

A Global Comparison of Schemas, Transparency, and Interoperability in Public-Sector AI Registers and Inventories

Dipto Das
Cornell University
Ithaca, New York, United States
dd749@cornell.edu

Shion Guha
University of Toronto
Toronto, Ontario, Canada
shion.guha@utoronto.ca

Abstract

Artificial intelligence (AI) registers and inventories are intended to make governmental AI visible, but their institutional scope, schemas, and reporting practices construct different representations of public-sector AI. We compare 8,368 records from country-specific and transnational inventories covering 72 countries. Across 23 harmonized fields, registers shared a descriptive core but rarely requested information about appeals, risks, legal bases, or external evaluation. We found that broad schemas often contained substantial missingness, schema similarity showed no significant patterned convergence, and multiple sources covering the same jurisdictions overlapped only selectively. Based on these findings, we synthesize a *layered visibility framework* that shows how register records reflect disclosure arrangements and why interoperability requires shared concepts, clear definitions, and preserved provenance.

ACM Reference Format:

Dipto Das and Shion Guha. 2026. A Global Comparison of Schemas, Transparency, and Interoperability in Public-Sector AI Registers and Inventories. In *Proceedings of ACM Conference on Digital Transformation (DXConf '26)*. ACM, New York, NY, USA, 8 pages. <https://doi.org/XXXXXXX.XXXXXXX>

1 Introduction

While AI is used in at least one area of government in 35 of 36 countries surveyed by the Organization for Economic Co-operation and Development, only three reported mandatory AI-use-case repositories, while another ten maintained repositories based on optional agency contributions [50]. This disparity is increasingly consequential as governments use AI to transform public services, internal operations, policymaking, and administrative decision-making, while its capabilities advance faster than the evidence and institutional safeguards needed to govern them [28]. Thus, governments face a fundamental digital-transformation problem: they are deploying AI faster than they are building shared infrastructure to identify, document, and scrutinize it. Public AI inventories and registers have emerged in response, but their institutional mandates, schemas, reporting units, and collection processes differ substantially. We refer to these collectively as AI-inventorying instruments. This

paper compares how these instruments are designed and what representations of governmental AI they produce.

AI registers extend structured disclosure from individual systems to portfolios used across public institutions. Although initiatives have emerged at municipal [25, 47], national [11], and transnational [52] levels, prior research has generally examined them individually. Consequently, we know comparatively little about how their schemas and reporting practices differ or whether sources covering the same jurisdiction enumerate the same systems. Rather than treating registers as interchangeable databases or censuses of governmental AI, we compare them as policy instruments whose reporting authorities, inclusion rules, schemas, classifications, and collection procedures shape what becomes publicly visible and governable [39]. Because policy models can circulate through learning, emulation, and transnational networks while being adapted to local conditions [16, 42, 59], we also examine whether schema similarities exhibit temporal or geographic clustering. We ask:

- **RQ1.** How do AI-inventorying instruments differ in schema breadth and field completion, and do their schemas exhibit temporal convergence or geographic clustering?
- **RQ2.** How do the portfolios represented by different AI-inventorying instruments vary in lifecycle composition and organizational concentration?
- **RQ3.** To what extent do country-specific and transnational instruments covering the same jurisdiction report the same AI systems?

We conducted a comparative quantitative study of 8,368 records from 17 country-specific and transnational instruments covering 72 countries, along with the European Union (EU)-level and global records. The instruments shared a descriptive core but showed no evidence of temporal convergence or clear geographic clustering. Broader schemas were not necessarily more complete, represented portfolios differed in lifecycle composition and organizational concentration, and sources covering the same jurisdictions overlapped only selectively. These findings show that AI registers do not directly mirror governmental AI activity but construct visibility through institutional scope, schema design, and reporting realization. We synthesize these relationships in a *layered visibility framework* explaining why similarly labeled instruments can produce different yet complementary representations of the same governmental AI landscape.

2 Literature Review

We situate AI registers within research on digital transformation and public-sector governance, documentation infrastructures, and policy diffusion, convergence, and interoperability. Together, these

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

DXConf '26, Ann Arbor, MI

© 2026 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 978-1-4503-XXXX-X/2018/06

<https://doi.org/XXXXXXX.XXXXXXX>

perspectives explain how institutional conditions, schemas, and reporting practices shape the public visibility of governmental AI.

2.1 AI Governance in Digital Transformation

Digital transformation involves more than adopting new technologies or converting existing information into digital formats. It entails broader changes to organizational structures, institutional relationships, and processes of value creation [63]. Within government, it similarly involves redesigning administrative processes and public services to pursue objectives such as efficiency, interoperability, transparency, and citizen responsiveness [44]. Public-sector AI adoption is one component of this transformation: as agencies introduce computational systems to automate tasks, it also reorganizes how information is collected, decisions are supported, services are delivered, and responsibilities are distributed among officials, technical systems, and external providers [57].

Public-sector AI governance is commonly articulated through principles such as transparency, accountability, fairness, safety, privacy, and trustworthiness [14, 21, 30]. However, these principles must be operationalized by bureaucratic organizations with different missions, legal responsibilities, resources, procurement arrangements, and technical capacities [22, 57, 61]. Algorithmic accountability is therefore not solely a property of system design but an institutional practice shaped by governing rules, professional judgment, organizational incentives, and relationships among agencies, technology providers, workers, and affected publics [3, 31, 46, 67]. It depends partly on whether institutions identify the systems they use, document their roles in administrative processes, trace the actors involved, and disclose this information to relevant publics [37].

Even when information is public, transparency does not necessarily produce accountability or make a system understandable, actionable, or contestable [1, 35]. Transparency is better understood as a situated communicative arrangement in which the content, format, audience, and institutional conditions of disclosure shape what recipients can know and do [20, 49]. From this perspective, public AI inventories are not merely websites or datasets but administrative mechanisms that determine which systems are reported, how they are classified, who maintains their records, and which aspects of public-sector AI become externally visible.

