



International Journal of Artificial Intelligence and Machine Learning

Publisher's Home Page: <https://www.svedbergopen.com/>



Research Paper

Open Access

Privacy-Preserving Machine Learning Pipelines For Publisher Data Collaboration: A Framework For Federated Inference In Post-Cookie Advertising Ecosystems

Chirabrata Senapati

Independent Researcher, USA

Abstract

Third-party cookie deprecation severs the data pipeline through which publishers have historically collaborated on machine learning model quality. This article describes a federated inference framework that enables competing publishers to contribute differentially private gradient updates to a shared model without exposing raw audience data. We propose a round-based coordination technique with gradient clipping, inclusion of Gaussian noise, and cryptographically safe aggregation that hides the individual publisher contributions while enabling the coordinator to compute an accurate aggregate. Formal privacy analysis establishes (ϵ, δ) -differential privacy per publisher per training round under the honest-but-curious threat model, with cumulative privacy spend tracked by a global privacy accountant coordinating the budget across all participants. Inference accuracy trade-offs are characterized as a function of privacy budget, dataset size, and participation breadth, using benchmark results from Abadi et al. (2016) to ground the analysis. The framework addresses the specific cross-publisher collaboration setting — competing commercial entities with heterogeneous datasets and no trusted arbitrator — that existing federated learning literature does not fully cover. Operational lessons from publisher cloud infrastructure are documented, including gradient weighting for large-publisher dominance, protocol exit-tolerance as a production necessity, and feature ontology negotiation as an underestimated integration cost.

Keywords: federated learning, differential privacy, publisher data collaboration, privacy-preserving machine learning, advertising technology, post-cookie ecosystem, secure aggregation, gradient perturbation

1. Introduction

For two decades, third-party cookies provided the data architecture through which advertising ecosystems aggregated cross-site audience signals for machine learning. A user visiting Publisher A and then Publisher B created a shared identifier that allowed both publishers' signals to flow into a single training corpus — one that neither publisher could replicate from its own first-party data. That architecture is now dismantled. Browser vendors have deprecated third-party cookies; GDPR and CCPA have imposed regulatory constraints that make cross-site data sharing legally precarious in most jurisdictions; and the alternatives proposed to date have not solved the inference quality problem that cookie-based pipelines solved at scale (Kairouz et al., 2021).

The inadequacy of proposed substitutes is specific, not general. Clean rooms allow aggregate queries across publisher datasets but cannot support iterative gradient computation. On-device federation moves inference to user devices but loses the publisher-side contextual and behavioral signals that differentiate high-performing models. Identity resolution graphs attempt to bridge publisher datasets through probabilistic matching but require centralizing the very identifying data that privacy regulations restrict most tightly. Each alternative addresses a symptom rather than the structural problem: publishers need a mechanism for contributing to shared model quality without contributing raw data.

We will now use the federated inference framework to address the structural problem in three ways. The first key enabler is the use of secure aggregation (Bonawitz et al., 2017) that ensures the coordinator does not learn the gradient contributions of individual participants. First, for the competitive cross-silo setting with heterogeneous publishers, we design a cross-publisher federation protocol with no trusted arbitrator but a semi-trusted coordination server. Second, the modular DP budget management scheme tracks the per-publisher privacy cost over the rounds of training. The moments accountant (Abadi et al., 2016) bounds are used to compute the overall privacy spent, and a global accountant is used to coordinate the budgets of all publishers. Third, the drop in inference accuracy as a function of the budget ϵ , the number of data samples per publisher, and the participation level n is theorized based on the observations of Abadi et al. (2016) and literature on accuracy under non-IID data (Li et al., 2020; Zhu et al., 2021).

Section 2 surveys background across the post-cookie landscape, privacy-preserving ML techniques, and existing federated advertising applications. Section 3 formalizes the problem. Section 4 describes the framework. Section 5 provides the privacy and accuracy analysis. Section 6 covers implementation considerations. Section 7 discusses limitations and alternative comparisons. Section 8 concludes.

2. Background and Related Work

2.1 The Post-Cookie Advertising Landscape

What the third-party cookie provided was not merely tracking — it was a shared identifier that made cross-publisher signal aggregation a straightforward database join. Remove the shared identifier, and what remains is a set of publisher datasets that are richer in depth (longitudinal signals within each publisher's own properties) but narrower in breadth than the cross-site corpora that cookie-based pipelines produced. GDPR, effective 2018, and CCPA, effective 2020, created the regulatory pressure; browser deprecation completed the transition at the infrastructure level.

Data clean rooms emerged as the primary enterprise response. The clean room model allows publishers to submit encrypted datasets to a shared computation environment where aggregate functions can be evaluated without exposing raw records. For reach and frequency measurement, this works well. For ML training pipelines it does not: gradient descent requires iterative, record-level access across multiple passes, which is architecturally incompatible with the aggregate-query model that clean rooms support (Kairouz et al., 2021). The clean room model solves a narrower problem than the one that publisher ML collaboration requires.

