

Productive but Not Fully Trustworthy: The Auditability of Information Systems in the Age of Probabilistic Generative AI

Ade Ayuni Syam¹

¹ Fakultas Sains dan Teknologi, Universitas Terbuka, Indonesia

Correspondence: y.ayuni12@gmail.com

Article Info

Article history:

Received May 20, 2026

Revised May 25, 2026

Accepted May 30, 2026

Keyword:

AI governance; Auditability;
Generative AI; Information
Systems; Probabilistic Systems.

ABSTRACT

Generative AI is increasingly embedded in organisational information systems, yet its productivity value raises a difficult audit problem. Model-assisted tools can accelerate document work, coding support, service responses, and administrative review, but their probabilistic outputs do not follow the stable procedural logic assumed in conventional audit trails. This study examines how Generative AI affects auditability once generated responses become part of organisational workflows. A directed literature review and qualitative content analysis were conducted on academic studies, technical guidance, and AI governance standards related to Generative AI, information systems, risk management, LLM security, and auditability. The analysis indicates that audit weakness appears less as a single technical defect than as a reconstruction problem across prompt, data, model, output, decision, and oversight conditions. Based on this finding, the study develops the Generative AI Auditability Stack, a conceptual framework that treats auditability as a distributed property of GenAI-based information systems. The framework clarifies why transparency and explainability alone are insufficient for model-assisted work, especially where proprietary systems limit internal visibility. Organisations can still strengthen reviewability by preserving prompt records, data lineage, model versioning, validation traces, and decision evidence across the lifecycle of use. The study contributes to information systems research by shifting the debate from adoption and productivity toward the auditability of probabilistic systems in organisational settings.



© 2026 The Authors. Published by Punggawa Legacy Center. This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by/4.0/>)

INTRODUCTION

Generative AI has moved from experimental chat interfaces into the ordinary machinery of organisational information systems (Feuerriegel et al., 2024). Customer-service platforms use language models to draft replies before agents approve them. Human-resource units rely on generative tools to summarise applicant documents. Software teams ask code assistants to repair functions that used to be reviewed line by line by senior developers. These uses look practical rather than spectacular, and that practical quality explains their fast adoption (Kanbach et al., 2024). Productivity increases because the system can prepare text, condense information, and suggest decisions before a user has finished collecting the relevant material (Fügener et al., 2022). Yet the same speed makes the system harder to examine. A response may appear coherent while the factual basis remains unclear, and a process that saves time can quietly weaken the chain of evidence behind organisational work.

The central difficulty lies in the kind of output produced by Generative AI. Conventional information systems usually rely on predefined rules, database transactions, or procedural workflows that can be traced through logs and business rules. A payroll system may still fail, but auditors can usually identify the relevant input, formula, database field, access record, or approval step. Generative AI changes that audit condition because the output emerges from probabilistic pattern generation rather than from a stable rule that can be read back in ordinary procedural terms. A similar prompt can produce a different answer, a missing context can alter a recommendation, and a fluent sentence can hide a false claim. Organisational users therefore face a system that behaves as if it knows, while its internal basis is difficult to verify.

Auditability becomes more than a compliance concern under these conditions. For information systems, auditability refers to the capacity to trace, review, verify, and assign responsibility for system behaviour after a process has taken place. The concept is familiar in transaction processing, enterprise systems, database administration, and cybersecurity governance, where logs, access controls, version records, and approval trails help reconstruct events. Generative AI unsettles these familiar mechanisms because the key event is no longer limited to data entry or data retrieval. The event also includes a prompt, a model version, a retrieved context, a generated output, and a human decision about whether to use that output. Each layer adds a possible source of error that may not appear in standard information-system audit records.

Recent discussions on trustworthy AI, AI risk management, and LLM security have begun to address parts of this problem. Risk frameworks draw attention to governance, validation, documentation, and human oversight, while security guidance warns that model-based applications can be manipulated through prompts or exposed through careless integration with organisational data. These contributions are important, but a gap remains for information systems research. Much of the debate treats Generative AI either as a tool for productivity or as a technical object located inside machine-learning research. The organisational system in which the model is embedded receives less careful treatment (Berente et al., 2021). A language model used inside a procurement platform, a university helpdesk, or a hospital administration system does not create risk in isolation; risk emerges through its connection with users, records, workflows, interfaces, and decisions.

