



Small-Language Models for Privacy-Aware Institutional Knowledge Support

Dr. Adrian Keller

Associate Professor, University of Vienna, Austria

Abstract

Institutional knowledge is increasingly distributed across policy manuals, internal reports, project repositories, standard operating procedures, meeting records, technical documentation, research files, email archives, knowledge bases, and other digitally stored resources. Generative artificial intelligence can make such information easier to retrieve and interpret, but conventional cloud-centered language-model deployments may create privacy, confidentiality, access-control, and governance concerns when sensitive institutional information is transmitted to or processed by external systems. Small-language models offer an alternative architectural direction because their reduced computational requirements can support local, edge, or institution-controlled deployment for narrowly defined knowledge tasks. Smaller model size, however, does not itself guarantee privacy, factual reliability, or secure institutional use. This study develops a Privacy-Aware Institutional Knowledge Support Framework (PIKSF) combining small-language models with retrieval-augmented generation, access-aware retrieval, data minimization, local inference, source attribution, human oversight, and continuous security governance. Because no real institutional deployment dataset was provided, the study adopts an explicitly simulation-based methodology.

Four hypothetical architectures are compared: cloud large-language-model processing with raw context, hosted language models with filtered retrieval, local small-language models with governed retrieval, and privacy-aware small-language models with controlled institutional knowledge. Six dimensions—data exposure control, authorization alignment, retrieval grounding, response reliability, operational efficiency, and governance observability—are incorporated into a Privacy-Aware Knowledge Support Index. The simulated score rises from 42 for externally processed raw-context systems to 92 for privacy-aware small-model architectures with controlled knowledge retrieval. The analysis indicates that the principal privacy advantage of small-language models is architectural rather than intrinsic: institutions can reduce unnecessary information movement, preserve stronger control over inference, and separate model reasoning from authoritative institutional knowledge. The study concludes that small-language models are especially promising for bounded institutional knowledge-support tasks when combined with retrieval governance, source-level permissions, secure deployment, documented evaluation, and human accountability.

Keywords: small-language models, institutional knowledge, privacy-aware AI, retrieval-augmented generation, local inference, enterprise knowledge management, data minimization, responsible artificial intelligence



1. Introduction

Institutions generate substantial quantities of knowledge through policies, operating procedures, administrative records, research reports, technical manuals, project documentation, correspondence, meeting records, legal guidance, training materials, and internal databases. Accessing this information efficiently is often difficult because documents are distributed across systems, stored in heterogeneous formats, updated at different times, and written for different organizational purposes. Retrieval-augmented generative systems offer a practical way to connect natural-language interfaces with external document stores instead of expecting a language model to contain all institutional knowledge in its parameters. RAG architectures retrieve relevant external information at inference time and use it to ground generated responses, thereby addressing some problems associated with static training knowledge and unsupported generation. Research on private-document RAG also emphasizes that keeping institutional information outside model training can create a clearer separation between proprietary knowledge and the base model.

The adoption of generative AI for internal knowledge support nevertheless introduces substantial privacy and confidentiality concerns. Institutional prompts may contain personal information, commercially sensitive material, unpublished research, employee records, legal advice, security procedures, or strategically important operational knowledge. Retrieval systems can also create new attack surfaces through document ingestion, embedding stores, access-control failures, prompt injection, malicious content, or inappropriate retrieval across departmental boundaries. NIST's AI and privacy materials recognize that AI systems can create privacy risks through data aggregation, inference, reconstruction, prompt-based attacks, and other forms of information exposure. The NIST Generative AI Profile likewise recommends systematic risk management throughout the design, deployment, monitoring, and governance lifecycle rather than relying on one technical safeguard.

Small-language models provide a potentially important architectural option within this environment. Unlike extremely large general-purpose models, smaller models can be optimized for bounded tasks and may require substantially fewer computational resources, making local, edge, or on-premises deployment more feasible. Microsoft Research's Phi family demonstrates that comparatively compact models can achieve strong performance on selected reasoning and language tasks when training quality and post-training are carefully designed. Broader research on SLMs similarly identifies on-device processing, lower latency, personalization, and reduced dependence on remote infrastructure as important deployment advantages. These characteristics can support privacy-aware institutional design because raw organizational knowledge may remain within controlled infrastructure rather than automatically being transmitted to an external general-purpose model.