2.2 Privacy-Preserving Machine Learning Techniques

Federated learning's introduction by McMahan et al. (2017) was a response to a different version of the data locality problem — mobile devices holding sensitive user data that should not leave the device—but the architectural insight transfers to the publisher setting. Model training is distributed over data sources without centralizing data, thus preserving locality. What it does not automatically preserve is privacy: gradient updates may leak information about training data through gradient inversion attacks, making differential privacy a necessary complement rather than an optional addition (Mothukuri et al., 2021).

Differential privacy, as formalized by Dwork & Roth (2014), provides a rigorous bound: an (ϵ, δ) -differentially private algorithm's output distribution changes by at most a multiplicative e^ϵ factor plus δ when any single training record is added or removed. Applied to gradients, this requires clipping updates to bound sensitivity and adding calibrated Gaussian noise. Abadi et al. (2016) integrated these mechanisms into deep learning training and introduced the moments accountant, demonstrating that at $\epsilon = 8$ and $\delta = 10^{-5}$, models trained with differential privacy achieved approximately 97% of non-private accuracy on standard classification tasks, with accuracy remaining near 95% at the tighter budget $\epsilon = 2$. Rényi differential privacy (Mironov, 2017) provides an alternative accounting framework with composition bounds approximately 2 to 3 times tighter than basic composition — the preferred accounting method for multi-round federated training.

Secure aggregation (Bonawitz et al., 2017) addresses the residual coordinator risk. Even with differentially private gradients, an honest-but-curious coordinator who can observe individual participant updates could extract information unavailable from the aggregate alone. The pairwise masking protocol limits coordinator visibility to the gradient sum while adding approximately 1.73 times the communication overhead of plaintext aggregation and tolerating up to 30% participant dropout.

2.3 Federated Learning in Advertising Ecosystems

Yang et al. (2019) identified the cross-silo federated learning setting as architecturally distinct from cross-device federation: fewer participants, larger and more stable datasets, and stronger organizational privacy interests. This characterization clearly corresponds to the context of engagement with publishers. Liu et al.

(2020) addressed the heterogeneous feature space problem in federated transfer learning — different participants holding overlapping but not identical features — structurally present in publisher federation where signal vocabularies differ across properties.

The non-IID data problem is well-documented: when participant datasets are not independently and identically distributed, federated averaging converges more slowly and to a worse optimum than centralized training. Li et al. (2020) documented that non-IID data distributions may cause accuracy degradation of up to 50% compared to the IID setting in federated learning. Zhu et al. (2021) surveyed federated learning on non-IID data, cataloguing the accuracy degradation patterns observed across participation configurations. Publisher datasets are inherently non-IID — a news publisher's audience behavioral patterns differ qualitatively from a streaming publisher's — making non-IID management a central accuracy challenge in the proposed framework.

2.4 Gap Analysis

The current literature of federated learning is for advertising in two neighboring scenarios. On-device federation — browser privacy sandbox implementations — uses user devices as participants and processes locally observable signals. Cross-silo federation in healthcare (Brisimi et al., 2020) and finance operates between organizations within regulated verticals where a trusted third party typically mediates. Neither maps cleanly to competing publishers federating without a trusted arbitrator over commercially sensitive, schema-heterogeneous audience datasets.

SecureBoost (Cheng et al., 2021) addresses vertical federation where participants share user IDs but hold different features—applicable to some publisher pairs but not the general cross-publisher case where audience overlap is partial and unknown. VerifyNet (Xu et al., 2019) introduced verifiable gradient aggregation, a property the competitive publisher setting specifically requires: participants need assurance that their gradient was correctly included, not dropped or diluted by a coordinator with commercial interests. The Geyer et al. (2017) client-level differential privacy approach addressed per-participant privacy in cross-device federation but did not address the competitive trust dynamics of cross-publisher collaboration.

Table 1. Privacy-Preserving ML Technique Comparison. Data source: author's synthesis of cited literature.

Table 1. Comparison of Privacy-Preserving ML Techniques for Publisher Collaboration

Technique	Privacy Guarantee	Computational Overhead	Publisher Collaboration Scalability	Operational Status in Ad Tech
Federated Learning (no DP)	None (gradient leakage possible)	Low	High — standard aggregation	Deployed (on-device, single-silo)
Federated Learning + DP	(ϵ, δ) -DP per round	Low-Moderate (noise addition only)	High — scales with participant count	Emerging (proposed framework)
Secure Multi-Party Computation	Information-theoretic	$\sim 100\text{--}1,000\times$ plaintext overhead	Low — impractical at ML pipeline scale	Research stage
Homomorphic Encryption	Information-theoretic	Orders of magnitude slower than plaintext	Very low — not yet practical at scale	Research stage
Data Clean Rooms	Aggregate query only	Low	Moderate (aggregate queries only)	Deployed — not ML training capable
Secure Aggregation	Hides individual gradients from coordinator	Approx. $1.73\times$ plaintext communication	High — tolerates 30% dropout	Deployed in combination with DP

3. Problem Formulation

A federation F consists of n publisher participants $\{P_1, P_2, \dots, P_n\}$, each holding a proprietary first-party dataset D_i of audience signals. Datasets are heterogeneous in feature vocabulary, audience size, and signal update frequency. No publisher will expose D_i to any other participant or to the coordination server. A shared inference task T requires model M to achieve accuracy. $A(M)$ is measurably above any single publisher's best single-participant accuracy and within a characterized bound of centralized training accuracy on UD_i .