A sharper account of auditability is needed because trust in Generative AI is often produced at the surface of interaction. Users may accept a response because the wording sounds complete, the interface feels familiar, or the system answers faster than a colleague. Such trust differs from auditability (Mökander et al., 2024). A system can be trusted by users and still fail an audit because the source of an output is unclear, the model version was not recorded, the prompt was not preserved, or the final decision cannot be separated from the generated recommendation. This distinction matters for organisations that adopt Generative AI in routine information work. Productivity may be visible immediately, while the cost of weak auditability appears later through disputed decisions, security incidents, poor data governance, or accountability gaps that only become visible after harm has occurred.

The study develops an information-systems account of Generative AI auditability by linking probabilistic output, organisational risk, and layered control. The guiding questions are how the probabilistic character of Generative AI changes audit conditions in information systems, what forms of risk arise from its integration into organisational workflows, and which auditability components are needed to make such systems traceable, verifiable, and accountable. The argument does not treat Generative AI as an inherently unreliable technology. A more precise claim is that productivity gains become fragile once prompt histories, data lineage, model versions, output validation, and human review are absent or poorly connected. Auditability should therefore be treated as a design property of Generative AI-based information systems, rather than a corrective measure added after deployment.

RESEARCH METHODS

A qualitative design was selected because the study deals with a conceptual and organisational problem rather than the technical performance of a particular model. The analysis focuses on how Generative AI affects auditability once language models are embedded in information systems used for routine organisational work. A directed literature review was combined with qualitative content analysis to keep the method simple, transparent, and still suitable for information systems research. The review did not aim to count the entire publication landscape. The narrower aim was to collect relevant academic and technical sources that explain the relation between probabilistic output, system risk, governance control, and audit practice. Such a design allows a focused reading of the problem without turning the article into a bibliometric mapping exercise.

The data consisted entirely of secondary sources. Academic sources were drawn from publications on Generative AI, information systems, AI governance, auditability, trustworthy AI, and LLM security (Van Slyke et al., 2023). Technical and institutional documents were also included because auditability in Generative AI-based systems is shaped by standards, risk frameworks, and security guidance rather than academic literature alone. NIST AI RMF, the NIST Generative AI Profile, OWASP guidance for LLM applications, and ISO/IEC 42001 were treated as reference documents for

comparing how risk, documentation, oversight, and accountability are translated into operational controls. The inclusion of these documents helped connect conceptual debates in information systems with practical requirements faced by organisations that deploy model-based applications.

The analytical process used qualitative coding. Each selected source was read for claims about the behaviour of Generative AI, the risks produced by probabilistic output, and the control mechanisms proposed to make such systems reviewable. Initial codes were developed from the research questions, especially traceability, explainability, verifiability, accountability, human oversight, data lineage, prompt logging, model versioning, and output validation. Coding remained flexible because several sources described similar problems with different terminology. For example, a security document may discuss prompt injection through the language of attack surfaces, while an information systems article may approach the same issue through system control and organisational responsibility. These differences were preserved rather than forced into a rigid coding scheme.

Analysis proceeded by comparing recurring ideas across academic and technical sources. The comparison paid attention to where the literature converged, where the guidance remained too broad, and where information systems research needed a more specific audit lens. The final stage organised the findings into the Generative AI Auditability Stack, a conceptual model that links prompt, data, model, output, decision, and oversight layers. The model was not designed as a measurement instrument. Its role is analytical. The stack clarifies which parts of a Generative AI-based information system must remain traceable if organisations expect the system to be reviewed after use, especially in workflows where generated outputs become part of administrative records or managerial decisions.

RESULTS AND DISCUSSION

Results

1. Probabilistic Output as an Audit Problem

Generative AI changes the audit object in information systems because the main event is no longer limited to a recorded transaction or a visible data modification. Across the coded materials, audit difficulty appeared most often as a problem of reconstruction. A generated response may come from an ordinary-language prompt, a model configuration set outside the organisation, retrieved internal documents, earlier conversational context, and filtering mechanisms that remain partly inaccessible to users. Conventional audit trails can still record access time, user identity, and system activity, yet such records do not explain why a particular answer was produced. A helpdesk system may show that an employee asked about refund approval, while the deeper audit question concerns the wording of the prompt, the policy text retrieved by the system, and the extent to which the generated answer followed or distorted that policy.