However, describing small-language models as “private” merely because they are smaller would be misleading. A locally deployed model can still reveal confidential information if access controls are weak, retrieval filters are incorrect, logs retain sensitive prompts, embeddings expose protected content, users can query information beyond their authorization, or malicious documents manipulate model behavior. Secure enterprise-RAG research demonstrates that retrieval and generation architecture must explicitly account for proprietary information, knowledge-source boundaries, and selective interaction with external models. Research on secure and privacy-aware RAG likewise identifies vulnerabilities



across ingestion, storage, retrieval, prompt construction, and generation rather than locating privacy risk in the language model alone.

The present study therefore investigates small-language models as components of a wider privacy-aware institutional knowledge architecture rather than as stand-alone privacy solutions. The paper develops a simulation-based framework integrating model size, deployment locality, retrieval grounding, institutional access control, data minimization, source provenance, human review, and lifecycle monitoring. Its central argument is that the strongest institutional design does not attempt to encode every organizational fact into the model. Instead, the language model should operate as a constrained reasoning and interaction layer over governed institutional knowledge stores whose access permissions, update processes, provenance, and retention policies remain independently manageable.

2. Objectives of the Study

2.1 To Develop a Privacy-Aware Institutional Knowledge Architecture

The first objective is to establish a framework in which small-language models support institutional search, explanation, summarization, and question answering without requiring unrestricted movement of sensitive organizational information into external generative-AI infrastructure. The framework emphasizes controlled inference, governed retrieval, source attribution, and separation between model capability and authoritative institutional data.

2.2 To Evaluate Privacy, Reliability, and Operational Trade-Offs

The second objective is to compare alternative knowledge-support architectures across privacy exposure, authorization alignment, retrieval grounding, response reliability, operational efficiency, and governance observability. This comparison recognizes that a privacy-preserving system that produces consistently unusable answers would not satisfy institutional needs, while a highly capable system that disregards confidentiality would be equally unsuitable.

2.3 To Identify Governance Requirements for Responsible SLM Deployment

The third objective is to determine which organizational controls are necessary to move from experimental small-model deployment toward accountable institutional use. Particular attention is given to role-based retrieval, document-level permissions, prompt and output logging, evaluation, incident response, human escalation, model updates, and source-level traceability.

3. Review of Literature

3.1 Small-Language Models and Localized AI Processing

Small-language models have attracted growing attention because model capability is no longer determined only by parameter scale. Microsoft Research's Phi-4 technical report describes a 14-billion-parameter model developed using a training strategy strongly focused on data quality, while the later Phi-4-reasoning work demonstrates how a model within this comparatively compact range can be specialized for reasoning-intensive tasks. These studies do not imply that smaller models are universally equivalent to frontier-scale systems, but they demonstrate that task-focused performance can be improved through careful data design, synthetic-data use, distillation, post-training, and specialization.



The practical importance of SLMs extends beyond benchmark scores. Reduced computational requirements make it more feasible to deploy language models on controlled institutional infrastructure, edge systems, or dedicated local devices. A recent survey of SLMs identifies privacy, response speed, personalization, and cloud independence as important motivations for on-device processing. Research examining collaboration between large and small models similarly distinguishes context-rich remote LLMs from specialized local SLMs that can reduce exposure of sensitive information.

These characteristics are particularly relevant for institutional knowledge support because many organizational questions are bounded rather than open-ended. Employees may ask which procedure applies to a particular workflow, which internal policy governs a request, how a technical process is documented, what evidence appears in a project archive, or where a specific institutional rule is stated. Such tasks may benefit less from unrestricted world knowledge than from high-quality retrieval over authoritative internal sources.