2.2 AI Registers as Disclosure Infrastructures

Documentation frameworks translate abstract commitments to transparency and accountability into standardized, inspectable artifacts. Datasheets [23], model cards [45], audit reports [55], impact assessments [65], data statements [43], and other reporting standards structure disclosures about provenance, intended use, limitations, risks, performance, and responsible actors. These artifacts can support coordination among developers, deploying organizations, regulators, auditors, and affected communities [10, 46, 62]. At the same time, documentation practices reflect institutional norms and assumptions about what information is relevant, which audiences matter, and what forms of responsibility should be recorded [4, 13, 54]. Therefore, documentation instruments such as AI inventories shape which system properties become legible, which actors are rendered responsible, and which accountability claims can subsequently be made [38, 53].

Critical data studies scholarship has shown how claims of scale, neutrality, completeness, and representativeness can obscure the social and organizational conditions under which datasets and systems are produced [12, 34]. Standardized documentation can similarly reduce complex concerns about fairness and accountability to questions of data sufficiency or technical performance while making data labor, contextual knowledge, and organizational practice less visible [2, 24, 66]. Thus, documentation can enable scrutiny while simultaneously delimiting its scope.

AI registers and inventories extend documentation from individual datasets or models to portfolios of algorithmic systems used across organizations. Municipal governments such as Amsterdam and Helsinki were among the earliest adopters of public algorithm registers, publishing structured information about systems used in public services [25, 47, 58]. Similar initiatives subsequently appeared in Nantes, Antibes, Lyon, and other municipalities, as well as at national and transnational levels across North America and Europe [32, 51]. Registers also function as institutional ontological design [11]: their categories define which systems and properties become administratively legible. Producing them depends on often-invisible work, including identifying relevant systems, interpreting inclusion rules, coordinating reporting across organizations, translating local practices into standardized fields, validating submissions, and updating records as systems move through their lifecycles [6, 54].

Hence, it is important to distinguish between the information an instrument is structurally capable of representing and the extent to which organizations populate its available fields. For example, a register that requests information about procurement, human involvement, or appeal mechanisms creates different possibilities for accountability than one limited to system names and short descriptions. Disclosure templates, reporting thresholds, machine-readable categories, and organizational incentives and interpretations shape both what institutions report and what remains difficult to disclose [26, 32, 48]. Transparency databases can consequently standardize reporting while still providing limited contextual explanation or falling short of their regulatory purposes [17, 24, 54, 60].

Because instruments differ in institutional coverage, reporting obligations, and treatment of planned, operational, and retired systems, the portfolios they disclose may also differ in lifecycle composition and organizational concentration. These distributions should be interpreted as properties of particular disclosure arrangements rather than as direct measures of governmental AI adoption. A public search interface or database-like inventory can be understood as the visible outcome of broader coordination among policy requirements, administrative labor, data schemas, and public communication [48, 68].

2.3 Policy Diffusion and Interoperability

AI registers can also be understood as policy instruments. Policy scholarship defines instruments as technical and social devices that organize relationships between governing institutions and those they seek to govern [39]. Because of their sociopolitical embeddedness, public AI registers and inventories should not be treated as neutral. Their reporting authorities, inclusion criteria, schemas, submission and update procedures, and public interfaces privilege

particular forms of knowledge, assign responsibilities, and structure administrative behavior. As a result, instruments that share the general objective of making governmental AI visible may embody substantially different models of accountability.

The growing adoption of AI registers worldwide also raises questions of diffusion and transfer: how policy spreads through learning, competition, coercion, and the construction of shared norms by professional and transnational communities [16]. International organizations, professional networks, civil-society initiatives, consultants, and other non-state actors can serve as transfer agents by circulating policy knowledge and facilitating the uptake of ideas, administrative arrangements, formal instruments, and informal norms of appropriate governance across various settings [42, 59].

However, diffusion does not necessarily produce uniformity. Jurisdictions may copy an existing template, emulate selected features, combine elements from several models, or reinterpret a shared idea in light of local administrative arrangements [29]. For example, similar labels such as “AI register,” “algorithm dashboard,” or “use-case inventory” can consequently refer to instruments with different reporting units, legal statuses, collection procedures, or public functions [52]. This makes schema interoperability—the capacity to represent and interpret corresponding concepts across heterogeneous instruments—a central concern. Instruments developed independently may nevertheless converge around a descriptive core as they confront similar administrative problems. Temporal proximity and participation in the same institutional networks may facilitate such convergence [18, 41], though particular diffusion mechanisms can vary. Observed similarity may also arise from common technical requirements, widely circulating governance norms, or the limited set of categories readily available for describing AI systems [15, 27].

Prior work has largely examined individual municipal, national, or jurisdictional initiatives, particularly in Western and Global North contexts [5, 11, 33, 47]. We know comparatively less about how schemas and reporting completion vary across instruments, whether schema similarities exhibit temporal, geographic, or institutional patterns, and how the portfolios represented by different instruments vary in composition. For jurisdictions covered by both country-specific and transnational repositories, it also remains unclear whether these sources identify the same systems or construct different representations of governmental AI. Cross-jurisdictional comparison therefore requires sufficient conceptual agreement about what fields represent and how their categories are defined, while preserving differences in provenance and context.

Such comparison requires distinguishing three analytical units. An *instrument* is a sociotechnical device with its own institutional purpose and reporting process [39]. In this study, an instrument may be a register, inventory, or transnational repository. A *schema* specifies the concepts, attributes, and relationships through which information can be represented [7]. A *source record* is an individual structured entry containing values for the fields defined by that schema [8], which was produced through a particular instrument’s collection process. These distinctions allow us to compare schema breadth separately from record-level completion.

3 Methods

Our analysis proceeded across three levels corresponding to the research questions. First, we compared instruments’ schema breadth and field completion and examined temporal and geographic patterns in schema similarity. Second, we compared the lifecycle composition and organizational concentration of the portfolios represented by each instrument. Third, for jurisdictions covered by both country-specific and transnational sources, we identified corresponding records and compared their overlap, disclosure completion, and lifecycle composition.