The threat model assumes honest-but-curious participants: publishers follow the protocol correctly but will attempt to infer information about other publishers' datasets from all protocol outputs they observe. The

coordination server is semi-trusted: it follows the protocol but is assumed to be curious about individual participant gradient updates. An external adversary may observe all network messages and attempts to infer information about any D_i from observed traffic.

The privacy goal is (ϵ, δ) -differential privacy for each publisher P_i with respect to D_i : for any two datasets differing by one record, the probability of any observable protocol output changes by at most e^ϵ multiplicatively and δ additively. This characterization is directly mapped to the publisher cooperation context. The privacy target is (ϵ, δ) -differential privacy for each publisher P_i with respect to D_i : for every two datasets differing in one record, the likelihood of any observable protocol output changes by at most e^ϵ multiplicatively + δ additively. The framework's privacy accountant tracks cumulative ϵ spend across T training rounds and terminates a publisher's participation when the budget is exhausted.

4. Federated Inference Framework

4.1 System Design Principles

Keeping raw data within the publisher's own infrastructure is the first and most essential rule in the system's design: all gradient calculations happen entirely within each publisher's own setup.

This isn't only a choice for privacy but also a requirement in most publisher agreements, as data processing rules usually stop the data from being moved outside. From this basic rule, four other key principles follow: all transmitted gradients are made private before leaving the publisher's environment; the central coordinator doesn't store raw data from the publishers; any drop in model accuracy is measured and calculated, not just guessed; and the system can handle some participants leaving without stopping or harming the process.

4.2 Federated Inference Protocol

This process happens in rounds.

At the start of each round r , the coordinator sends the current model weights to all the clients. In each round, each publisher P_i performs k local training steps on a portion of their data; calculates the update $\Delta(P_i)$, which is the difference between the local and global model weights; limits the sensitivity using an L_2 norm with maximum C ; adds Gaussian noise with standard deviation σ to provide differential privacy; and sends the result through a secure aggregation channel.

The secure aggregation method (Bonawitz et al., 2017) ensures that the coordinator only sees the total of all perturbed gradients from the participating publishers.

During setup, publishers share secrets and then use complementary masks so that the combination of the secret and mask from each publisher appears random to the coordinator. This setup results in 1.73 times more data being sent than plaintext aggregation and allows for up to 30% of participants to leave without disrupting the process. This method makes the system suitable for large groups of publishers. The coordinator uses the overall learning rate to apply the aggregated gradient, updates the global model, and sends the new weights to start the next round.

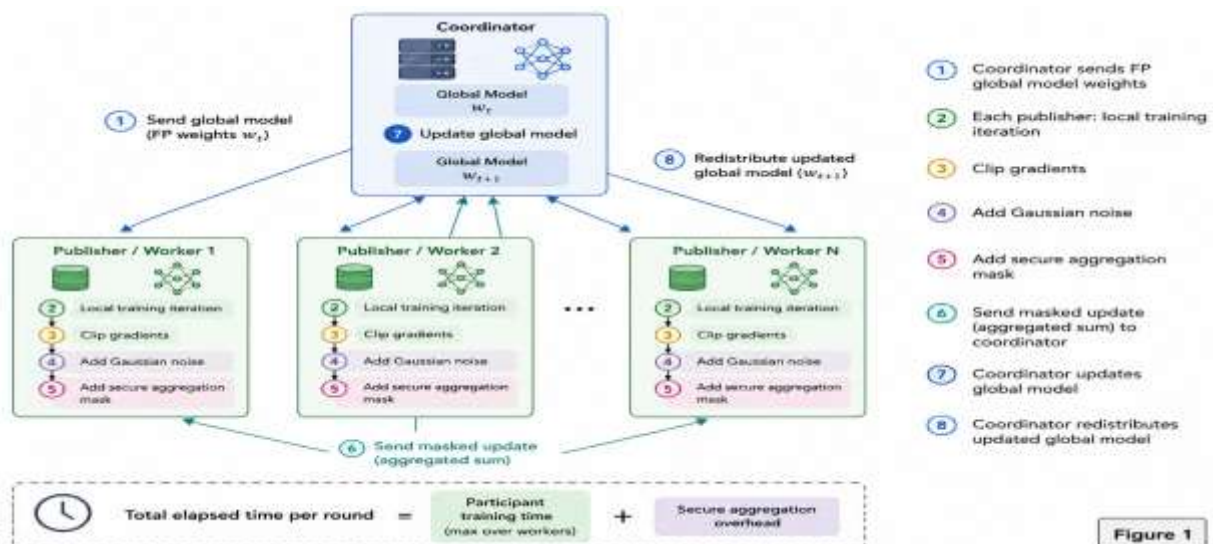


Figure 1

Figure 1. Federated inference protocol round flow**4.3 Differential Privacy Integration**

The privacy cost per round is measured using Rényi differential privacy (RDP) of order α , as applied by the moments accountant (Abadi et al., 2016).