The architectural benefit of a small model therefore lies in task specialization plus deployment control. A model can be selected or tuned for internal question answering while the institutional corpus remains external to the model. Updates to policy documentation can then be reflected through retrieval-index changes rather than repeated model retraining.

This separation also helps reduce memorization pressure. Private-document RAG research explicitly highlights inference-time retrieval as a way to provide context without embedding all proprietary information into fine-tuned model parameters. EnronQA was developed partly because enterprise-style private retrieval had become important enough to require realistic benchmarking beyond public corpora.

3.2 Retrieval-Augmented Institutional Knowledge and Security

RAG is especially relevant to institutional knowledge because organizational information is dynamic. Policies change, projects evolve, personnel responsibilities shift, technical manuals receive revisions, and compliance requirements are updated. A pretrained model cannot reliably represent such changes unless information is supplied externally.

RAG systems retrieve relevant records and provide them to the generator during inference. This approach can improve grounding, provide source evidence, and permit knowledge updates without retraining. However, enterprise RAG introduces distinct security problems because retrieval itself can expose information.

Secure Multifaceted-RAG research addresses this problem by using a local open-source generator for internal knowledge and selectively involving external models only when a security filter determines that prompts are safe. In an automotive report-generation setting, the proposed architecture was designed specifically to improve completeness while reducing proprietary-data leakage concerns.

Other work goes further by treating privacy at the storage and retrieval layers. Privacy-aware RAG research has explored encrypting both textual content and embeddings, while provably secure RAG architectures propose encrypted storage and controlled access to retrieved knowledge. These approaches demonstrate that embedding databases should not automatically be treated as harmless technical indexes; they belong within the security boundary of the institutional knowledge system.



Distributed RAG research also demonstrates an alternative in which personal or private knowledge can remain local while general knowledge is handled separately. DRAGON, for example, was designed to combine cloud and device-side retrieval processes while limiting leakage of private documents. Such work reinforces the architectural principle that different knowledge classes need not be processed in the same environment.

Security challenges remain significant. RAG systems can be attacked through poisoned documents, adversarial prompts, unauthorized retrieval, manipulated metadata, malicious instructions embedded in files, or disclosure of sensitive retrieved content. Recent risk-assessment work maps vulnerabilities across preprocessing, storage, retrieval, and generation, showing that secure institutional deployment requires controls across the complete pipeline rather than only around the final model endpoint.

3.3 Privacy Governance, Institutional Trust, and Research Gap

Privacy-aware institutional AI requires both technical safeguards and governance. The NIST Privacy Framework treats privacy as an enterprise risk-management problem and emphasizes understanding data processing, governance responsibility, controls, and ongoing risk assessment. Its evolving AI-related materials specifically recognize that AI processing can create privacy risks even where traditional security controls are present.

The NIST AI Risk Management Framework similarly organizes AI governance around the functions Govern, Map, Measure, and Manage. This lifecycle perspective is highly relevant to institutional knowledge systems because model performance and privacy risk may change when documents are added, access permissions are modified, model versions are replaced, retrieval algorithms are updated, or user behavior changes. Continuous monitoring is therefore more appropriate than one-time predeployment approval.

Institutional knowledge systems introduce an additional governance challenge: authorization must follow the knowledge, not merely the application. A user may legitimately access the AI assistant but still lack authorization to view confidential human-resources records, legal documents, restricted research data, or executive strategy material. Retrieval systems therefore need source-level or document-level authorization filtering before information enters the model context.

This requirement is different from ordinary chatbot authentication. Logging into the system establishes user identity, but privacy-aware retrieval requires the system to evaluate whether the identified user may access each candidate knowledge item.

A second challenge concerns response provenance. Institutional users should be able to distinguish generated interpretation from authoritative source material. Retrieval-supported responses should therefore provide citations or source references where possible and should abstain or escalate when evidence is insufficient or contradictory.