3.1 Source Identification and Eligibility

Using the Dux Distributed Global Search Python package¹, we searched combinations of each country’s name² and terms, such as *AI register*, *AI inventory*, *algorithm register*, and *automated decision system inventory*, between June 26, 2026 and July 11, 2026. We included resources that (1) operated at the national or federal level, sought countrywide coverage, or cataloged multiple jurisdictions through a common transnational framework; (2) enumerated identifiable AI systems, algorithms, projects, or use cases; and (3) were publicly accessible or publicly documented. We excluded city- or agency-specific registers, general AI policy documents, and reports containing only aggregate adoption statistics. Countrywide and transnational sources could nevertheless include records from local or agency-level institutions.

3.2 Data Provenance and Corpus Construction

Our provenance audit identified two principal source families. We identified 17 country-specific sources from Argentina, Australia, Brazil, Canada, Chile, China, Colombia, Costa Rica, France, India, Mexico, the Netherlands, Norway, Singapore, the United Kingdom (UK), the United States (US), and Uruguay. Australia’s use case library³ requires special membership, and China’s Beian⁴ and India’s AIKosh⁵ provided public lookup interfaces, but we could not identify an official release of their underlying datasets or APIs and therefore excluded these three as standalone sources. For several other countries without directly downloadable datasets, including Argentina, Brazil, and Norway, we converted publicly available RSS⁶ or Omeka feeds and HTML pages into tabular form. The second family consisted of repositories compiled by the Public Sector Tech Watch (PSTW), the Policy Innovation Lab (PIL), and the US Agency for International Development (USAID)⁷. From these repositories, we extracted records covering 65 countries, along with EU-level and global records, where AI was identified as the primary technology. Across both source families, the corpus contained information about AI use in 72 countries, along with EU-level and global records (Figure 1). Some countries, including France, Mexico, and the US, appeared in both families. Although AIKosh was excluded as a standalone source, we retained records about Indian AI applications found in the transnational repositories. Following prior work [56, 64], we used GPT-5 to translate non-English documents into English.

¹ <https://pypi.org/project/ddgs/> ² <https://www.un.org/en/about-us/member-states>

³ <https://www.govai.gov.au/explore/use-cases>

⁴ <https://beian.miit.gov.cn/Integrated/index>

⁵ <https://aikosh.indiaai.gov.in/home/use-cases/all> ⁶ Really Simple Syndication: an automated XML web feed for content distribution. ⁷ The PIL and USAID inventories were retrieved from the same source and treated as combined.

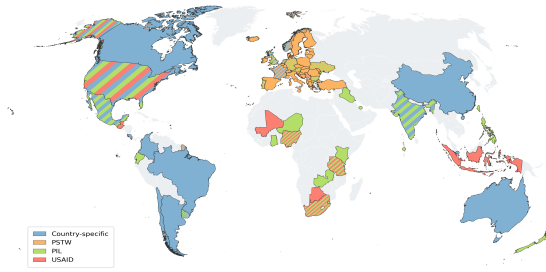


Figure 1: Geographic coverage of AI register sources. Diagonal stripes indicate countries covered by multiple sources.

Because the sources varied in both schema and reporting unit—describing systems, use cases, projects, algorithms, or applications—we treated each row as an observation produced by a particular inventorying process rather than as a directly comparable measure of AI adoption. We retained the original source fields and constructed a harmonized analytical layer containing 23 conceptually comparable fields. These fields covered provenance and identification; system description, responsible organization, lifecycle, administrative function, and technology; intended users and development responsibility; data, procurement, and human involvement; legal basis, risk, and evaluation; source-code availability; and public contact mechanisms. Fields were mapped at the conceptual rather than literal-name level. For example, “*human review required*”, “*human intervention*”, and “*human involved in the final decision*” were mapped to a shared human-involvement field while preserving their source-specific wording and definitions. We harmonized administrative-function and technology categories only when correspondence between source categories was sufficiently direct. We similarly mapped lifecycle values, where possible, into planned, in development, pilot, operational, completed or retired, and unknown. We did not impute absent information, retained ambiguous values in their original form and excluded from cross-source comparisons.

The resulting corpus contained 8,368 records, which did not necessarily represent unique systems because the same system could appear in multiple sources. Each record therefore received a provenance key combining its source instrument, original identifier, and reporting unit. To analyze missingness, we distinguished *structural omission*, where an instrument did not request a field, from *reporting omission*, where the field existed but was unpopulated. This distinction separated *structural visibility*, the fields a register made available for disclosure, from *reporting realization*, the extent to which those fields were completed in individual records.

3.3 Comparative Data Analysis

3.3.1 Register Design and Disclosure. We represented each instrument as a 23-field binary vector indicating whether each harmonized field was present in its schema. We calculated pairwise Jaccard similarity and used average-linkage agglomerative clustering based on Jaccard distance to order an exploratory register-by-field heatmap. We included adoption date and geographic region as annotations to aid interpretation.

To examine diffusion-consistent convergence, we used node-label permutation tests [36], which preserve the dependence structure created by instrument pairs sharing the same nodes. The temporal test compared mean Jaccard similarity between pairs adopted within two years of one another and pairs adopted more than two years apart. For each test, we permuted the corresponding adoption year across instrument nodes, recalculated the difference in group means, and used a one-sided test of whether temporally proximate exhibited greater schema similarity.

For each field present in an instrument’s schema, we calculated completion as the proportion of eligible records containing a reported value. We visualized schema presence and field completion as separate register-by-field matrices, distinguishing narrow schemas, broad and consistently populated schemas, and broad schemas with substantial reporting omission.

3.3.2 Lifecycle Composition and Organizational Concentration. Because governmental level, institutional sector, administrative function, and technology type were represented using highly source-specific vocabularies, cross-instrument portfolio analysis focused on the harmonized lifecycle variable and within-instrument organizational concentration. We calculated record-weighted lifecycle totals to describe the assembled corpus. We then compared the lifecycle distributions of instruments with usable harmonized lifecycle information using pairwise Jensen–Shannon divergence (JSD) [40]. JSD ranges from zero for identical distributions to one for distributions with no overlap.

