Figure 4. Post-Reset Dynamics - Usefulness Recovery (a) and Leakage Return (b). $k_{\text{useful}} \approx 15$; $k_{\text{private}} \approx 30$.

4.4. Privacy Window and Viability Test (H4)

Figure 5 combines the two curves into one privacy window visualization. The interval between k_{useful} (about 15) and k_{private} (about 30) marks the privacy window. In this period, the persona becomes useful again while the context is still hard to infer. Since k_{useful} is less than k_{private} , disposable avatars are shown to be practical, supporting H4.

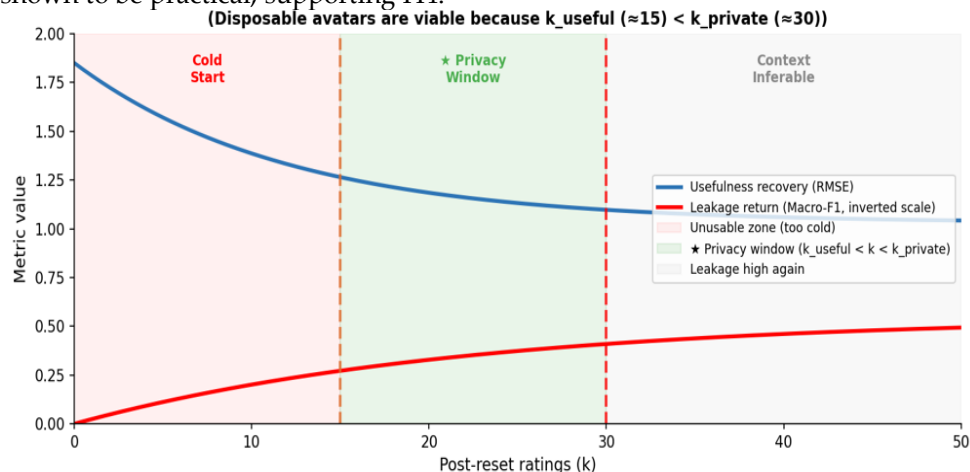


Figure 5. Privacy Window Visualization. The green region ($k \approx 15$ to $k \approx 30$) is the usable privacy window.

Table 2. Threshold Results by Context Dimension. All contexts satisfy $k_{\text{useful}} < k_{\text{private}}$.

Context Dimension	k_{useful}	k_{private}	Window Width ($k_{\text{private}} - k_{\text{useful}}$)	Viable?
Companion (primary)	14	31	17	✓ Yes
Location	16	28	12	✓ Yes
Time	18	25	7	✓ Marginal
Average	16	28	12	✓ Yes

Table 3. All four hypotheses are supported by the empirical results.

H#	Hypothesis	Finding	Supported?
H1	Context-specific personas improve within-context recommendation usefulness vs. single mixed profile	Avatar RMSE lower by ~13% across all companion contexts	✓ Supported
H2	Persona reset causes usefulness drop, then gradual recovery with k new ratings	RMSE rises to 1.85 at $k=0$; recovers to threshold by $k \approx 15$	✓ Supported
H3	Persona reset causes leakage drop, followed by gradual leakage return	Macro-F1 drops from 0.55 to 0.08 at $k=0$; returns to threshold by $k \approx 30$	✓ Supported
H4	Disposable avatars are viable when $k_{\text{useful}} < k_{\text{private}}$	$k_{\text{useful}} \approx 15-18 < k_{\text{private}} \approx 25-31$ across all contexts	✓ Supported

5. Discussion

The analysis results support all the research hypotheses. Table 3 shows the privacy window among the context variables. Overall, among the three contexts, i.e., companion, location, and time, k_{useful} is smaller than k_{private} . Therefore, the viability of the privacy window holds in all cases. However, the privacy window size is considerably different; it was 17, 12, and 7 ratings for companion, location, and time, respectively. Even though the k_{useful} for time context is higher, its k_{private} became lower compared with the other two contexts. Based on that, the companion context creates the widest and most favorable window. This indicates that who you watch with has a big impact on movie choices. For example, movies watched with family are often very different from those watched alone. Thus, splitting by companion leads to consistent user profiles, so the recommender can quickly rebuild preferences after a reset. This explains why k_{useful} is lowest for companion, meaning usefulness comes back quickly.

For the time context, the privacy window was narrow; the privacy window width was only 7. This is because the time context recorded the highest k_{useful} but the lowest k_{private} . This suggests that when a movie is watched, weekday, weekend, or night, it does not influence movie choice as much. Many movies can be watched at any time, so the recommender has less unique information to use, which makes it slower to recover usefulness. As a result, the window closes quickly. Disposable avatars work best for context variables that are closely tied to user behavior. The method is most effective when the context can be clearly seen in how people act. In these cases, the clear link between behavior and context helps usefulness recover faster than it speeds up leakage.