A third challenge concerns the trade-off between SLM efficiency and reasoning capacity. Collaborative LLM–SLM research explicitly recognizes that local SLMs may provide privacy benefits while having weaker performance on tasks requiring broad reasoning. The research gap is therefore not whether institutions should universally replace LLMs with SLMs, but how they can design architectures that route tasks according to privacy sensitivity, complexity, and knowledge requirements.



The present study addresses this gap by treating the SLM as one layer within a governed institutional knowledge-support architecture and by evaluating privacy readiness together with reliability, access control, and operational practicality.

4. Research Methodology

4.1 Research Design and Synthetic Analytical Corpus

The study adopts a simulation-based methodological design because no verified institutional deployment logs or confidential enterprise datasets were provided. No actual employees, institutional documents, personal records, or proprietary databases were analyzed.

A synthetic corpus of 400 hypothetical knowledge-support interactions was constructed and divided equally across four system architectures:

- Cloud LLM with Raw Institutional Context;
- Hosted LLM with Access-Filtered RAG;
- Local SLM with Governed RAG; and
- Privacy-Aware SLM with Controlled Institutional Knowledge.

Each architecture was evaluated through six dimensions selected from current SLM, RAG-security, enterprise AI, and privacy-risk literature.

Table 1: Privacy-Aware Institutional Knowledge Support Framework

Evaluation Dimension	Core Question	Illustrative Control	Scale
Data Exposure Control	How effectively does the architecture minimize unnecessary transmission of sensitive institutional information?	Local inference, minimization, controlled connectors, restricted logging	0–100
Authorization Alignment	Does retrieved knowledge respect the user's institutional permissions?	Identity-aware retrieval, document-level access control, role mapping	0–100
Retrieval Grounding	Are answers supported by current and authoritative institutional sources?	Governed RAG, source ranking, freshness controls, provenance	0–100
Response Reliability	Does the system appropriately answer, abstain, qualify uncertainty, or escalate?	Evaluation thresholds, citation checking, human escalation	0–100
Operational Efficiency	Can the system operate with acceptable latency, cost, and infrastructure requirements?	Model right-sizing, caching, local inference, constrained tasks	0–100
Governance Observability	Can institutions audit model, retrieval, access, and output behavior?	Logs, model inventory, monitoring, incident management, evaluation records	0–100



Source: Author-developed simulation framework informed by NIST AI and Privacy Frameworks and contemporary research on SLMs and secure enterprise RAG.

4.2 Architecture Classification

The Cloud LLM with Raw Context model represents a deployment in which institutional information is inserted directly into prompts sent to a remotely hosted general-purpose language model with limited retrieval segmentation. Such an architecture may offer strong general language capability but creates a larger external data-processing surface.

The Hosted LLM with Access-Filtered RAG architecture introduces retrieval over institutional documents, removes irrelevant context, and filters candidate documents using organizational permissions before sending the selected context to a hosted model. This model reduces unnecessary exposure compared with raw-context prompting but still depends upon external model processing.

The Local SLM with Governed RAG architecture places the generator inside institution-controlled infrastructure and connects it to permission-aware retrieval. Institutional documents do not need to leave the local processing boundary for ordinary knowledge-support tasks. The smaller model is assumed to be optimized for bounded institutional question answering rather than unrestricted general reasoning.

The Privacy-Aware SLM with Controlled Institutional Knowledge architecture adds data classification, least-privilege retrieval, local inference, source-level permissions, secure embeddings, prompt-injection defenses, provenance, abstention thresholds, human escalation, continuous evaluation, and controlled routing to larger models only when a request can be safely externalized.

This classification reflects emerging collaborative-model research in which local small models and larger external models may be used selectively rather than as mutually exclusive technologies.

4.3 Analytical Procedure and Research Integrity

A Privacy-Aware Knowledge Support Index (PAKSI) was developed on a 0–100 scale by combining the six dimensions with equal conceptual weight. Equal weighting was selected for methodological transparency and should not be treated as an empirically validated scoring system.

No inferential statistical significance tests were conducted because simulated observations cannot establish population-level causal effects.