To examine whether disclosed records were broadly distributed or concentrated among a small number of reporting organizations, we calculated the normalized Herfindahl–Hirschman Index (HHI) for each instrument [9]. The normalized HHI measure ranges from zero for an equal distribution to one for complete concentration. We also reported the shares contributed by the largest and five largest organizations to aid interpretation.

3.3.3 Cross-Source Representation. For jurisdictions represented in both country-specific and transnational sources, we generated potentially corresponding record pairs using normalized system titles, responsible-organization names, URLs, dates, and semantic similarity between descriptions. This candidate-generation procedure produced 490 pairs, which we manually reviewed and confirmed 81 one-to-one matches. We accepted a candidate only when comparison of the available identifying fields indicated that both records described the same underlying system or use case; ambiguous candidates were not treated as matches. Records without a verified counterpart were classified as source-unique for that particular source pair, which does not establish that the underlying system was absent from the other source. For each source pair, we calculated the number of verified shared records, the number and proportion of records unique to each source, Jaccard overlap, and directional coverage. Then, we compared disclosure completion and lifecycle composition between shared and source-unique records.

4 Results

We present the results in three parts. First, we compare schema breadth, field completion, and patterns of schema similarity. Second, we examine differences in lifecycle composition and organizational

concentration across instruments. Third, we assess the overlap between country-specific and transnational sources and compare the records they share with those unique to each source.

4.1 Register Design and Schema Similarity

Across all instruments, we identified 23 distinct disclosure fields. The schemas shared a basic descriptive core: every instrument disclosed a system name and responsible organization, 15 identified the technology type, and 14 included a description and lifecycle status. Fields related to contestability and external accountability were considerably less common. Only one instrument (the US) included an appeal mechanism, while two instruments included risks (the Netherlands and the US), legal basis (France and the Netherlands), and external evaluation (France and the US). Schema breadth ranged from five fields in Argentina to 21 in the US.

Across the 120 unique instrument pairs, Jaccard schema similarity ranged from 0.21 to 0.85 (mean $J = 0.53$). The most similar pair, Chile and Colombia, represented two Latin American countries. However, the second-most similar pair, Mexico–PSTW ($J = 0.83$), crossed geographic regions and source types. Hierarchical ordering revealed no clear geographic groupings (Figure 2).

Schema presence did not necessarily mean that the corresponding information was consistently reported. Across 187 eligible schema–field combinations, 115 (61.5%) were complete for every eligible record, whereas 24 (12.8%) had completion rates below 50%. Mean completion across available fields ranged from 54.9% in the US inventory to 100% in the Argentina, Colombia, Mexico, and Singapore inventories. These comparisons must be considered alongside schema breadth: the US inventory contained 21 disclosure fields, whereas the Argentina and Singapore inventories contained only five and eight, respectively. Thus, broad schemas did not necessarily produce more complete disclosure.

The 82 instrument pairs adopted within two years of one another had a mean similarity of 0.540, compared with 0.504 among the 38 pairs adopted more than two years apart. This difference was small and not statistically significant (mean difference = 0.036, one-sided permutation $p = 0.184$), providing no evidence that temporal proximity was associated with greater schema similarity.

4.2 Lifecycle Divergence and Organizational Concentration

Across instruments, responsible organization and lifecycle stage were reported for 98.06% and 88.85% of records, respectively. Among records with a reported lifecycle stage, operational systems constituted the largest category (49.5%), followed by systems in development (26.9%), pilots (16.6%), retired systems (6.5%), and planned systems (0.3%). Lifecycle distributions varied sharply across instruments. Operational systems constituted all 10 usable records in Costa Rica, 97.5% of records in Mexico, and 91% in the Netherlands. In the second source family, 79% and 42.56% of systems reported by PIL/USAID and PSTW, respectively, were in use. Colombia presented a contrasting pattern, with 65.5% of its records classified as in development. Systems in development also constituted 50% of Uruguay's, 44% of Canada's, and 44.6% of the US inventory. Pilots were especially prominent in PSTW (40.9%).

Across 14 eligible instruments and 91 instrument pairs, lifecycle JSD ranged from 0.013 between Costa Rica and Mexico to 0.845 between Colombia and Costa Rica, with a mean of 0.250. These two comparisons should be interpreted cautiously because Costa Rica contributed only 10 usable lifecycle records. Other large divergences involved Colombia's development-heavy portfolio, including Colombia–Mexico (0.798), Colombia–Chile (0.711), and Colombia–Netherlands (0.698). Among the most similar larger portfolios were France and the UK (0.025) and the UK and PSTW (0.038).

Organizational concentration also varied substantially. Normalized HHI values ranged from 0 to 0.17 (median = 0.018). Singapore's technology-offer portfolio had the highest concentration ($HHI = 0.17$): the National University of Singapore accounted for 46.7% of its records, and the five most represented organizations collectively accounted for 93.3%. Uruguay (0.07), PIL/USAID (0.06), and the US (0.05) exhibited the next-highest values. In the US inventory, the largest organization, the National Aeronautics and Space Administration, contributed 14.6% of records, while the five largest organizations collectively accounted for 58.2%. In contrast, PSTW contained records from 1,415 organizations in the EU region and had a normalized HHI of 0.0006. Its largest organization, the Portuguese Agency for Administrative Modernization, accounted for only 1.1% of records, and its five largest organizations together accounted for 4.2%. Mexico (0.002), Chile (0.003), and the Netherlands (0.008) also exhibited comparatively dispersed portfolios.