Probabilistic output weakens the usual assumption that repeated system use will produce stable and reviewable results. A deterministic system may return the same tax calculation or inventory status as long as the input and rule set remain unchanged. Generative AI does not provide the same audit comfort. A small variation in wording, a missing organisational context, or an updated model version can shift the answer without producing an obvious system error. The coded academic sources tended to describe this problem through reliability, hallucination, uncertainty, and user trust (Huang et al., 2025). Technical guidance approached the same difficulty through documentation, testing, monitoring, and human oversight (Das et al., 2025). The two streams meet at a practical point. Generated output cannot be treated as a self-contained artefact because its reliability depends on conditions that precede the visible response.

The analysis therefore suggests that auditability should move beyond output checking. Treating a generated answer as a final object to be accepted or rejected misses the chain of conditions that shaped the answer. Hallucination offers a useful example, although the broader issue is not limited to false content. A wrong claim may originate from weak retrieval, ambiguous prompting, incomplete organisational data, model behaviour, or absent human validation before the text enters a report or decision note. An audit process that looks only at the final answer would identify the error but fail to locate the system condition that made the error possible. A more useful audit lens follows the output backward through the prompt, the data context, the model environment, the validation process, and the later organisational use.

Generative AI-based information systems can fail while appearing technically functional. The model responds, the interface loads, and the user receives an answer, yet the organisation may have no

adequate record of the conditions that produced the response. Failure becomes less visible than in conventional systems because no crash, rejection, or transaction mismatch may occur. A plausible response can travel into an email, a service ticket, a policy summary, or a managerial note before anyone asks whether the underlying claim was grounded. The reviewed technical documents sharpen this issue by treating documentation and monitoring as control requirements rather than administrative accessories. Academic discussions add a further caution: user trust can grow faster than the organisation's capacity to verify what the system has done.

This pattern explains why auditability in Generative AI-based information systems needs a layered structure. The problem is not located only inside the model, even though model behaviour remains important. Audit weakness may begin with an unrecorded prompt, continue through unclear data retrieval, become harder to detect through fluent output, and finally become consequential through a human decision that treats the generated text as reliable. The emerging analytical category is therefore layered auditability. Prompt records, data lineage, model versioning, output validation, and human review are connected parts of one audit chain rather than separate technical controls (Holldack et al., 2026). The next analytical step concerns how specific risks attach themselves to that chain and why the proposed Generative AI Auditability Stack must follow the architecture of organisational use.

2. From Isolated AI Risks to Information-System Audit Chains

The coding process indicates that the most recurring analytical pattern was not the presence of separate AI risks, but the conversion of those risks into disruptions of the audit chain. The audit chain refers to the connected record through which an organisation can relate a user instruction, data context, model behaviour, generated output, and subsequent decision. Security documents often name risks through technical labels, while academic sources describe their effects through trust, reliability, and governance problems. A direct listing of threats would miss the deeper issue. In organisational use, risk emerges through the route taken by an output before it becomes part of work. A generated answer may begin with a weak prompt, draw from an outdated document, pass through an unrecorded model version, and later appear in a service note that carries institutional authority.

Prompt injection illustrates how a technical risk becomes an audit-chain disruption. Treated narrowly, prompt injection appears as an attack against a language model. Treated through an information-systems lens, the risk becomes a failure of traceability across interface, document retrieval, and output handling. A hidden instruction inside a retrieved file may redirect the model's response while the user believes the system is following an approved organisational source. Later review may reveal that the answer was wrong, but the organisation still needs to know which document carried the instruction, whether the retrieval mechanism exposed the model to that document, and whether the generated response entered a workflow without validation. The incident is less about a clever prompt than weak reconstruction capacity inside the system.

Hallucination creates a different form of audit-chain disruption because false content can move from an interaction window into organisational records. Academic literature commonly treats hallucination as a reliability problem, but information systems research has to ask where the false claim travels after generation. A model-generated policy summary may be copied into an email to a student or inserted into a procurement note before anyone checks the source. The practical issue is record contamination rather than textual error alone. Once unsupported content enters a ticket, memorandum, or decision file, later users may treat it as part of the organisation's knowledge base. Technical guidance on validation and monitoring becomes relevant here because generated statements need to be checked before they acquire administrative weight.

