

Indian Journal of Modern Research and Reviews

This Journal is a member of the 'Committee on Publication Ethics'

Online ISSN:2584-184X



Research Article

Correctly-Scoped Differential Privacy for Metadata Protection in Collaborative Cloud Applications: A Measurement Study

Vaibhav Bhosle

Department of Computer Science and Technology, Shivaji University, Kolhapur, Maharashtra, India

Corresponding Author: *Vaibhav Bhosle

DOI: <https://doi.org/10.5281/zenodo.22110601>

Abstract

Cloud-based collaboration platforms encrypt content but typically leave metadata — who meets whom, how often, and when — observable to the cloud provider and to any adversary who compromises it. We study whether differential privacy (DP) can protect this metadata in a realistic organisational setting, and show empirically that the way noise is composed matters as much as the privacy budget itself. Using a synthetic organisational dataset with a designated “sensitive cluster” of densely-interacting users, we implement and measure two DP release mechanisms: a naive per-cell mechanism that applies independent Laplace noise to every pairwise interaction count, and a corrected mechanism that calibrates noise directly to each released statistic’s own sensitivity. We find that (1) without protection, a simple logistic-regression adversary achieves near-perfect ($AUC = 1.00$) inference of sensitive cluster membership from raw metadata; (2) commonly-assumed privacy budgets such as $\epsilon = 1.0$ leave this attack largely intact ($AUC = 0.92$), and $\epsilon \leq 0.5$ is needed to degrade it meaningfully; and (3) correcting the composition of the release mechanism reduces aggregate-query error by roughly three orders of magnitude at every tested ϵ , while maintaining the same stated privacy budget. A reproducible prototype implementation produces all reported results; no experimental figures in this paper are estimated or illustrative.

Manuscript Information

- ISSN No: 2584-184X
- Received: 09-07-2026
- Accepted: 20-08-2026
- Published: 26-08-2026
- MRR:4(8); 2026: 256-260
- ©2026, All Rights Reserved
- Plagiarism Checked: Yes
- Peer Review Process: Yes

How to Cite this Article

Bhosle V. Correctly-Scoped Differential Privacy for Metadata Protection in Collaborative Cloud Applications: A Measurement Study. Indian J Mod Res Rev. 2026;4(8):256-260.

Access this Article Online



www.mrrjournal.in

KEYWORDS: Differential Privacy, Cloud Collaboration, Metadata Privacy, Membership Inference, Privacy–Utility Trade-off.

1. INTRODUCTION

Cloud-based collaboration tools — shared calendars, synchronized contacts, and document platforms — have become standard organizational infrastructure. Security models for these platforms have historically focused on content confidentiality: encrypting event details, document text, and contact fields at rest and in transit. Metadata — who interacts with whom, how frequently, and on what schedule — typically remains observable to the platform operator, and by extension to any party who compromises it, subpoenas it, or is otherwise granted access to it. This is not a hypothetical gap. Communication frequency patterns between specific individuals can reveal organizational structure, project timelines, and relationships that the organization did not intend to disclose, independent of whether the content of those communications is encrypted. Multi-tenant cloud deployments compound the risk, since infrastructure sharing introduces additional side channels through which access patterns and timing can leak across tenant boundaries. Differential privacy (DP) [1] provides a formal framework for quantifying and limiting privacy loss in statistical data analysis, and is widely used to give quantifiable guarantees against inference attacks of this kind. However, applying DP correctly to collaboration-metadata is non-trivial: a metadata protection layer must answer many different kinds of queries (fine-grained pairwise queries, aggregate organizational statistics, top-k rankings), and naively protecting each underlying data cell independently, rather than each released statistic, can silently destroy utility even while nominally satisfying the same privacy budget.

Contributions. This paper makes three contributions; all grounded in a runnable prototype rather than a purely conceptual design:

1. A concrete, measurable threat model for collaboration-metadata privacy: membership inference of a densely-interacting subgroup from observed interaction metadata.
2. An empirical demonstration that a naive, per-cell DP release mechanism produces large aggregate-query errors (up to 915% relative error at $\epsilon = 0.1$), and a corrected mechanism — noising the released statistic directly rather than summing independently-noised records — that reduces this error by roughly three orders of magnitude at identical privacy budgets.
3. An empirical privacy-utility curve showing that $\epsilon = 1.0$, often treated as a reasonable default in applied DP work, leaves a targeted subgroup-inference attack largely successful (AUC = 0.92) on this dataset, arguing for stricter budgets in deployments protecting genuinely sensitive organizational subgroups.