4.3 Cross-Source Representation and Overlap

Cross-source comparison was possible for eight pairs across seven countries and the two transnational repositories, PSTW and PIL/USAID. Norway exhibited the greatest overlap with PSTW: 48 records appeared in both sources (Jaccard overlap = 0.245). Although 73.8% of PSTW's Norwegian records also appeared in Norway's country-specific inventory, PSTW contained counterparts for only 26.8% of the larger Norwegian inventory. France exhibited more modest overlap with PSTW: its 17 verified matches produced a Jaccard overlap of 0.093 and represented 14.2% of the French inventory and 21.3% of PSTW's France records. Overlap was substantially lower for the Netherlands and the UK. PIL/USAID exhibited almost no overlap with the corresponding country-specific inventories: only one of its 10 US records matched the federal inventory, and no matches were verified for Mexico, the UK, or Uruguay. The latter country subsets contained only two to four PIL/USAID records, so these null overlaps should be interpreted in light of their limited size. Appendix Table 1 reports the results for every source pair.

Shared records were not consistently more complete than source-unique records. Within the French inventory, shared records had a mean disclosure-completion rate of 87.1%, compared with 82.3% among source-unique records. A similar difference appeared in Norway (70.8% vs 62.3%). In the Netherlands, however, shared records were less complete than source-unique records (77.8% vs 84.3%), while the difference in the UK was negligible (100.0% vs 99.6%). On the PSTW side of these comparisons, shared and source-unique records generally differed by no more than two percentage points.

Lifecycle composition also differed in some source pairs. In the Netherlands register, six of the nine shared records (66.7%) were post-development, compared with 3.8% of source-unique records;

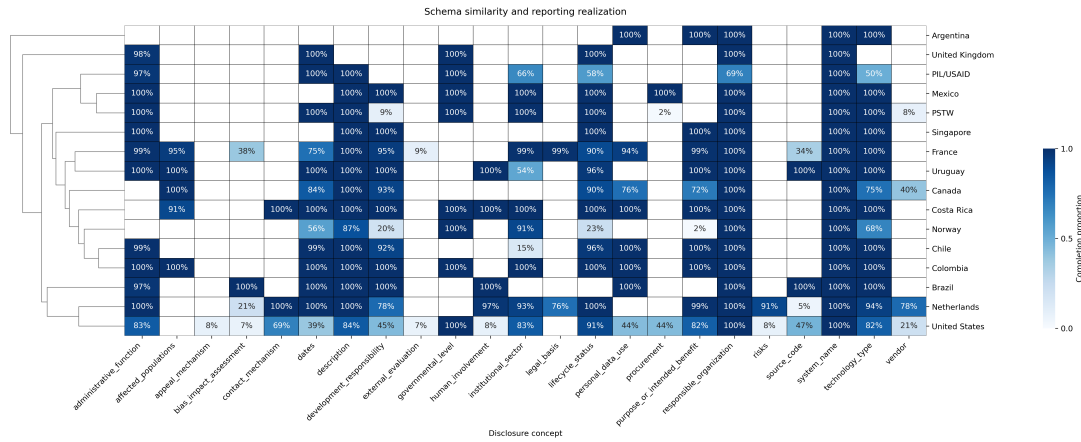


Figure 2: Schema similarity and record-level completion across AI inventories. Blank cells indicate fields' absence from schema.

91.5% of source-unique Dutch records were operational. Within PSTW's Norway subset, 68.8% of shared records were classified as pilots, compared with 35.3% of source-unique records. Overall, country-specific and transnational sources intersected selectively rather than reproducing the same portfolios.

5 Discussion

While most AI registers shared a descriptive core, they differed in scope, schema, field completion, and the organizations and systems represented. Schema breadth did not necessarily produce greater transparency: broad schemas often contained substantial missingness, whereas narrower schemas could achieve higher completion by requesting less information. The lack of strong evidence for temporal or regional groupings in the hierarchical ordering complicates straightforward accounts of policy diffusion. As local administrative settings translate and adapt policy scope, terminology, and implementation [42], schema similarity alone cannot establish borrowing, and the diffusion of a common policy form, such as register schema, need not produce convergence in operational design. Our cross-source comparisons further demonstrate that visibility is distributed across a source ecology, as country-specific and transnational sources frequently disclosed different systems within the same jurisdiction. Low overlap between such sources may reflect different definitions or institutional reach rather than inaccuracy. Here, neither source type should be treated automatically as ground truth, but together be considered as evidence of representational complementarity. In such cases, rather than simply aggregating AI registers, triangulation can expand visibility when provenance and source-specific definitions are preserved. Hence, AI registers should be understood not as censuses, but as institutionally produced and complementary representations of public-sector AI. Drawing on infrastructure and policy instrumentation scholarship, we propose the **Layered Visibility Framework**, which conceptualizes how this representation is produced through four interconnected layers.

The first layer, *institutional scope*, determines which entities and activities enter a register. Since reporting entities may decide differently to include pilot or retired systems and use mandatory,

voluntary, centrally compiled, or independently assembled reporting arrangements, differences in reporting mandates, participation, institutional coverage, collection model, or units of observation may shape organizational concentration and lifecycle composition independent of underlying patterns of governmental AI.

The second layer, *schema affordances*—the possibilities for disclosure enabled and constrained by its design [19]—determines what an instrument can represent. Most schemas included names, responsible organizations, descriptions, technology types, and lifecycle statuses, suggesting a shared descriptive core. In contrast, only a few registers' schemas explicitly requested information about risks, legal bases, external evaluations, and appeal processes. Because information absent from a schema cannot be disclosed systematically, schema design establishes the boundaries of accountable knowledge [38, 53]. The observed variation in the number of fields shows that instruments identically labeled as AI registers may enable substantially different forms of scrutiny.

The third layer, *reporting realization*, concerns whether these possibilities are realized in individual records. Schema presence and record-level disclosure were not equivalent. For example, the US had the broadest schema but many fields (e.g., appeal process) had a low completion rate, whereas several narrower instruments fully populated their eligible fields. Thus, registers can be broad but thin or narrow but dense, while the expectation is that they be broad and consistently populated. Counting fields alone overstates the visibility afforded by broad but incomplete schemas, while completion rates alone favor schemas that request less information.

These layers produce the fourth: the *public representation* of governmental AI. Lifecycle composition and organizational concentration show that register records constitute portfolios produced through particular disclosure arrangements. Selective cross-source overlap also demonstrates that different instruments can produce different representations of the same jurisdiction. Hence, register records are not unmediated representations of governmental AI.