Data-related risks add another layer because Generative AI systems often blur the boundary between retrieval, generation, and disclosure (Chen et al., 2025). A user may paste confidential material into a public model, while an internal assistant may retrieve documents beyond the user's proper access level. The immediate concern is privacy or security, yet the audit problem runs deeper. An organisation must be able to reconstruct what data was exposed, which permission rule allowed access, how the model incorporated the material, and whether sensitive information appeared in the output. Without data lineage and access records that connect to model use, later investigation may identify leakage without knowing whether the failure came from user behaviour, system design, retrieval configuration, or governance policy (Damaris et al., 2025).

Accountability gaps appear at the end of the chain, where generated output becomes part of organisational judgment (Bartsch et al., 2025). A decision maker may describe the model as merely advisory, while the written record may rely heavily on generated reasoning. Vendors may control model updates, administrators may configure access, and users may decide whether to accept a response. Responsibility becomes distributed in a way that conventional audit trails rarely capture. The convergence between academic and technical sources reaches beyond shared topics. Both streams point toward the same audit requirement. Technical evidence must remain connected to organisational responsibility. Prompt records, retrieval logs, model documentation, validation notes, and approval traces need to be linked because responsibility cannot be assigned from the final output alone.

3. The Emergence of Layered Auditability

The coding did not produce auditability as one flat theme. The recurring pattern was a separation of audit evidence into control points located before, during, and after generation. This pattern explains why the Generative AI Auditability Stack is better treated as an analytical device than as a neat description of organisational reality. In practice, prompt, data, model, output, decision, and oversight layers may overlap, especially in retrieval-augmented systems where a user instruction and a retrieved document shape the response at the same time (Gao et al., 2025). The stack remains useful because it helps locate where evidence disappears. A conventional audit record may capture login activity, but model-assisted work requires a wider trail that connects user instruction, data context, generated output, and later organisational use.

Prompt and data layers sit close to the origin of system behaviour. A prompt record matters because ordinary language can carry hidden assumptions, incomplete instructions, or sensitive material. Data lineage matters because a generated response may rely on documents whose authority, age, and access status vary across repositories. Weakness at these early layers produces uncertainty before any answer appears. A university helpdesk assistant that draws from an outdated tuition policy may generate a confident reply, while the deeper failure lies in retrieval governance rather than language generation alone. The model and output layers then add another form of uncertainty because vendor updates, configuration changes, and validation gaps can alter what the organisation later needs to reconstruct (Kadir, 2026a, 2026c).

Decision and oversight layers complete the stack because organisational risk rarely ends at the generated response. A draft becomes consequential once a staff member accepts it, edits it, or copies it into a record. Human review should leave traces that clarify whether the output was checked, modified, rejected, or escalated. Oversight also connects technical evidence to policy, role assignment, and incident handling. The reviewed standards and security documents describe these controls in broad language, but the coding suggests a sharper requirement for information systems. Documentation is useful only when it records the disputed part of the workflow. Monitoring matters only when it can connect a generated output with the condition that made later organisational reliance possible.

Discussion

1. Auditability Beyond Transparency

Debates on AI governance often give transparency a central role, but Generative AI makes transparency a difficult operational remedy. Large language models cannot be inspected easily at the level of internal computation, and many organisations rely on proprietary systems whose design choices remain outside local control. A demand for full transparency may be normatively attractive while offering limited guidance for a helpdesk, procurement office, or academic administration unit already using model-assisted tools. Auditability offers a more practical lens because it asks what evidence must be preserved for later review, even where complete internal visibility is unavailable. The concern shifts from seeing the entire model to reconstructing the conditions under which a specific output entered organisational work.

The distinction also refines the relation between auditability and explainability (Meske et al., 2022). Explainability makes system behaviour intelligible to users, developers, or reviewers. Auditability asks whether that behaviour can be examined, challenged, and connected to responsibility. A short explanation attached to a generated answer may reassure a user, but an auditor still needs prompt history, retrieved sources, model version, access records, and validation traces. Conversely, an organisation may lack parameter-level explanations and still maintain a defensible audit posture if

surrounding evidence is preserved. Fluent explanations can even deepen the problem if they encourage acceptance while leaving source quality, retrieval logic, and decision use poorly documented. The issue concerns the distance between apparent intelligibility and institutional reviewability.