Future empirical validation should compare architectures using real institutional question-answer workloads, permission-sensitive retrieval tests, latency measurement, adversarial prompt testing, sensitive-information leakage assessment, citation accuracy, retrieval recall, abstention quality, user evaluation, energy consumption, and cost.

Particular attention should be given to false authorization. A system that returns the correct answer to the wrong user constitutes a privacy failure even if conventional accuracy metrics classify the output as successful.



5. Analysis

5.1 Comparative Performance of Institutional AI Architectures

The simulated analysis demonstrates that increasing model privacy requires more than replacing a large model with a smaller one.

Table 2: Simulated Privacy-Aware Institutional Knowledge Support Scores

Architecture	Data Exposure Control	Authorization Alignment	Retrieval Grounding	Response Reliability	Operational Efficiency	Governance Observability	PAKS I
Cloud LLM + Raw Institutional Context	28	36	45	68	48	27	42
Hosted LLM + Access-Filtered RAG	52	69	74	76	61	46	63
Local SLM + Governed RAG	88	87	85	73	83	76	82
Privacy-Aware SLM + Controlled Knowledge	95	94	92	87	89	95	92

Source: Simulated analytical values developed for methodological illustration. These are not measurements from an operational institution.

The cloud-based raw-context model achieves a comparatively stronger response-reliability score than its privacy score because a large general-purpose model may provide broad linguistic and reasoning capability. Its overall readiness remains low because sensitive context is unnecessarily exposed to a larger processing boundary and authorization may be difficult to enforce at source level.

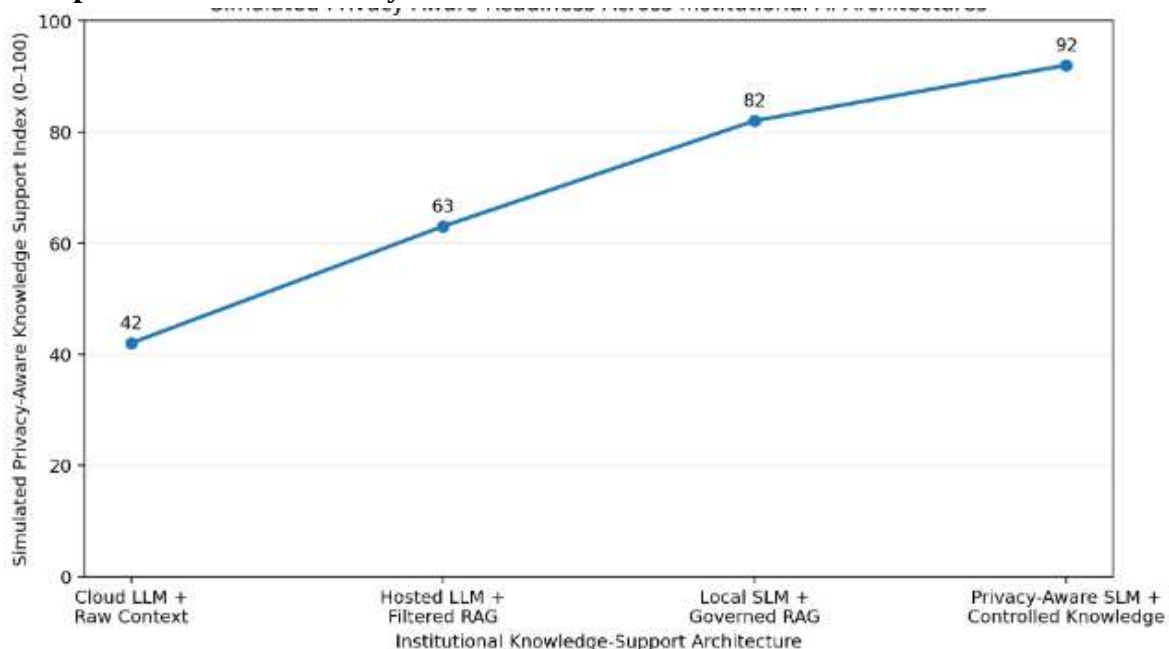
The hosted RAG model improves substantially by retrieving only relevant information and introducing access filtering. The architecture demonstrates that retrieval governance can materially improve privacy even when generation remains external.