The framework extends the evaluation of registers beyond transparency to representational quality and interoperability. Instead of imposing a universal template, registers could combine a clearly

defined core—such as system ID, responsible organization, lifecycle stage, and last-updated date—with modules covering procurement, data use, human involvement, risk, evaluation, and contestability. Distinguishing structural absence from missing reporting, publishing category definitions, describing coordination processes, documenting schema changes, and preserving source links would support comparison without erasing institutional differences.

6 Future Work and Conclusion

AI registers are becoming part of the administrative infrastructure through which governments render digital transformation knowable and governable. However, similarly labeled instruments can produce distinct and potentially complementary representations of governmental AI. Besides being necessarily partial, the publicly accessible instruments we studied capture particular points in time, which limited our record linkage capacity and our ability to map source-specific vocabularies. Future work should track how registers and schemas change over time to accommodate mappings for heterogeneous categories while preserving corresponding institutional meanings. Interoperability among these instruments depends on shared concepts, clear definitions, preserved provenance, and traceable differences for accountable digital transformation.

Acknowledgments

We used OpenAI's GPT-5.6 to translate non-English texts and assist with language editing.

References

- [1] Mike Ananny and Kate Crawford. 2018. Seeing without knowing: Limitations of the transparency ideal and its application to algorithmic accountability. *new media & society* 20, 3 (2018), 973–989.
- [2] Cecilia Aragon, Shion Guha, Marina Kogan, Michael Muller, and Gina Neff. 2022. *Human-centered data science: an introduction*. MIT Press.
- [3] Reuben Binns. 2018. Algorithmic accountability and public reason. *Philosophy & technology* 31, 4 (2018), 543–556.
- [4] William Boag, Harini Suresh, Bianca Lepe, and Catherine D'Ignazio. 2022. Tech worker organizing for power and accountability. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*. 452–463.
- [5] Corinne Cath and Fieke Jansen. 2021. Dutch Comfort: The limits of AI governance through municipal registers. *arXiv preprint arXiv:2109.02944* (2021).
- [6] Shreya Chappidi, Jennifer Cobbe, Chris Norval, Anjali Mazumder, and Jatinder Singh. 2025. Accountability Capture: How Record-Keeping to Support AI Transparency and Accountability (Re) shapes Algorithmic Oversight. In *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, Vol. 8. 554–566.
- [7] Peter Pin-Shan Chen. 1976. The entity-relationship model—toward a unified view of data. *ACM transactions on database systems (TODS)* 1, 1 (1976), 9–36.
- [8] Edgar F Codd. 1970. A relational model of data for large shared data banks. *Commun. ACM* 13, 6 (1970), 377–387.
- [9] Daniel Cracau and José E Durán Lima. 2016. On the normalized herfindahl-hirschman index: a technical note. *International Journal on Food System Dynamics* 7, 4 (2016), 382–386.
- [10] Dipto Das, Matthew Tamura, Syed Ishtiaque Ahmed, and Shion Guha. 2026. "Is This Not Enough?": Asymmetries in Institutional Accountability and Collective Sensemaking in the Case of Canada's Algorithmic Visa Triage System. In *Proceedings of the ACM SIGCAS/SIGCHI Conference on Computing and Sustainable Societies* (Virtual Event, 2026-07-27/2026-07-31) (COMPASS '26). Association for Computing Machinery, New York, NY, USA. doi:10.1145/3811242.3819106
- [11] Dipto Das, Christelle Tessono, Syed Ishtiaque Ahmed, and Shion Guha. 2026. Bureaucratic Silences: What the Canadian AI Register Reveals, Omits, and Obscures. In *Proceedings of the 2026 ACM Conference on Fairness, Accountability, and Transparency (FAccT '26)* (Montréal, Canada) (FAccT '26). Association for Computing Machinery, New York, NY, USA. doi:10.1145/3805689.3812336
- [12] Emily Denton, Alex Hanna, Razvan Amironesei, Andrew Smart, and Hilary Nicole. 2021. On the genealogy of machine learning datasets: A critical history of ImageNet. *Big Data & Society* 8, 2 (2021), 20539517211035955.
- [13] Daria Dergacheva, Vasilisa Kuznetsova, Rebecca Scharlach, and Christian Katzenbach. 2023. One Day in Content Moderation: Analyzing 24 h of Social Media Platforms' Content Decisions through the DSA Transparency Database. (2023).
- [14] Natalia Díaz-Rodríguez, Javier Del Ser, Mark Coeckelbergh, Marcos López De Prado, Enrique Herrera-Viedma, and Francisco Herrera. 2023. Connecting the dots in trustworthy Artificial Intelligence: From AI principles, ethics, and key requirements to responsible AI systems and regulation. *Information Fusion* 99 (2023), 101896.
- [15] Paul J DiMaggio, Walter W Powell, et al. 1983. The iron cage revisited: Institutional isomorphism and collective rationality in organizational fields. *American sociological review* 48, 2 (1983), 147–160.
- [16] Frank Dobbin, Beth Simmons, and Geoffrey Garrett. 2007. The global diffusion of public policies: Social construction, coercion, competition, or learning? *Annu. Rev. Sociol.* 33, 1 (2007), 449–472.
- [17] Chiara Patricia Drolsach and Nicolas Pröllochs. 2024. Content moderation on social media in the EU: Insights from the DSA Transparency Database. In *Companion Proceedings of the ACM Web Conference 2024*. 939–942.
- [18] Zachary Elkins and Beth Simmons. 2005. On waves, clusters, and diffusion: A conceptual framework. *The Annals of the American Academy of Political and Social Science* 598, 1 (2005), 33–51.
- [19] Sandra K Evans, Katy E Pearce, Jessica Vitak, and Jeffrey W Treem. 2017. Explicating affordances: A conceptual framework for understanding affordances in communication research. *Journal of computer-mediated communication* 22, 1 (2017), 35–52.
- [20] Florian Eyert and Paola Lopez. 2023. Rethinking transparency as a communicative constellation. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*. 444–454.
- [21] Heike Felzmann, Eduard Fosch Villaronga, Christoph Lutz, and Aurelia Tamò-Larrieux. 2019. Transparency you can trust: Transparency requirements for artificial intelligence between legal norms and contextual concerns. *Big Data & Society* 6, 1 (2019), 2053951719860542.
- [22] Ana Gagua, Haiko van der Voort, Nihit Goyal, and Alexander Verbraeck. 2025. Responsible AI Governance in the Public Sector: Explaining Contextual Dynamics through a Realist Synthesis Review. In *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*, Vol. 8. 990–1002.
- [23] Timnit Gebru, Jamie Morgenstern, Briana Vecchione, Jennifer Wortman Vaughan, Hanna Wallach, Hal Daumé III, and Kate Crawford. 2021. Datasheets for datasets. *Commun. ACM* 64, 12 (2021), 86–92.
- [24] Samuel Groesch, Alena Birrer, Natascha Just, and Florian Saurwein. 2025. Big data, small answers: How the DSA Transparency Database falls short of its regulatory objectives. *Telecommunications Policy* (2025), 103088.
- [25] M Haataja, L van de Fliert, and P Rautio. 2020. Public AI Registers: Realising AI transparency and civic participation in government use of AI. *Whitepaper Version 1* (2020).
- [26] Charlotte Högborg. 2024. Stabilizing translucencies: Governing AI transparency by standardization. *Big Data & Society* 11, 1 (2024), 20539517241234298.
- [27] Katharina Holzinger and Christoph Knill. 2005. Causes and conditions of cross-national policy convergence. *Journal of European public policy* 12, 5 (2005), 775–796.
- [28] Independent International Scientific Panel on Artificial Intelligence. 2026. Preliminary Report of the Independent International Scientific Panel on AI: Evidence-Based Assessment of Opportunities, Risks and Impacts of Artificial Intelligence. <https://www.un.org/independent-international-scientific-panel-ai/en/preliminary-report>
- [29] Simon Jarvers and Orestis Papakyriakopoulos. 2026. Engaged AI Governance: Addressing the Last Mile Challenge Through Internal Expert Collaboration. In *The 2026 ACM Conference on Fairness, Accountability, and Transparency*. 4415–4438.
- [30] Anna Jobin, Marcello Ienca, and Effy Vayena. 2019. The global landscape of AI ethics guidelines. *Nature machine intelligence* 1, 9 (2019), 389–399.
- [31] Deborah G Johnson. 2021. Algorithmic accountability in the making. *Social Philosophy and Policy* 38, 2 (2021), 111–127.
- [32] Rishabh Kaushal, Jacob Van De Kerkhof, Catalina Goanta, Gerasimos Spanakis, and Adriana Iamnitchi. 2024. Automated transparency: A legal and empirical analysis of the digital services act transparency database. In *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency*. 1121–1132.
- [33] Nigel Kingsman, Emre Kazim, Ali Chaudhry, Airlie Hilliard, Adriano Koshiyama, Roseline Polle, Giles Pavey, and Umar Mohammed. 2022. Public sector AI transparency standard: UK Government seeks to lead by example. *Discover Artificial Intelligence* 2, 1 (2022), 2.
- [34] Rob Kitchin and Gavin McArdle. 2016. What makes Big Data, Big Data? Exploring the ontological characteristics of 26 datasets. *Big data & society* 3, 1 (2016), 2053951716631130.
- [35] René F Kizilec. 2016. How much information? Effects of transparency on trust in an algorithmic interface. In *Proceedings of the 2016 CHI conference on human factors in computing systems*. 2390–2395.
- [36] David Krackhardt. 1988. Predicting with networks: Nonparametric multiple regression analysis of dyadic data. *Social networks* 10, 4 (1988), 359–381.