A stronger information systems perspective treats auditability as a socio-technical relation. Logs, version records, and validation checkpoints matter, but their value depends on whether people use them during review, correction, and accountability processes (Laux, 2024). A system may generate detailed records that no team examines, while a governance policy may require review without leaving any system trace that review occurred. Transparency and explainability remain useful, yet neither can replace the practical need to connect technical artefacts with organisational routines. Some organisations respond to opacity through disclaimers or interface warnings. Those measures may reduce confusion, but they do not explain why a response was generated or how it entered a decision file.

2. Auditability-by-Design in Information Systems

Generative AI makes auditability-by-design a necessary principle for information systems development, although implementation will rarely be clean. Retrofitted controls usually arrive after users have already built habits around the system. Staff may paste model responses into email templates, treat summaries as reliable starting points, or use generated code inside internal tools before documentation practices mature. Auditability-by-design addresses this problem by embedding review capacity into the architecture of use. Prompt retention, retrieval records, model versioning, validation workflow, and decision traces should become ordinary parts of system behaviour, but their design has to follow workflow risk rather than a universal compliance template.

The principle becomes harder when organisations depend on third-party models, cloud platforms, or embedded AI features inside software they do not fully control. Vendor dependency may limit access to model behaviour, logging granularity, retention policy, and configuration history. Smaller organisations may have even less leverage to request audit functions from providers. These constraints do not weaken the need for auditability; they clarify where information systems governance must intervene (Kadir, 2026b). Procurement requirements, API logging, access policy, data-retention rules, and service-level agreements become part of the audit architecture. A system analyst cannot treat GenAI integration as a simple interface problem because later review depends on evidence that may sit outside the organisation's direct technical control.

Auditability-by-design should also vary according to how generated output enters organisational authority. A private drafting assistant carries a different audit burden from a model connected to student services, finance, procurement, or healthcare administration. The relevant distinction is not the label of the tool, but the route through which generated text becomes a record, recommendation, or decision. Design choices decide whether reviewers can later identify the prompt, check the retrieved source, compare the model version, and see whether a human accepted or changed the answer. The stack does not solve every problem created by probabilistic models. Its value lies in helping organisations locate where audit evidence is lost before productivity becomes difficult to review.

CONCLUSION

The directed reading of academic studies, technical guidance, and AI governance standards indicates that Generative AI creates a distinctive audit problem for information systems. Model-assisted work can accelerate organisational tasks that involve text, code, documents, and service communication, yet the speed of generation often exceeds the organisation's ability to reconstruct the conditions behind the output. The deeper risk is not limited to inaccurate content. A generated response becomes difficult to govern once its prompt history, data context, model condition, validation record, and subsequent human use are scattered across systems or left undocumented. Under these conditions, productivity becomes fragile because the organisation gains faster work without an equally strong capacity to verify how that work was produced and why it was accepted.

The Generative AI Auditability Stack reframes auditability as a distributed property of GenAI-based information systems. Audit evidence must remain connected across the interaction that shapes the response, the data environment that supplies its context, the model condition under which it is produced, and the organisational decision that gives the response practical force. This framing avoids reducing auditability to transparency or explainability alone. Proprietary models may never provide full internal visibility, but organisations can still preserve the surrounding evidence needed for review,

challenge, correction, and accountability. Further research should examine the stack in real deployments, particularly in systems where generated text enters service records, academic administration, procurement files, software development workflows, or public-sector decisions before its evidentiary basis has been fully tested.