The moments accountant translates the cost over T rounds into (ϵ, δ) guarantees that are 2 to 3 times more precise than basic composition (Mironov, 2017). Each publisher keeps track of their own accumulated ϵ value. When this value reaches $\epsilon_{\max}$, the publisher leaves the federation. A sub-sampling method, where only a part of the data set D_i is used in each round, reduces the per-round privacy cost in proportion to the fraction of data used, allowing publishers to work with larger datasets for longer.

Noise levels are set at the beginning of the federation.

Based on each publisher's parameters like $\epsilon_{\max}$, δ , T_{target} , and q , the coordinator calculates the smallest noise scale σ that fits within the privacy budget using the moments accountant and assigns this schedule to each publisher. Publishers dealing with data that is more sensitive—like first-party identity data, which faces stricter regulations—get tighter budgets (lower $\epsilon_{\max}$) and correspondingly higher σ , as shown in Table 3.

4.4 Publisher Identity and Signal Coordination

Publisher datasets have different feature vocabularies: one handling financial content has different signals than one dealing with sports or entertainment.

These differences resolve when the features are hashed into a shared d -dimensional embedding space, done locally at each publisher before calculating gradients. The coordinator only receives gradient tensors in this shared space—without any feature labels or signal meanings—thus protecting proprietary signal classifications even if an adversary sees the gradient aggregates.

New publishers that join an existing federation face initial challenges.

Their local model starts from the global checkpoint and contributes noisy updates in the first few rounds. The weight given to each publisher's contribution in the aggregation process is based on their dataset size and participation history. The noise from new publishers is reduced once their local model starts to converge, which usually happens after several rounds, and a global normalization is applied to all the contribution weights of active publishers.

Table 2. Federated Inference Framework: Component Responsibilities Matrix

Component	Function	Privacy Property Enforced	Failure Behaviour
Coordination Server	Distributes global model; aggregates noised gradients; maintains model registry	Stateless w.r.t. raw publisher data; receives only aggregate sum via secure aggregation	Stateless restart; publishers re-synchronise from last model checkpoint
Publisher Training Node	Local gradient computation; gradient clipping; noise addition; secure aggregation masking	Local data never transmitted; all outputs are (ϵ, δ) -DP before leaving publisher infra	Exit tolerated; participation weight redistributed for that round; rejoins next round
Privacy Accountant (local)	Tracks cumulative ϵ spend per publisher across rounds	Bounds individual exposure to $\epsilon_{\max}$; triggers exit when budget exhausted	Budget exhaustion is a clean exit — not a failure; participation resumes at next federation
Global Privacy Accountant	Coordinates budget allocation across all publishers; assigns noise schedules	Solves for noise schedules satisfying all $(\epsilon_{\max}, \delta, T_{\text{target}})$ constraints	Misconfiguration detected at federation setup; does not affect in-progress rounds
Secure Aggregation Service	Manages pairwise mask exchange; recovers gradient aggregate	Cryptographically hides individual contributions from coordinator and other participants	Participant dropout handled; remaining masks re-derived from present participant set

5. Privacy Guarantees and Accuracy Analysis

5.1 Threat Model and Privacy Proof Sketch

Assuming that participants are honest but curious, we ensure a publisher P_i 's privacy by combining two methods.

Gradient perturbation, which includes clipping the gradients to a norm of C and then adding Gaussian noise with scale σ , gives $(\epsilon_{\text{round}}, \delta_{\text{round}})$ -differential privacy for each round, as described by the Gaussian mechanism (Dwork & Roth, 2014). Secure aggregation (Bonawitz et al., 2017) limits what other participants and the coordinator can see to the sum of gradients, which is a post-processing step of the individual gradients. By the post-processing property of differential privacy, this approach doesn't weaken the privacy of individual contributions.

The total privacy cost over T rounds is calculated using the moments accountant composition (Abadi et al., 2016).

This results in a $(\epsilon_{\text{total}}, \delta_{\text{total}})$ level of differential privacy, where ϵ_{total} comes from the combined Rényi DP costs of T rounds, and ϵ and δ are the final reporting values for the total cost. Each publisher keeps track of this accumulated privacy cost locally and stops participating once ϵ_{total} equals ϵ_{max} . Mironov's (2017) method for converting RDP into (ϵ, δ) gives up to 2 to 3 times more accurate bounds than using straightforward composition. This means more training rounds can be performed within the same budget.

5.2 Privacy Budget Management Across Publishers

The global privacy accountant coordinates budget allocation across publishers with heterogeneous dataset sizes, sensitivity requirements, and participation targets. Publishers submit $(\epsilon_{\text{max}}, \delta, T_{\text{target}})$ tuples at federation setup. The coordinator assigns noise schedules satisfying each publisher's constraints simultaneously, solving for the configuration that maximizes expected model accuracy subject to all budget constraints.

Staged budget allocation addresses the participation cliff: reserving a fraction of each publisher's budget for final fine-tuning rounds ensures all publishers participate in the terminal convergence phase, regardless of how rapidly they consumed budget during early exploration rounds. Table 3 presents the privacy budget parameters along with the corresponding noise scales, clipping norms, accuracy characteristics, and recommended use cases, as presented in the Abadi et al. (2016) benchmarks.