The local SLM architecture receives substantially stronger exposure and authorization scores because model inference and institutional retrieval remain under the organization's infrastructure controls. Its simulated response-reliability score is slightly lower than that of the hosted larger model, reflecting the realistic possibility that smaller models may underperform on broad or highly complex reasoning tasks.

The final architecture achieves the strongest result because it does not depend on model size as its primary safeguard. Privacy comes from the combined effect of controlled inference, permission-aware retrieval, minimized context, source provenance, secure storage, task routing, abstention behavior, logging, and institutional governance.

5.2 Privacy-Aware Knowledge-Support Readiness

Graph 1: Simulated Privacy-Aware Readiness Across Institutional AI Architectures



The graph shows an illustrative increase in the Privacy-Aware Knowledge Support Index from 42 for cloud-based raw-context processing to 63 for hosted access-filtered RAG, 82 for locally deployed SLMs with governed retrieval, and 92 for the integrated privacy-aware architecture.

The 21-point increase between the first two models illustrates the contribution of controlled retrieval. The additional increase from 63 to 82 illustrates the potential value of local generation when institutional confidentiality is important.

The final increase to 92 reflects the more important conclusion that locality must be combined with governance. A local language model with unrestricted access to every internal document would still create serious privacy and security risk.

5.3 Task Routing and Institutional Knowledge Boundaries

The strongest architectural model does not require every request to be processed by the same model.



Routine policy lookup, procedure explanation, internal document summarization, standard-form assistance, and source-grounded question answering may be suitable for local SLM processing. Some complex tasks may require stronger reasoning than the local model can provide. In such cases, the institution can route only a sanitized or abstracted version of the problem to an external model, provided that the routing process prevents confidential content from leaving the institutional boundary.

Collaborative SLM–LLM research supports this broader architectural direction by examining how local specialized models and remote general models can complement one another. The institutional knowledge repository should remain authoritative regardless of which model handles the reasoning stage.

This means that source access, document versioning, permission checks, and freshness controls should be governed independently from the generative model. Such separation reduces the risk that institutional knowledge becomes inseparably embedded inside a model whose later behavior is difficult to audit or selectively update.

6. Discussion

6.1 Privacy Advantage Comes from Architecture, Not Model Size Alone

The central implication of the analysis is that institutions should avoid treating “small model” and “private model” as equivalent concepts. Smaller models create privacy opportunities because they are easier to deploy within controlled infrastructure and can reduce dependence on remote inference. However, privacy protection ultimately depends on the complete processing architecture. A small model that receives unrestricted confidential documents, stores every prompt indefinitely, or ignores institutional access permissions can create serious privacy failures despite operating entirely on local hardware.

This distinction is consistent with NIST's privacy-risk approach, which treats privacy as a consequence of data processing rather than merely a property of a particular technology. AI systems can generate privacy risk through inference, aggregation, reconstruction, or inappropriate data flows even where traditional identifiers have been removed.

For institutional deployment, the appropriate design principle is therefore data locality plus data minimization plus authorization. Local inference reduces unnecessary external processing. Data minimization reduces the quantity of sensitive information entering the model context. Authorization controls determine whether the current user is permitted to access the retrieved material at all. These three safeguards address different risks and cannot substitute for one another.

6.2 RAG Should Function as a Governance Boundary

The second major implication concerns retrieval-augmented generation.

RAG is frequently described as a technique for improving factual accuracy, but in institutional settings it should also be understood as a governance boundary. The retrieval layer decides which knowledge becomes visible to the model.

This makes retrieval one of the most important privacy and authorization enforcement points in the complete system. Secure enterprise-RAG research reinforces this conclusion by showing that proprietary knowledge requires filtering and controlled interaction with external model components.



Research on encrypted and privacy-aware retrieval further demonstrates that knowledge stores and embedding indexes themselves require protection.