REFERENCES

- Bartsch, S. C., Nguyen, L. H., Schmidt, J.-H., Du, G., Adam, M., Benlian, A., & Sunyaev, A. (2025). The Present and Future of Accountability for AI Systems: A Bibliometric Analysis. *Information Systems Frontiers*, 27(6), 2463–2484. <https://doi.org/10.1007/s10796-025-10636-9>
- Berente, N., Gu, B., Recker, J., & Santhanam, R. (2021). Managing Artificial Intelligence. *MIS Quarterly*, 45(3), 1433–1450. <https://doi.org/10.25300/MISQ/2021/16274>
- Chen, K., Zhou, X., Lin, Y., Feng, S., Shen, L., & Wu, P. (2025). A survey on privacy risks and protection in large language models. *Journal of King Saud University Computer and Information Sciences*, 37(7), 163. <https://doi.org/10.1007/s44443-025-00177-1>
- Damaris, R., Rosadi, S. D., & Bratadana, I. M. D. (2025). DATA GOVERNANCE FOR ARTIFICIAL INTELLIGENCE IMPLEMENTATION IN THE FINANCIAL SECTOR: AN INDONESIAN PERSPECTIVE. *Journal of Central Banking Law and Institutions*, 4(3), 445–472. <https://doi.org/10.21098/jcli.v4i3.430>
- Das, B. C., Amini, M. H., & Wu, Y. (2025). Security and Privacy Challenges of Large Language Models: A Survey. *ACM Computing Surveys*, 57(6), 1–39. <https://doi.org/10.1145/3712001>
- Feuerriegel, S., Hartmann, J., Janiesch, C., & Zschech, P. (2024). Generative AI. *Business & Information Systems Engineering*, 66(1), 111–126. <https://doi.org/10.1007/s12599-023-00834-7>
- Fügener, A., Grahl, J., Gupta, A., & Ketter, W. (2022). Cognitive Challenges in Human–Artificial Intelligence Collaboration: Investigating the Path Toward Productive Delegation. *Information Systems Research*, 33(2), 678–696. <https://doi.org/10.1287/isre.2021.1079>
- Gao, Z., Jian, Z., & Mousavirad, S. J. (2025). Reinforcement learning-driven feature selection enhanced by an evolutionary approach tuning for criminal suspect identification. *Scientific Reports*, 15(1), 41879. <https://doi.org/10.1038/s41598-025-25920-6>
- Holldack, F., Banh, L., & Strobel, G. (2026). Agentic information systems. *Electronic Markets*, 36(1), 5. <https://doi.org/10.1007/s12525-025-00861-0>
- Huang, L., Yu, W., Ma, W., Zhong, W., Feng, Z., Wang, H., Chen, Q., Peng, W., Feng, X., Qin, B., & Liu, T. (2025). A Survey on Hallucination in Large Language Models: Principles, Taxonomy, Challenges, and Open Questions. *ACM Transactions on Information Systems*, 43(2), 1–55. <https://doi.org/10.1145/3703155>
- Kadir, Z. K. (2026a). Beyond Firearms: Mapping Global Patterns of Non-Firearm Homicide in Comparative Criminological Perspective. *Punggawa Global Research: Jurnal Multidisiplin*, 1(2), 1–10.
- Kadir, Z. K. (2026b). Narrative Closure in Honor killing Cases: How Judgments Stabilise Meaning, Eliminate Ambiguity, and Produce Sentencing Certainty. *Punggawa Law Review*, 1(1), 1–10.
- Kadir, Z. K. (2026c). Neurocriminology and the Next Generation of Criminological Theory: Integration, Limits, and Ethical Risks. *Punggawa Global Research: Jurnal Multidisiplin*, 1(1), 1–8.
- Kanbach, D. K., Heiduk, L., Blueher, G., Schreiter, M., & Lahmann, A. (2024). The GenAI is out of the bottle: generative artificial intelligence from a business model innovation perspective. *Review of Managerial Science*, 18(4), 1189–1220. <https://doi.org/10.1007/s11846-023-00696-z>
- Laux, J. (2024). Institutionalised distrust and human oversight of artificial intelligence: towards a democratic design of AI governance under the European Union AI Act. *AI & SOCIETY*, 39(6), 2853–2866. <https://doi.org/10.1007/s00146-023-01777-z>
- Meske, C., Bunde, E., Schneider, J., & Gersch, M. (2022). Explainable Artificial Intelligence: Objectives, Stakeholders, and Future Research Opportunities. *Information Systems Management*, 39(1), 53–63. <https://doi.org/10.1080/10580530.2020.1849465>
- Mökander, J., Schuett, J., Kirk, H. R., & Floridi, L. (2024). Auditing large language models: a three-layered approach. *AI and Ethics*, 4(4), 1085–1115. <https://doi.org/10.1007/s43681-023-00289-2>

Van Slyke, C., Johnson, R., & Sarabadani, J. (2023). Generative Artificial Intelligence in Information Systems Education: Challenges, Consequences, and Responses. *Communications of the Association for Information Systems*, 53(1), 1–21. <https://doi.org/10.17705/1CAIS.05301>