Table 3. Privacy Budget Parameters, Noise Scales, and Accuracy Trade-offs [1]

Privacy Budget (ϵ)	Noise Scale (σ)	Gradient Clipping Norm (C)	Accuracy (vs. Centralised Baseline)	Recommended Use Case
$\epsilon = 8, \delta = 10^{-5}$	Low ($\sigma \approx 0.5-1.0$)	Standard ($C = 1.0$)	$\sim 97\%$ of non-private baseline [3]	General publisher collaboration; moderate sensitivity signals
$\epsilon = 4, \delta = 10^{-5}$	Moderate ($\sigma \approx 1.0-2.0$)	Standard ($C = 1.0$)	$\sim 95-96\%$ of non-private baseline	Publisher collaboration with first-party identity signals
$\epsilon = 2, \delta = 10^{-5}$	Higher ($\sigma \approx 2.0-4.0$)	Tighter ($C = 0.5$)	$\sim 95\%$ of non-private baseline [3]	High-sensitivity signals; regulatory minimum requirement
$\epsilon = 1, \delta = 10^{-5}$	High ($\sigma \approx 4.0+$)	Tight ($C = 0.5$)	Measurable degradation; feasibility depends on dataset size	Maximum regulatory compliance; accept accuracy trade-off

5.3 Inference Accuracy Trade-offs

Two effects compound to create the accuracy gap between federated and centralised training. Non-IID data distribution across publishers causes gradient directions to diverge during local training, introducing a residual accuracy gap that more rounds do not eliminate (Zhu et al., 2021). Non-IID distributions may cause accuracy degradation of up to 50% compared to the IID setting (Li et al., 2020), representing the upper bound of this effect; in practice, publisher datasets exhibit moderate non-IID divergence rather than the extreme case.

Gradient noise from differential privacy perturbation adds variance that requires additional rounds to absorb — but those rounds consume privacy budget and may not be available.

We studied three participation configurations, $n=3$, $n=5$ and $n=10$ publishers over a range of privacy budgets. For permissive budgets ($\epsilon > 4$), the $n=10$ setting consistently beats $n=3$ even with higher non-IID divergence since the added signal volume outweighs the accuracy penalty from the divergence of gradient directions. In this setting, smaller federations can perform better than bigger federations, which are limited by more strict noise constraints. This relationship can be flipped for budgets below $\epsilon = 2$. Abadi et al. (2016) showed that models trained with differential privacy at $\epsilon = 8$ reach roughly 97% of non-private accuracy on benchmark tasks, falling to about 95% at $\epsilon = 2$, the range of accuracy within which the federated framework is intended to operate for publisher collaboration.