Institutional RAG should therefore inherit permissions from underlying systems wherever possible. If a user cannot open a restricted document in the original repository, an AI assistant should not reveal its content simply because semantic retrieval ranks it highly. Authorization must be evaluated before protected content enters the model context. The same principle applies to multi-document reasoning. A generated answer may indirectly disclose information when it synthesizes several documents that individually appear harmless.

Institutions should therefore evaluate not only direct retrieval but also inference risk across combined sources. RAG can also strengthen accountability because responses can be linked to identifiable source documents. Source attribution allows users to verify whether a generated explanation reflects current policy, outdated documentation, or conflicting internal records.

This is particularly important in institutional environments where generated text may otherwise appear more authoritative than the evidence justifies. A privacy-aware assistant should be able to say that evidence is insufficient, that documents disagree, or that the user lacks access to the source required to answer a question. Such behavior may reduce apparent conversational fluency, but it increases institutional trustworthiness.

6.3 Human Oversight, Evaluation, and Deployment Governance

The third implication is that small-language-model deployment should be treated as an institutional capability requiring continuing governance rather than as a software installation.

The NIST AI RMF emphasizes continuous monitoring and risk management because AI-system behavior can change as models, components, data, and operating conditions change. This requirement is especially relevant for RAG systems because institutional knowledge bases change continuously. A model that performed well during initial testing may later receive new document types, new departments, revised access-control structures, or conflicting policy versions.

Consequently, institutions should maintain evaluation datasets representing genuine internal tasks and privacy boundaries. Testing should include routine questions, ambiguous requests, unauthorized information requests, prompt-injection attempts, outdated documents, conflicting sources, absent evidence, and questions requiring human escalation. Evaluation should not focus only on whether the generated answer matches an expected answer. It should also assess whether the correct source was retrieved, whether protected information was excluded, whether uncertainty was represented appropriately, whether the model fabricated unsupported institutional rules, and whether the system correctly refused requests outside the user's authorization.

Human oversight remains particularly important in high-consequence institutional contexts. An SLM can assist employees in locating procedures or summarizing internal material, but final responsibility for legal, financial, medical, disciplinary, admissions, or other consequential institutional decisions should remain with appropriately authorized personnel unless a separately validated governance framework justifies greater automation. The practical value of SLMs is therefore not that they eliminate human expertise. Their strongest role is to reduce the friction between institutional users



and governed organizational knowledge while preserving clear accountability for consequential interpretation and decision-making.

7. Conclusion

Small-language models provide a promising architectural foundation for privacy-aware institutional knowledge support because their reduced computational requirements can make local or institution-controlled deployment more practical. When combined with retrieval-augmented generation, they can provide employees with natural-language access to current organizational knowledge without requiring every policy, report, procedure, or confidential document to be incorporated into the model's trained parameters. The present study nevertheless demonstrates that small model size should not be treated as a privacy guarantee. Privacy arises from how information is classified, retrieved, transmitted, processed, logged, retained, and governed throughout the complete system.

The simulation-based comparison produced a Privacy-Aware Knowledge Support Index of 42 for cloud large-model processing with raw institutional context, 63 for hosted models using access-filtered RAG, 82 for locally deployed small-language models with governed retrieval, and 92 for an integrated privacy-aware SLM architecture. These values are methodological illustrations rather than empirical deployment results, but they clarify the structural contribution of each layer. Retrieval filtering reduces unnecessary exposure, local generation reduces dependence on external processing, and controlled institutional knowledge adds authorization, provenance, minimization, monitoring, and escalation. The strongest architecture therefore emerges not from replacing one model with another but from redesigning the complete knowledge-support pipeline around institutional privacy requirements.

The study ultimately concludes that small-language models are most valuable when institutions use them for bounded, source-grounded, privacy-sensitive knowledge tasks within a governed architecture. Future empirical research should evaluate such systems across healthcare organizations, universities, public agencies, research institutions, financial organizations, and other knowledge-intensive environments using real permission structures and adversarial privacy testing. Comparative studies should examine model quality, retrieval accuracy, privacy leakage, authorization failures, latency, cost, energy use, and user trust. The long-term objective should not be to maximize autonomous language-model capability inside institutions, but to create knowledge systems that make organizational information more accessible while preserving confidentiality, evidence provenance, accountability, and human control.