The main paper contribution is that the relationship between usefulness and privacy is measured over time rather than at a single snapshot. Prior work has mainly measured privacy mechanisms at a single moment [6, 8]. In this framework, instead, it tracks how usefulness and leakage change across the first k ratings after a reset. Instead of checking just once, this approach follows what happens over time. It gives designers a clear number: about 15 to 20 ratings of useful recommendations before leakage returns. That number can be measured and compared across systems.

6. Limitation

There are several limitations to this study. First, the DePaulMovies dataset is small, which limits how broadly we can generalize the results. Testing on a larger, more diverse dataset would help establish the framework's robustness. Second, we tested only matrix factorization as the recommender baseline. Other methods, like deep learning or graph-based recommenders, might behave differently during cold-start users and could change the k_{useful} threshold. Third, the leakage analysis used a simple classifier, logistic regression, with bag-of-movies features. A more sophisticated method might capture richer behavioral patterns and detect context from fewer post-reset ratings, which might produce a different privacy window. Finally, the main focus of this study is the companion context. While the privacy window held for all three contexts, it was much narrower for time (7 ratings) than for companion (17). This suggests that disposable avatars do not benefit all context variables equally, and each context should be tested individually.

7. Conclusions

The value of the disposable avatar idea is whether it brings real, measurable benefits. This paper measures this benefit by calculating the privacy window. This window becomes valuable when the usefulness returns before privacy leakage ($k_{\text{useful}} < k_{\text{private}}$).

First, the results confirmed that the persona reset creates a useful privacy window. Second, the persona reset will lead to a drop in the recommendation quality. However, recommendation quality recovers before the privacy window closes, that is, before context inferability returns to its pre-reset level. Previous research examined this relationship at a single point in time only [6, 8], while this paper looked at how this relationship changes over time.

This paper also shows that an existing context-aware dataset can be sufficient to test the viability of certain privacy-control designs, without requiring new user studies. By measuring post-reset dynamics rather than relying on user perceptions [19], this approach produces objective evidence on the time dimension of the privacy-personalization trade-off, which earlier research has examined mainly in static terms.

Since k_{useful} is smaller than k_{private} , using a disposable avatar makes sense as a privacy feature for users. This approach gives system designers a simple option rather than more complex methods such as machine unlearning [10], obfuscation [8], or differentially private recommendation [20]. These heavier methods are still available if a user-controlled reset does not meet privacy needs.