- [37] Joshua A Kroll. 2021. Outlining traceability: A principle for operationalizing accountability in computing systems. In *Proceedings of the 2021 ACM Conference on fairness, accountability, and transparency*. 758–771.
- [38] Blaine Kuehnert, Rachel Kim, Jodi Forlizzi, and Hoda Heidari. 2025. The “Who”, “What”, and “How” of Responsible AI Governance: A Systematic Review and Meta-Analysis of (Actor, Stage)-Specific Tools. In *Proceedings of the 2025 ACM Conference on Fairness, Accountability, and Transparency*. 2991–3005.
- [39] Pierre Lascoumes and Patrick Le Galès. 2007. Introduction: Understanding public policy through its instruments—From the nature of instruments to the sociology of public policy instrumentation. *Governance* 20, 1 (2007), 1–21.
- [40] Jianhua Lin. 1991. Divergence measures based on the Shannon entropy. *IEEE Transactions on Information theory* 37, 1 (1991), 145–151.
- [41] Martino Maggetti and Fabrizio Gilardi. 2016. Problems (and solutions) in the measurement of policy diffusion mechanisms. *Journal of Public Policy* 36, 1 (2016), 87–107.
- [42] David Marsh and Jason C Sharman. 2013. Policy diffusion and policy transfer. In *New directions in the study of policy transfer*. Routledge, 32–51.
- [43] Angelina McMillan-Major, Emily M Bender, and Batya Friedman. 2024. Data statements: From technical concept to community practice. *ACM Journal on Responsible Computing* 1, 1 (2024), 1–17.
- [44] Ines Mergel, Noella Edelmann, and Nathalie Haug. 2019. Defining digital transformation: Results from expert interviews. *Government information quarterly* 36, 4 (2019), 101385.
- [45] Margaret Mitchell, Simone Wu, Andrew Zaldivar, Parker Barnes, Lucy Vasserman, Ben Hutchinson, Elena Spitzer, Inioluwa Deborah Raji, and Timnit Gebru. 2019. Model cards for model reporting. In *Proceedings of the conference on fairness, accountability, and transparency*. 220–229.
- [46] Erina Seh-Young Moon, Devansh Saxena, Dipto Das, and Shion Guha. 2025. The Datafication of Care in Public Homelessness Services. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. 1–16.
- [47] Maya Murad. 2021. *Beyond the “Black Box”: Enabling Meaningful Transparency of Algorithmic Decision-Making Systems through Public Registers*. Ph. D. Dissertation. Massachusetts Institute of Technology.
- [48] Esther Nieuwenhuizen. 2024. Algorithm Registers: A Box-Ticking Exercise or Meaningful Tool for Transparency? *Information Policy* 29, 4 (2024), 415–433.
- [49] Chris Norval, Kristin Cornelius, Jennifer Cobbe, and Jatinder Singh. 2022. Disclosure by Design: Designing information disclosures to support meaningful transparency and accountability. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*. 679–690.
- [50] OECD. 2026. Digital Government Outlook 2026: From Foundations to Transformational Impact. <https://doi.org/10.1787/0496b2bc-en>
- [51] Open Government Partnership. 2021. Building Public Algorithm Registers: Lessons Learned from the French Approach. <https://www.opengovpartnership.org/stories/building-public-algorithm-registers-lessons-learned-from-the-french-approach/>. [Accessed: June 11, 2026].
- [52] Yulu Pi, Wenlong Li, and Jatinder Singh. 2026. Understanding the Role of Algorithm Registers in AI Governance Through Comparative Analysis of China and the UK. *arXiv preprint arXiv:2606.00035* (2026).
- [53] Lindsay Poirier. 2021. Reading datasets: Strategies for interpreting the politics of data signification. *Big Data & Society* 8, 2 (2021), 20539517211029322.
- [54] Lindsay Poirier. 2022. Accountable data: The politics and pragmatics of disclosure datasets. In *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*. 1446–1456.
- [55] Inioluwa Deborah Raji, Andrew Smart, Rebecca N White, Margaret Mitchell, Timnit Gebru, Ben Hutchinson, Jamila Smith-Loud, Daniel Theron, and Parker Barnes. 2020. Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing. In *Proceedings of the 2020 conference on fairness, accountability, and transparency*. 33–44.
- [56] Nils Rehlinger. 2026. A Fine-Grained Linguistic Evaluation of Low-Resource Luxembourgish–English MT. In *Proceedings for the Ninth Workshop on Technologies for Machine Translation of Low Resource Languages (LoResMT 2026)*. 138–150.
- [57] Devansh Saxena, Karla Badillo-Urquiola, Pamela J Wisniewski, and Shion Guha. 2021. A framework of high-stakes algorithmic decision-making for the public sector developed through a case study of child-welfare. *Proceedings of the ACM on Human-Computer Interaction* 5, CSCW2 (2021), 1–41.
- [58] Renée Sieber, Ana Brandusescu, and Jonathan van Geuns. 2026. Building AI Governance in Municipalities from the Ground Up.
- [59] Diane Stone. 2004. Transfer agents and global networks in the ‘transnationalization’ of policy. *Journal of European public policy* 11, 3 (2004), 545–566.
- [60] Amaury Trujillo, Tiziano Fagni, and Stefano Cresci. 2025. The DSA Transparency Database: Auditing self-reported moderation actions by social media. *Proceedings of the ACM on Human-Computer Interaction* 9, 2 (2025), 1–28.
- [61] Michael Veale, Max Van Kleek, and Reuben Binns. 2018. Fairness and accountability design needs for algorithmic support in high-stakes public sector decision-making. In *Proceedings of the 2018 chi conference on human factors in computing systems*. 1–14.
- [62] Janet Vertesi, Danah Boyd, Alex S Taylor, and Benjamin Shestakofsky. 2026. Reckoning with the Political Economy of AI: Avoiding Decoys in Pursuit of