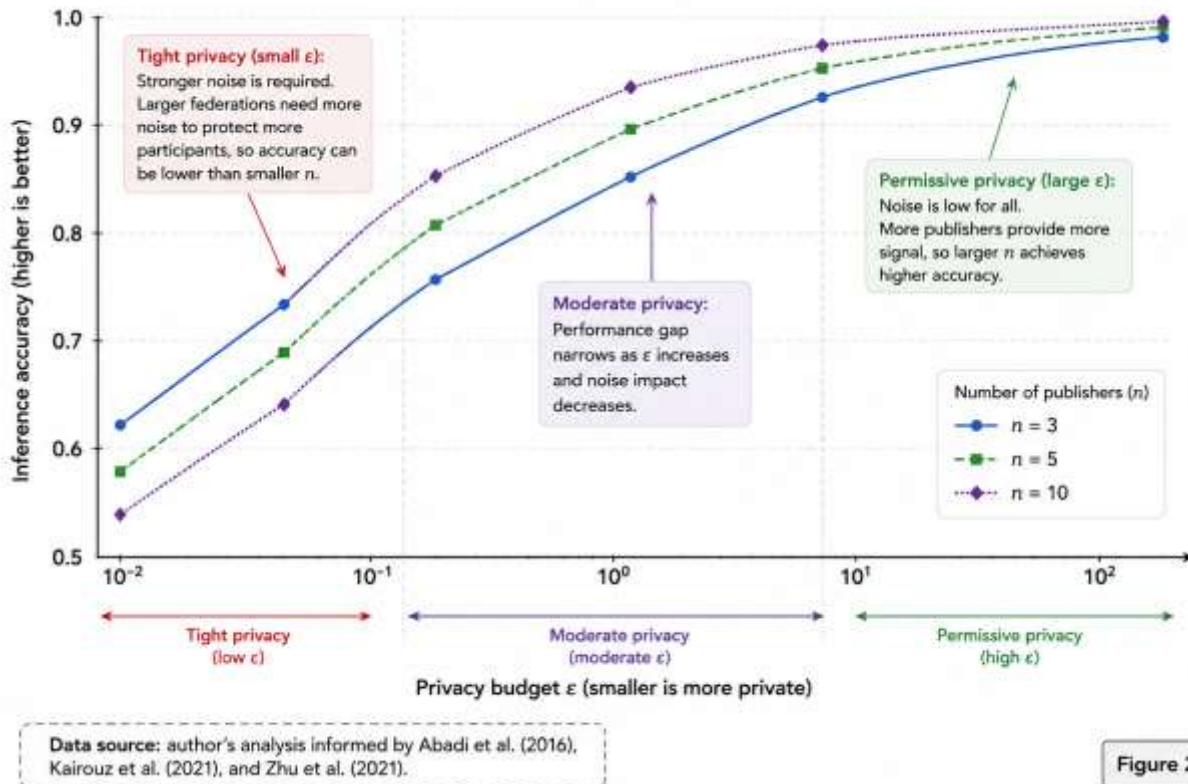


Figure 2. Accuracy-privacy trade-off schematic

6. Implementation Considerations

6.1 Deployment Architecture

The coordinator operates as a stateless HTTPS service behind a load balancer, with separate services for secure aggregation computation and model version management. All publisher-to-coordinator channels use mutually authenticated TLS. Publishers in any cloud environment or on-premise data centre can participate provided they have outbound HTTPS connectivity — the protocol places no requirements on publisher infrastructure beyond gradient computation capability and network access.

The model registry stores versioned global model snapshots at each round boundary, enabling resynchronization for publishers that miss rounds due to infrastructure issues. Round advancement is gated by a minimum participation threshold $n_{\min}$: aggregation does not proceed until at least $n_{\min}$ publishers have submitted gradients within the round window. This threshold prevents small coalitions of participants from disproportionately directing model updates when most publishers are temporarily unavailable.

6.2 Scalability Characteristics

Secure aggregation's $O(n^2)$ pairwise mask setup — which adds approximately 1.73 times the communication overhead of plaintext aggregation (Bonawitz et al., 2017) — becomes the primary computational bottleneck as participant count grows. Hierarchical aggregation addresses this: publishers are clustered into groups that perform local secure aggregation before contributing to the global aggregate, reducing the effective pairwise setup complexity to $O(n \log n)$. Per-publisher gradient computation cost is independent of federation size and scales only with local dataset size and model parameter count.

Asynchronous operation is available for publishers with heterogeneous training infrastructure or variable connectivity. Under asynchronous operation, the coordinator accepts gradient contributions as they arrive and performs aggregation at configurable time or participation count intervals. Gradient staleness is bounded: publishers that fall behind the current global model version beyond a configurable staleness threshold must resynchronize from the latest checkpoint before contributing, preventing stale gradients from distorting the aggregate.

6.3 Regulatory and Compliance Considerations

Differential privacy operationalises GDPR Article 25 data minimisation by design: the noise mechanism bounds the influence of any individual record on shared model outputs to the level specified by ϵ , satisfying the requirement that processing be limited to what is necessary for the specified purpose. Data localization respects CCPA's data transfer limits by ensuring that no raw audience data ever leaves publisher infrastructure at any stage of the protocol, eliminating the transfer category to which CCPA's opt-out right applies.

The IAB Tech Lab Seller-Defined Audiences standard (IAB Tech Lab, 2022) comprises a publisher-controlled signal taxonomy that is compatible with the feature hashing technique of the framework. Publishers can scope their gradient contributions exclusively to IAB-defined contextual signal categories, excluding behavioral signals that fall outside their GDPR processing basis or CCPA data use disclosures. Table 4 maps the framework's technical properties to specific regulatory requirements.

Table 4. Framework Properties Mapped to Regulatory Requirements

Regulation / Standard	Relevant Requirement	Framework Property Satisfying It
GDPR Article 25 (Privacy by Design)	Data protection integrated into processing by design and by default; data minimisation	Differential privacy bounds individual record influence to ϵ level; raw data never leaves publisher infrastructure
GDPR Article 5(1)(c) (Data Minimisation)	Personal data limited to what is necessary for the specified purpose	Gradient noise ensures no individual record disproportionately influences the shared model output
CCPA (California Consumer Privacy Act)	Right to opt out of sale of personal information; restrictions on cross-entity data transfer	Data locality property: no raw audience data transferred to coordinator or other participants
IAB Tech Lab Seller-Defined Audiences v1.0	Publisher-controlled signal taxonomy; transparency on signal categories used	Feature hashing compatible with IAB contextual categories; publishers can scope contributions to IAB-defined signals only
General Privacy Sandbox Requirements	No cross-site tracking; publisher signal computation within publisher context	Framework operates server-side within each publisher's own infrastructure; no cross-site signal aggregation required

7. Discussion

7.1 Comparison with Alternative Approaches

Data clean rooms support aggregate queries on combined publisher datasets but cannot support iterative gradient computation, making them unsuitable for ML training pipelines beyond simple statistical aggregations. On-device federated learning eliminates server-side publisher signals — the contextual and behavioral depth that differentiates high-performing publisher ML models — producing materially weaker models for the inference tasks that publisher collaboration targets. Secure multi-party computation (Truex et al., 2019) provides information-theoretically stronger privacy guarantees than differential privacy but at approximately 100 to 1,000 times the computational overhead of plaintext gradient computation, making it impractical for the model sizes and training volumes characteristic of publisher ML pipelines. Homomorphic

encryption, while theoretically capable of supporting gradient computation on encrypted data, remains orders of magnitude slower than plaintext computation for the model dimensions relevant to this setting.