References

1. Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. *Proceedings of NAACL-HLT 2019*, 4171–4186. <https://doi.org/10.18653/v1/N19-1423>
2. Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., Levy, O., Lewis, M., Zettlemoyer, L., & Stoyanov, V. (2019). RoBERTa: A robustly optimized BERT pretraining approach. *arXiv preprint arXiv:1907.11692*.
3. Sanh, V., Debut, L., Chaumond, J., & Wolf, T. (2019). DistilBERT, a distilled version of BERT: Smaller, faster, cheaper and lighter. *arXiv preprint arXiv:1910.01108*.



4. Lan, Z., Chen, M., Goodman, S., Gimpel, K., Sharma, P., & Soricut, R. (2019). ALBERT: A lite BERT for self-supervised learning of language representations. *International Conference on Learning Representations (ICLR)*.
5. Howard, J., & Ruder, S. (2018). Universal language model fine-tuning for text classification. *Proceedings of ACL 2018*, 328–339. <https://doi.org/10.18653/v1/P18-1031>
6. Peters, M. E., Neumann, M., Iyyer, M., Gardner, M., Clark, C., Lee, K., & Zettlemoyer, L. (2018). Deep contextualized word representations. *Proceedings of NAACL-HLT 2018*, 2227–2237. <https://doi.org/10.18653/v1/N18-1202>
7. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.
8. Sutskever, I., Vinyals, O., & Le, Q. V. (2014). Sequence to sequence learning with neural networks. *Advances in Neural Information Processing Systems*, 27, 3104–3112.
9. McMahan, B., Moore, E., Ramage, D., Hampson, S., & y Arcas, B. A. (2017). Communication-efficient learning of deep networks from decentralized data. *Proceedings of AISTATS 2017*, 1273–1282.
10. Konečný, J., McMahan, H. B., Yu, F. X., Richtárik, P., Suresh, A. T., & Bacon, D. (2016). Federated learning: Strategies for improving communication efficiency. *arXiv preprint arXiv:1610.05492*.
11. 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>
12. Dwork, C. (2006). Differential privacy. In *Proceedings of the 33rd International Colloquium on Automata, Languages and Programming*, 1–12. Springer. https://doi.org/10.1007/11787006_1
13. 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>
14. Carlini, N., Liu, C., Erlingsson, Ú., Kos, J., & Song, D. (2019). The secret sharer: Evaluating and testing unintended memorization in neural networks. *28th USENIX Security Symposium*, 267–284.
15. Shokri, R., Stronati, M., Song, C., & Shmatikov, V. (2017). Membership inference attacks against machine learning models. *2017 IEEE Symposium on Security and Privacy*, 3–18. <https://doi.org/10.1109/SP.2017.41>
16. Kairouz, P., McMahan, H. B., Avent, B., Bellet, A., Bennis, M., Bhagoji, A. N., et al. (2019). Advances and open problems in federated learning. *arXiv preprint arXiv:1912.04977*.
17. Al-Rfou, R., Choe, D., Constant, N., Guo, M., Jones, L., Meier-Hellstern, K., & Shazeer, N. (2019). Character-level language modeling with deeper self-attention. *Proceedings of AAAI 2019*, 33, 3159–3166.
18. Ackerman, M. S., Pipek, V., & Wulf, V. (Eds.). (2003). *Sharing Expertise: Beyond Knowledge Management*. MIT Press.



19. Alavi, M., & Leidner, D. E. (2001). Review: Knowledge management and knowledge management systems: Conceptual foundations and research issues. *MIS Quarterly*, 25(1), 107–136. <https://doi.org/10.2307/3250961>
20. Nonaka, I., & Takeuchi, H. (1995). *The Knowledge-Creating Company: How Japanese Companies Create the Dynamics of Innovation*. Oxford University Press.