Accountability. In *The 2026 ACM Conference on Fairness, Accountability, and Transparency*. 2186–2205.

- [63] Gregory Vial. 2021. Understanding digital transformation: A review and a research agenda. *Managing digital transformation* (2021), 13–66.
- [64] Longyue Wang, Chenyang Lyu, Tianbo Ji, Zhirui Zhang, Dian Yu, Shuming Shi, and Zhaopeng Tu. 2023. Document-level machine translation with large language models. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*. 16646–16661.
- [65] Elizabeth Anne Watkins, Emanuel Moss, Jacob Metcalf, Ranjit Singh, and Madeleine Clare Elish. 2021. Governing algorithmic systems with impact assessments: Six observations. In *Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society*. 1010–1022.
- [66] Hilde Weerts, Raphaële Xenidis, Fabien Tarissan, Henrik Palmer Olsen, and Mykola Pechenizkiy. 2024. The neutrality fallacy: When algorithmic fairness interventions are (not) positive action. In *Proceedings of the 2024 ACM conference on fairness, accountability, and transparency*. 2060–2070.
- [67] Maranke Wieringa. 2020. What to account for when accounting for algorithms: a systematic literature review on algorithmic accountability. In *Proceedings of the 2020 conference on fairness, accountability, and transparency*. 1–18.
- [68] Amy Winecoff and Miranda Bogen. 2025. Improving governance outcomes through AI documentation: Bridging theory and practice. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. 1–18.

A Appendix

Table 1: Overlap between country-specific and transnational sources

Country-specific	Transnational source	Records (C/T)	Shared	Jaccard	Coverage (%) (C/T)
France	PSTW	120/80	17	0.093	14.2/21.3
Mexico	PIL/USAID	119/4	0	0.000	0.0/0.0
Netherlands	PSTW	1507/168	9	0.005	0.6/5.4
Norway	PSTW	179/65	48	0.245	26.8/73.8
UK	PIL/USAID	133/3	0	0.000	0.0/0.0
UK	PSTW	133/135	6	0.023	4.5/4.4
US	PIL/USAID	3611/10	1	0.0003	0.03/10.0
Uruguay	PIL/USAID	28/2	0	0.000	0.0/0.0

Notes: C/T denotes country-specific/transnational. Coverage is the percentage of records in each source with a verified counterpart. Jaccard overlap is the number of shared records divided by the total distinct records across the source pair.