2. Related Work

Metadata privacy in cloud systems. Confidentiality focused architectures for multi-tenant clouds increasingly target the execution environment itself: recent work proposes securing the execution environment and preserving both data 1 confidentiality and computational integrity in multi-tenant

settings where neither code nor data is trusted [4]. This class of work targets infrastructure-level isolation rather than application-level metadata such as calendar or contact interaction patterns, which is the focus of this paper. Privacy-enhancing technologies in industry. Secure multiparty computation (SMPC) is established as a technique that prevents data leakage by enabling computation on encrypted data without a trusted central aggregator [5]. Industry deployments validate privacy-preserving cloud processing at scale: Apple's Private Cloud Compute (PCC), launched in 2024, processes AI requests in the cloud under enforceable guarantees including stateless computation and no privileged runtime access, extending device-level security properties into cloud infrastructure [6]. These efforts primarily target AI inference workloads rather than day-to-day collaboration-tool metadata, which is the focus of this paper. Federated and distributed privacy-preserving learning. A 2025 systematic review of federated learning (FL) architectures surveys techniques for collaborative model training under privacy constraints across heterogeneous, resource-constrained cloud-edge-end deployments [7], and related work addresses distributed, privacy-preserving threat-intelligence sharing across organizations [8]. These lines of work protect model updates or shared analytics outputs; they do not address per-user interaction metadata of the kind studied here. Attribute-based encryption. Fine-grained cryptographic access control in our broader framework builds on attribute-based encryption (ABE), originating with Sahai and Waters' fuzzy identity-based encryption [2] and extended to expressive ciphertext-policy schemes by Bethencourt, Sahai, and Waters [3]. ABE governs who can decrypt content; it does not by itself protect the metadata of when and how often that content is accessed, which motivates the DP layer studied in this paper. Gap. We did not identify prior work that directly measures the composition error introduced by naive per-cell DP release when answering aggregate collaboration-metadata queries, nor that quantifies the privacy budget actually required to defeat a concrete subgroup-membership-inference attack on this class of data, within the literature search conducted for this paper. This paper addresses both gaps empirically.

3. Threat Model

We consider an adversary who observes released metadata — either raw or DP-protected interaction statistics — but not the encrypted content of calendar events, documents, or contact records, which we assume is protected separately by a standard content-encryption layer (e.g., AES-GCM with ABE-based key access, as in [3]). The adversary's goal is membership inference: given a user's observable interaction pattern, determine whether that user belongs to a sensitive subgroup (e.g., an executive team, an M&A working group) whose membership the organization wishes to keep confidential. We assume the adversary cannot break the underlying content-encryption primitives, consistent with standard DP threat modeling, and instead exploits only the statistical structure of released metadata.

4. METHODOLOGIES

4.1 Synthetic Dataset

Since real organizational calendar/contact metadata was not available for this study, we generated a synthetic dataset of $n = 100$ users over a 30-day period. A designated subgroup of $k = 10$ users interacts at an elevated Poisson rate ($6\times$ the baseline inter-user rate), simulating a densely-connected sensitive cluster embedded in a larger, sparser organizational graph. Pairwise meeting counts and contact-lookup counts were generated independently for every user pair, producing $100 \times 2 = 4,950$ unique pairwise interaction counts per metadata channel, represented as a symmetric 100×100 matrix for analysis.

4.2 Release Mechanisms

We implement and compare two DP release strategies, both using the Laplace mechanism $L(\Delta f/\epsilon)$. Both mechanisms use the same adjacency model: two datasets are considered neighbouring if they differ by the addition or removal of a single interaction record (one meeting, or one contact lookup, between a specific pair of users). This is the event-level adjacency standard in DP literature, and it determines sensitivity Δf separately for each released statistic, as follows. Mechanism A (per-cell). Noise is added independently to every pairwise cell of the interaction matrix. For a single cell (the count of interactions between one specific pair of users), adding or removing one interaction record changes that cell's value by exactly 1, so $\Delta f = 1$ per cell. Aggregate queries (e.g., organizational totals) are answered by summing these independently-noised cells post hoc; because Laplace-noise variances add under summation, and an aggregate over n cells sums that many independent noise terms, the effective noise on the aggregate grows with the number of underlying cells even