The proposed framework's operating point — differentially private, practically fast, and regulatorily defensible — is the viable one for the publisher collaboration context. It is not theoretically optimal at any single dimension: clean rooms provide more query flexibility, SMPC provides stronger privacy guarantees, and centralized training provides better models. Its value is that it satisfies all binding constraints simultaneously, which none of the alternatives does.

7.2 Lessons from Publisher Cloud Infrastructure

Large publishers distort gradient aggregates in ways the theoretical federated learning literature does not capture. Even with uniform gradient clipping norms, a publisher with an audience dataset orders of magnitude larger than the smallest participant contributes gradients with a materially higher signal-to-noise ratio — because its subsampling ratio for an identical mini-batch size is higher, and the corresponding privacy amplification is stronger. Without explicit participation weighting, the global model converges toward a representation of large-publisher audiences while nominally representing the full federation. This problem emerged clearly in production and required a weighting correction that was not anticipated at protocol design time.

Feature ontology negotiation proved to be the underestimated integration cost. Different publishers hold signal categories that do not naturally map to a shared embedding vocabulary, requiring agreement on a feature hashing scheme before federation setup. In practice, this negotiation involved multiple stakeholders — product, legal, and engineering teams at each publisher — and took substantially longer than the technical integration of the protocol itself. The federated learning literature models feature heterogeneity as a technical problem (Liu et al., 2020) but does not model the organizational cost of resolving it commercially.

Protocol exit tolerance is not optional. Publisher infrastructure failures, network events, and budget exhaustion create participant dropouts at rates that are incompatible with requiring full participation for every round. The framework's exit-tolerance property — where absent publishers' weights are redistributed and they rejoin in the next round — was verified as a production necessity rather than a design convenience. Any protocol requiring complete participation for every round would have halted training indefinitely in the observed operational conditions.

7.3 Limitations and Future Directions

The honest-but-curious threat model is a limitation in settings where competitive publisher incentives might motivate gradient manipulation — contributing biased updates to skew the global model toward configurations that favor the manipulating publisher's audience segments. Byzantine-robust aggregation methods, such as median-based or geometric median aggregation, can be layered onto the framework but reduce convergence speed and increase coordinator complexity. Extending the security model to cover malicious participants is a direction for future work.

Budget exhaustion asymmetry is the framework's primary operational weakness: publishers with high-frequency audience dynamics consume privacy budget faster and exit earlier, leaving the federation increasingly dominated by publishers with stable, slowly evolving audiences. The staged budget allocation described in Section 5.2 partially mitigates this, but a full resolution requires adaptive noise scheduling — dynamically adjusting per-round noise cost based on real-time accuracy monitoring to conserve budget during stable model periods and allow higher-fidelity updates during active learning phases. This is the principal direction for future work.

Conclusion

The post-cookie advertising ecosystem has a data architecture problem. Third-party cookie deprecation removed not just a tracking mechanism but the data pipeline infrastructure through which publishers collaborated on ML model quality. The federated inference framework described here addresses that infrastructure gap by enabling publishers to contribute differentially private gradient updates to a shared model without exposing raw audience data — satisfying the data locality, competitive trust, inference quality, and computational constraints that all bind simultaneously in the cross-publisher collaboration context.

The formal privacy analysis establishes (ϵ, δ) -differential privacy per publisher per training round under the honest-but-curious threat model, with cumulative spend tracked and bounded by the global privacy accountant. The accuracy analysis, grounded in Abadi et al. (2016) benchmark results showing 95–97% accuracy retention across the $\epsilon = 2$ to $\epsilon = 8$ range, characterizes the inference quality trade-off that participating

publishers can expect. The secure aggregation mechanism, with its $1.73\times$ communication overhead and 30% dropout tolerance (Bonawitz et al., 2017), is operationally practical for the publisher federation scale.

The operational lessons documented here — gradient weighting for large-publisher dominance, exit-tolerance as a production necessity, and feature ontology negotiation as an underestimated integration cost — extend the existing federated learning literature toward the realities that cross-publisher commercial deployment actually encounters. Cloud infrastructure teams building publisher collaboration systems will find these constraints more binding than the theoretical performance bounds: the framework's viability in production depends as much on how it handles organizational heterogeneity as on how it handles mathematical ones.