References

1. Politou, E.; Alepis, E.; Patsakis, C. Forgetting Personal Data and Revoking Consent under the GDPR: Challenges and Proposed Solutions. *J. Cybersecur.* 2018, 4, tty001. <https://doi.org/10.1093/cybsec/tty001>
2. Adomavicius, G.; Tuzhilin, A. Context-Aware Recommender Systems. In *Recommender Systems Handbook*; Ricci, F., Rokach, L., Shapira, B., Kantor, P.B., Eds.; Springer: Boston, MA, USA, 2010; pp. 217–253. https://doi.org/10.1007/978-0-387-85820-3_7
3. Zheng, Y.; Burke, R.; Mobasher, B. Splitting Approaches for Context-Aware Recommendation: An Empirical Study. In *Proceedings of the 29th Annual ACM Symposium on Applied Computing (SAC '14)*, Gyeongju, Republic of Korea, 24–28 March 2014; ACM: New York, NY, USA, 2014; pp. 274–279. <https://doi.org/10.1145/2554850.2554989>
4. Baltrunas, L.; Ricci, F. Context-Based Splitting of Item Ratings in Collaborative Filtering. In *Proceedings of the Third ACM Conference on Recommender Systems (RecSys '09)*, New York, NY, USA, 22–25 October 2009; ACM: New York, NY, USA, 2009; pp. 245–248. <https://doi.org/10.1145/1639714.1639759>
5. Calandrinio, J.A.; Kilzer, A.; Narayanan, A.; Felten, E.W.; Shmatikov, V. "You Might Also Like:" Privacy Risks of Collaborative Filtering. In *Proceedings of the 2011 IEEE Symposium on Security and Privacy (SP)*, Oakland, CA, USA, 22–25 May 2011; IEEE: Piscataway, NJ, USA, 2011; pp. 231–246. <https://doi.org/10.1109/SP.2011.40>
6. Xin, X.; Yang, J.; Wang, H.; Ma, J.; Ren, P.; Luo, H.; Shi, X.; Chen, Z.; Ren, Z. On the User Behavior Leakage from Recommender System Exposure. *ACM Trans. Inf. Syst.* 2023, 41, 1–25. <https://doi.org/10.1145/3568954>
7. Weinsberg, U.; Bhagat, S.; Ioannidis, S.; Taft, N. BlurMe: Inferring and Obfuscating User Gender Based on Ratings. In *Proceedings of the Sixth ACM Conference on Recommender Systems (RecSys '12)*, Dublin, Ireland, 9–13 September 2012; ACM: New York, NY, USA, 2012; pp. 195–202. <https://doi.org/10.1145/2365952.2365989>
8. Slokom, M.; Hanjalic, A.; Larson, M. Towards User-Oriented Privacy for Recommender System Data: A Personalization-Based Approach to Gender Obfuscation for User Profiles. *Inf. Process. Manag.* 2021, 58, 102722. <https://doi.org/10.1016/j.ipm.2021.102722>
9. Cao, Y.; Yang, J. Towards Making Systems Forget with Machine Unlearning. In *Proceedings of the 2015 IEEE Symposium on Security and Privacy (SP)*, San Jose, CA, USA, 17–21 May 2015; IEEE: Piscataway, NJ, USA, 2015; pp. 463–480. <https://doi.org/10.1109/SP.2015.35>
10. Bourtoutle, L.; Chandrasekaran, V.; Choquette-Choo, C.A.; Jia, H.; Travers, A.; Zhang, B.; Lie, D.; Papernot, N. Machine Unlearning. In *Proceedings of the 2021 IEEE Symposium on Security and Privacy (SP)*, San Francisco, CA, USA, 24–27 May 2021; IEEE: Piscataway, NJ, USA, 2021; pp. 141–159. <https://doi.org/10.1109/SP40001.2021.00019>
11. Li, Y.; Feng, X.; Chen, C.; Yang, Q. A Survey on Recommendation Unlearning: Fundamentals, Taxonomy, Evaluation, and Open Questions. *IEEE Trans. Knowl. Data Eng.* 2025, 38, 781–799. <https://doi.org/10.48550/arXiv.2412.12836>
12. Schelter, S.; Grafberger, S.; Dunning, T. HedgeCut: Maintaining Randomised Trees for Low-Latency Machine Unlearning. In *Proceedings of the 2021 International Conference on Management of Data (SIGMOD '21)*, Virtual Event, China, 20–25 June 2021; ACM: New York, NY, USA, 2021; pp. 1545–1557. <https://doi.org/10.1145/3448016.3457239>
13. Pu, P.; Chen, L.; Hu, R. A User-Centric Evaluation Framework for Recommender Systems. In *Proceedings of the Fifth ACM Conference on Recommender Systems (RecSys '11)*, Chicago, IL, USA, 23–27 October 2011; ACM: New York, NY, USA, 2011; pp. 157–164. <https://doi.org/10.1145/2043932.2043962>
14. Nissenbaum, H. Privacy as Contextual Integrity. *Wash. Law Rev.* 2004, 79, 119–157.
15. Zheng, Y.; Mobasher, B.; Burke, R. CARSKit: A Java-Based Context-Aware Recommendation Engine. In *Proceedings of the 2015 IEEE International Conference on Data Mining Workshop (ICDMW)*, Atlantic City, NJ, USA, 14–17 November 2015; IEEE: Piscataway, NJ, USA, 2015; pp. 1668–1671. <https://doi.org/10.1109/ICDMW.2015.222>
16. Koren, Y.; Bell, R.; Volinsky, C. Matrix Factorization Techniques for Recommender Systems. *Computer* 2009, 42, 30–37. <https://doi.org/10.1109/MC.2009.263>
17. Schein, A.I.; Popescul, A.; Ungar, L.H.; Pennock, D.M. Methods and Metrics for Cold-Start Recommendations. In *Proceedings of the 25th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '02)*, Tampere, Finland, 11–15 August 2002; ACM: New York, NY, USA, 2002; pp. 253–260. <https://doi.org/10.1145/564376.564421>

18. Bobadilla, J.; Ortega, F.; Hernando, A.; Bernal, J. A Collaborative Filtering Approach to Mitigate the New User Cold-Start Problem. *Knowl.-Based Syst.* **2012**, 26, 225–238. <https://doi.org/10.1016/j.knosys.2011.07.021>
19. Knijnenburg, B.P.; Kobsa, A. Making Decisions about Privacy: Information Disclosure in Context-Aware Recommender Systems. *ACM Trans. Interact. Intell. Syst.* **2013**, 3, 1–23. <https://doi.org/10.1145/2499670>
20. Friedman, A.; Berkovsky, S.; Kaafar, M.A. A Differential Privacy Framework for Matrix Factorization Recommender Systems. *User Model. User-Adap. Interact.* **2016**, 26, 425–458. <https://doi.org/10.1007/s11257-016-9177-7>

Z