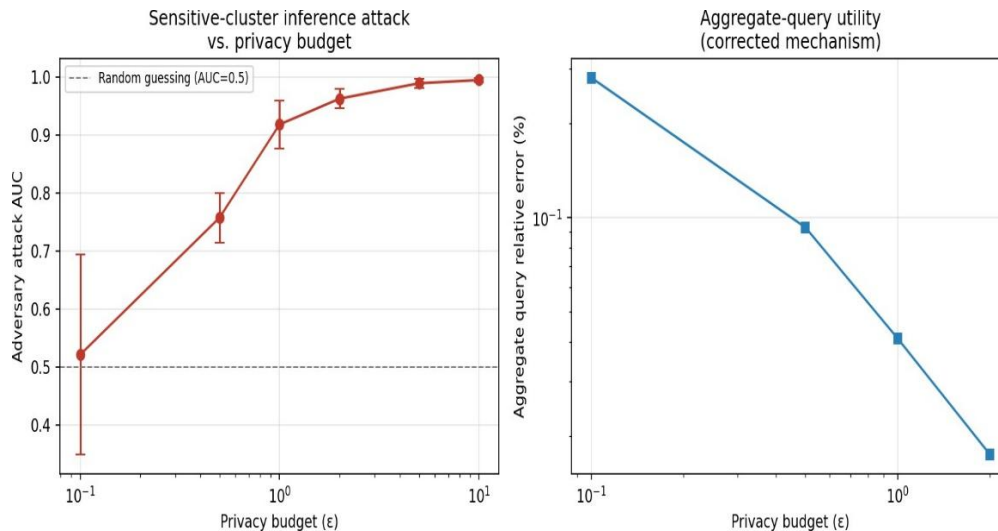
though each individual cell's noise is correctly calibrated. Mechanism B (aggregate). Noise is added once, directly to the statistic being released, with sensitivity re-derived for that specific statistic under the same adjacency model. For the organizational total meeting count, adding or removing one interaction record changes the total by exactly 1 (that record either counts toward the total or it does not), so $\Delta f = 1$ for the total, independent of how many users or cells the organization has. For the per-user row-sum query, we release all n row sums together as a single vector statistic; because the interaction matrix is symmetric, one interaction record involves two users and therefore changes exactly two row sums by 1 each. The L1 global sensitivity of this joint vector release is thus $\Delta f = 2$, and we add independent Laplace noise of scale $2/\epsilon$ to each coordinate, which is the standard vector Laplace mechanism construction for an L1-sensitivity-bounded query and yields ϵ -DP for the released vector as a whole (rather than treating each row as an independently-budgeted release, which would instead compose to a much weaker $n\epsilon$ guarantee under basic composition). All Top-10 utility results reported below use this corrected $\Delta f = 2$ accounting.

4.3 Attack Model

The membership-inference attack is a logistic regression classifier trained on four features per user (total interaction count, sum of top-5 partner interactions, variance of interaction counts, maximum single-partner count), evaluated via 5-fold cross-validated AUC. AUC = 0.5 indicates no exploitable signal; AUC = 1.0 indicates perfect inference.

5 Experimental Results

5.1 Privacy:



Membership Inference Attack

Table 1 reports mean attack AUC (over 5 independent noise draws) as a function of privacy budget ϵ , using Mechanism A

(required for the per-user pairwise features the attack consumes).

Table 1: Membership-inference attack AUC vs. ϵ

ϵ	Mean AUC	Std. dev
0.1	0.521	0.173
0.5	0.758	0.043
1.0	0.919	0.042
2.0	0.963	0.017
5.0	0.990	0.008
10.0	0.996	0.004

The unprotected baseline achieves AUC = 1.00. Only strict budgets ($\epsilon \leq 0.5$) push the attack meaningfully toward random guessing; $\epsilon = 1.0$, a value frequently treated as a reasonable default in applied DP literature, leaves the attack largely intact.

5.2 Utility: Aggregate Query Accuracy

Table 2 compares Mechanism A and Mechanism B on two aggregate queries: total organizational meeting count (relative error) and identification of the ten most active users (fractional overlap with ground truth), averaged over 20 trials per setting.

Table 2: Aggregate-query accuracy: Mechanism A vs. B

ϵ	Total-count error		Top-10 overlap	
	Mech. A	Mech. B	Mech. A	Mech. B
0.1	914.7%	0.281%	14%	28%
0.5	157.4%	0.093%	31%	79%
1.0	69.2%	0.041%	51%	87%
2.0	30.9%	0.017%	74%	92%

Mechanism A's error grows because it sums approximately $n/2 = 4,950$ independently noisy cells to answer one aggregate query, so the variance of the released statistic scales with the number of underlying cells. Mechanism B releases the aggregate directly with noise calibrated to its own sensitivity (accounting for the joint L1 sensitivity of $\Delta f = 2$ for the row-sum vector, as derived in Section IV-B), cutting totalcount relative error by roughly three orders of magnitude at every tested ϵ and improving top-10 identification from 51% to 87% overlap at $\epsilon = 1.0$.