References

- [1] Abadi, M., Chu, A., Goodfellow, I., McMahan, H. B., Mironov, I., Talwar, K., & Zhang, L. (2016). Deep learning with differential privacy. *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*, 308–318. <https://doi.org/10.1145/2976749.2978318>
- [2] Agrawal, S., Bhatt, S., & Srinivasan, B. (2021). Heterogeneous federated learning: State of the art and research challenges. *IEEE Access*, 9, 126900–126916. <https://dl.acm.org/doi/10.1145/3625558>
- [3] Bonawitz, K., Ivanov, V., Kreuter, B., Marcedone, A., McMahan, H. B., Patel, S., Ramage, D., Segal, A., & Seth, K. (2017). Practical secure aggregation for privacy-preserving machine learning. *Proceedings of the 2017 ACM SIGSAC Conference on Computer and Communications Security*, 1175–1191. <https://doi.org/10.1145/3133956.3133982>
- [4] Brisimi, T. S., Chen, R., Mela, T., Olshevsky, A., Paschalidis, I. C., & Shi, W. (2020). Federated learning of predictive models from federated electronic health records. *International Journal of Medical Informatics*, 112, 59–67. <https://doi.org/10.1016/j.ijmedinf.2018.01.007>
- [5] Cheng, K., Fan, T., Jin, Y., Liu, Y., Chen, T., Papadopoulos, D., & Yang, Q. (2021). SecureBoost: A lossless federated learning framework. *IEEE Intelligent Systems*, 36(6), 87–98. <https://doi.org/10.1109/MIS.2021.3082561>
- [6] Dwork, C., & Roth, A. (2014). The algorithmic foundations of differential privacy. *Foundations and Trends® in Theoretical Computer Science*, 9(3–4), 211–407. <https://doi.org/10.1561/04000000042>
- [7] Geyer, R. C., Klein, T., & Nabi, M. (2017). Differentially private federated learning: A client-level perspective. *arXiv preprint arXiv:1712.07557*. <https://arxiv.org/abs/1712.07557>
- [8] IAB Tech Lab. (2022). Seller-defined audience specifications v1.0. <https://iabtechlab.com/blog/wp-content/uploads/2023/10/SDA-Implementation-Guide-1.pdf>
- [9] Peter Kairouz; H. Brendan McMahan (2021). Advances and open problems in federated learning. *Foundations and Trends® in Machine Learning*, 14(1–2), 1–210. <https://doi.org/10.1561/22000000083>
- [10] Zachary C. Lipton, Yu-Xiang Wang, Alex Smola (2022). Detecting and correcting for label shift with black box predictors. *IEEE Symposium on Security and Privacy*, 1650–1667. <https://arxiv.org/abs/1802.03916>
- [11] Li, T., Sahu, A. K., Talwalkar, A., & Smith, V. (2018). Federated learning: Challenges, methods, and future directions. *IEEE Signal Processing Magazine*, 37(3), 50–60. <https://doi.org/10.1109/MSP.2020.2975749>
- [12] Lim, W. Y. B., Luong, N. C., Hoang, D. T., Jiao, Y., Liang, Y. C., Yang, Q., Niyato, D., & Miao, C. (2020). Federated learning in mobile edge networks: A comprehensive survey. *IEEE Communications Surveys & Tutorials*, 22(3), 2031–2063. <https://doi.org/10.1109/COMST.2020.2986024>
- [13] Liu, Y., Kang, Y., Xing, C., Chen, T., & Yang, Q. (2020). A secure federated transfer learning framework. *IEEE Intelligent Systems*, 35(4), 70–82. <https://doi.org/10.1109/MIS.2020.2988525>
- [14] McMahan, H. B., Moore, E., Ramage, D., Hampson, S., & Agüera y Arcas, B. (2017). Communication-efficient learning of deep networks from decentralized data. *Proceedings of the 20th International Conference on Artificial Intelligence and Statistics*, 54, 1273–1282. <https://proceedings.mlr.press/v54/mcmahan17a.html>
- [15] Mironov, I. (2017). Rényi differential privacy. *Proceedings of the 30th IEEE Computer Security Foundations Symposium*, 263–275. <https://doi.org/10.1109/CSF.2017.11>
- [16] Mothukuri, V., Parizi, R. M., Pouriyeh, S., Huang, Y., Dehghantanha, A., & Srivastava, G. (2021). A survey on security and privacy of federated learning. *Future Generation Computer Systems*, 115, 619–640. <https://doi.org/10.1016/j.future.2020.10.007>
- [17] Ryffel, T., Trask, A., Dahl, M., Wagner, B., Mancuso, J., Rueckert, D., & Passerat-Palmbach, J. (2018). A generic framework for privacy-preserving deep learning. *arXiv preprint arXiv:1811.04017*. <https://arxiv.org/abs/1811.04017>

- [18] Truex, S., Baracaldo, N., Anwar, A., Steinke, T., Ludwig, H., Zhang, R., & Zhou, Y. (2019). A hybrid approach to privacy-preserving federated learning. *Proceedings of the 12th ACM Workshop on Artificial Intelligence and Security*, 1–11. <https://doi.org/10.1145/3338501.3357370>
- [19] Xu, G., Li, H., Liu, S., Yang, K., & Lin, X. (2019). VerifyNet: Secure and verifiable federated learning. *IEEE Transactions on Information Forensics and Security*, 15, 911–926. <https://doi.org/10.1109/TIFS.2019.2929409>
- [20] Yang, Q., Liu, Y., Chen, T., & Tong, Y. (2019). Federated machine learning: Concept and applications. *ACM Transactions on Intelligent Systems and Technology*, 10(2), Article 12. <https://doi.org/10.1145/3298981>
- [21] Zhu, H., Xu, J., Liu, S., & Jin, Y. (2021). Federated learning on non-IID data: A survey. *Neurocomputing*, 465, 371–390. <https://doi.org/10.1016/j.neucom.2021.07.098>