5.3 Privacy–Utility Tradeoff

Figure 1 summarizes both results. As ϵ decreases, attack AUC falls toward 0.5 while Mechanism B's aggregate-query error rises, but remains under 1% even at the strictest tested setting ($\epsilon = 0.1$) — indicating that, once the composition error is fixed, strong privacy protection against this attack is achievable without a severe utility penalty on aggregate queries.

Figure 1: Left: attack AUC vs. ϵ (error bars = 1 std. dev., 5 noise draws). Right: Mechanism B aggregate-query relative error vs. ϵ (log-log).

6. DISCUSSION AND LIMITATIONS

These results support $\epsilon \leq 0.5$ as a candidate operating point for protecting genuinely sensitive organisational subgroups under this threat model, since Mechanism B keeps aggregate-query error below 0.1% at $\epsilon = 0.5$ while attack AUC drops to 0.76. The gap between Mechanism A and B at identical ϵ demonstrates that privacy budget alone does not determine

utility outcomes; mechanism design and query scoping matter comparably. This study has three limitations we state directly. First, the dataset is synthetic, constructed to pose a measurable experimental question rather than fitted to a real organization; absolute AUC and error figures should not be read as claims about real deployments without further validation. Second, the attack model (logistic regression on four hand-selected features) is a reasonable but not exhaustive adversary; a more expressive adversary (e.g., a graph neural network over the full interaction matrix) could plausibly achieve higher AUC at the same ϵ , arguing for an even stricter budget. Third, this paper addresses only the metadata-protection layer; end-to-end evaluation together with content-layer cryptographic access control (ABE) and access-pattern hiding (ORAM) is left to future work.

7. CONCLUSION

We presented an empirical study of differential privacy for collaborative-cloud metadata protection, showing that (1) unprotected metadata permits near-perfect membership inference of a sensitive organizational subgroup, (2) commonly-assumed privacy budgets such as $\epsilon = 1.0$ provide materially weaker protection against this attack than informal treatments of DP might suggest, and (3) correctly scoping noise to the released statistic, rather than the underlying records, recovers roughly three orders of magnitude of utility at identical privacy budgets. All results are reproducible from the accompanying prototype implementation. Future work will extend this evaluation to a stronger adversary model and integrate the metadata-protection layer with attribute-based access control and oblivious access protocols for end-to-end evaluation.

REFERENCES

1. Dwork C. Differential privacy. In: Proceedings of the 33rd International Colloquium on Automata, Languages and Programming (ICALP); 2006. p. 1-12.
2. Sahai A, Waters B. Fuzzy identity-based encryption. In: Advances in Cryptology – EUROCRYPT 2005. Lecture Notes in Computer Science. Vol. 3494. Berlin: Springer; 2005. p. 457-473.
3. Bethencourt J, Sahai A, Waters B. Ciphertext-policy attribute-based encryption. In: Proceedings of the IEEE Symposium on Security and Privacy (S&P); 2007. p. 321-334.
4. Natsos D, Symeonidis AL. Towards a defense-in-depth approach for securing collaborative cloud infrastructures. In: Software Engineering and Advanced Applications (SEAA 2025). Lecture Notes in Computer Science. Vol. 16083. Cham: Springer; 2025. p. 409-418. doi:10.1007/978-3-032-04207-1_27.
5. Anand C. Exploring practical considerations and applications for privacy enhancing technologies. ISACA White Paper. 2024.
6. Apple Security Research. Private Cloud Compute: A new frontier for AI privacy in the cloud. Apple Security Blog.

2024. Available from:
<https://security.apple.com/blog/private-cloud-compute/>
7. Zhan S, Huang L, Luo G, Zheng S, Gao Z, Chao HC. A review on federated learning architectures for privacy-preserving AI: Lightweight and secure cloud-edge-end collaboration. *Electronics*. 2025;14(13):2512.
 8. Troncoso-Pastoriza JR, Mermoud A, Bouyé R, Marino F, Bossuat JP, Lenders V, et al. Orchestrating collaborative cybersecurity: A secure framework for distributed privacy-preserving threat-intelligence sharing. *arXiv*. 2022;2209.02676.

Creative Commons License

This article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution–Non-commercial–No Derivatives 4.0 International (CC BY-NC-ND 4.0) License. This license permits users to copy and redistribute the material in any medium or format for non-commercial purposes only, provided that appropriate credit is given to the original author(s) and the source. No modifications, adaptations, or derivative works are permitted.

About the Author



Vaibhav Bhosle is associated with the Department of Computer Science and Technology at Shivaji University, Kolhapur, Maharashtra, India. His academic interests include computer science, emerging technologies, data privacy, cloud computing, cybersecurity, and privacy-preserving systems. He is engaged in academic and research activities focused on technological innovation and contemporary challenges in computing.