

Data production and ecosystem

Data governance and the Datasphere: Literature review¹

By Datasphere Initiative²

Introduction

In recent years, the term “data governance” has garnered growing attention. It has moved from being a niche topic, addressed solely as a technical aspect of data-sharing projects, or within enterprise information communication technology (ICT) disciplines as a companion to data management, to becoming the overarching container for thinking about both data protection and access to data.

Data governance has emerged as a key framework within which to address both the opportunities and risks of data collection, sharing, and use. This reflects a growing recognition of the importance of data within wider processes of governance, and of the potential power data has as both a resource for progress and a catalyst of harm when misused.

However, the relatively rapid convergence of interests from policymakers, technologists, activists, and practitioners on “data governance” comes with some challenges. Different agendas, conceptualizations, concerns, and areas of emphasis collide, and there is, as yet, no coherent field of data governance research.

By providing an initial map of who is writing about data governance, and the kinds of topics being addressed, this paper offers the groundwork for a response to the call from de La Chapelle and Porciuncula (2021, p. 3) for work on data governance that can “reframe the discussion, harness emerging innovative approaches, and engage in a much needed global, multistakeholder, and cross-sectoral debate.”

To support the reframing, this paper also looks at the emerging conceptual framework of the Datasphere, which is understood as “the complex system encompassing all types of data and their dynamic interactions with human groups and norms” (de La Chapelle & Porciuncula, 2022, p. 3). The conceptual shift this introduces invites a move from discussing relatively flat notions of “data governance,” to “governance of the Datasphere”: Which brings into focus the interaction of datasets, norms, and human groups.

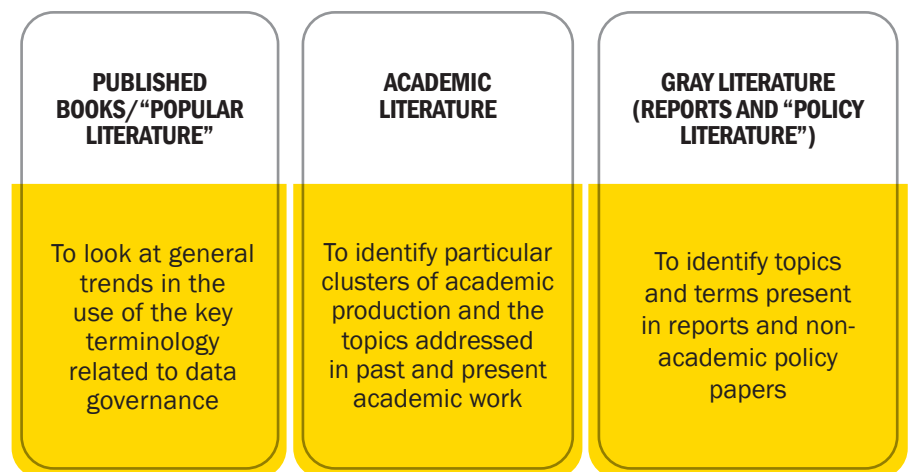
¹ The edited version of the homonymous work published by the Datasphere Initiative (DI). Available at: <https://www.thedatasphere.org/datasphere-publish/data-governance-and-the-datasphere/>

² This report was prepared by Tim Davies, as the output of his consultancy and fellowship work. It received guidance and input from Carolina Rossini, director for Research and Partnerships at DI. Tim is the director of Research at Connected by data, a non-profit company based in the United Kingdom (UK). He has an MSc in Social Science of the Internet from the Oxford Internet Institute and has been a fellow of the Berkman Klein Center for Internet and Society, of Harvard University. Copyright of The Datasphere Initiative Foundation (2022).

Methodology

This review deploys several overlapping strategies for providing an overview of current writing on data governance. Whilst the analysis that follows draws primarily on academic literature, published books – via the Google Books *corpus* – and gray literature – via a *corpus* based on the Datasphere Governance Atlas (Datasphere Initiative, 2022) – are used to provide complementary insights.

Figure 1 – TYPE OF REVIEWED DOCUMENT AND PURPOSE OF THE REVIEW

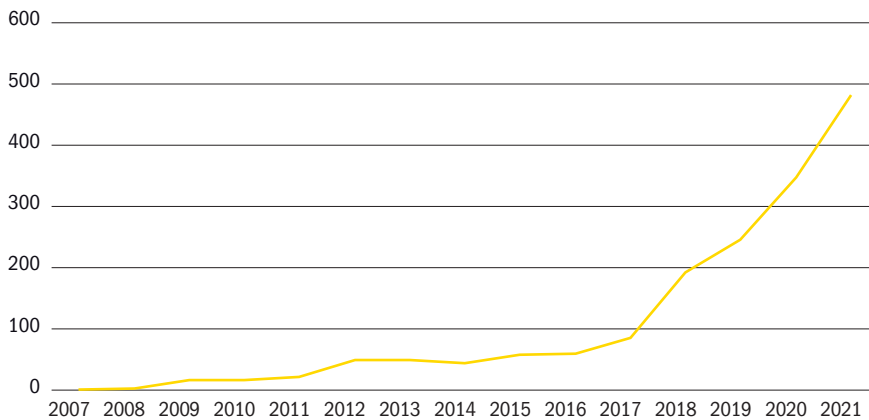


Source: Prepared by the authors.

Data governance is a growing field bringing together formerly distinct areas of focus

Given the proliferation of current work and writing on data governance – the recent Datasphere Governance Atlas (Datasphere Initiative, 2022) counts no less than 261 organizations that focus to some extent on data governance topics –, it may be surprising to note that the term “data governance” has only entered the research and policy lexicon at scale in the last decade. Use of the term in the titles and abstracts of academic papers increased almost five-fold between 2015 and 2021 and looks set to increase even further in 2022.

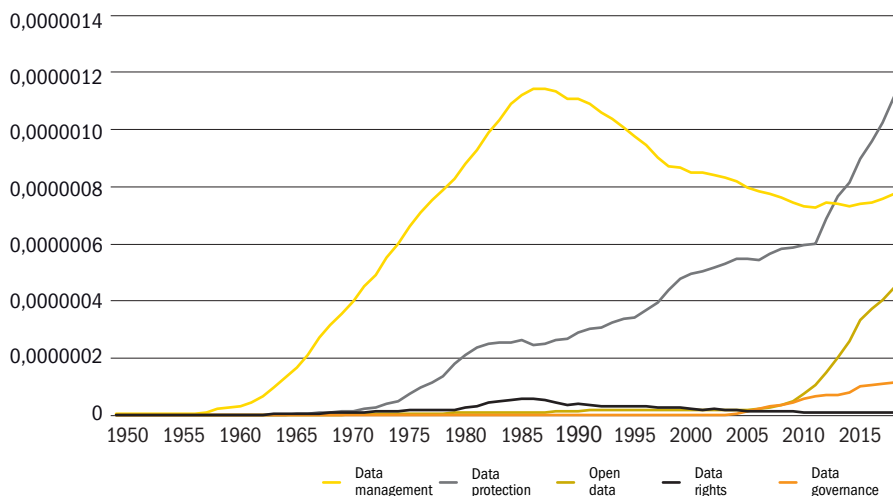
Chart 1 – NUMBER OF PUBLICATIONS RECORDED IN THE DIMENSIONS DATASET THAT INCLUDE THE TERM “DATA GOVERNANCE” IN THEIR TITLE, ABSTRACT, OR KEYWORDS, BY PUBLICATION YEAR (2007-2021)



Source: Prepared by the authors.

The rapid development of data governance discourse does not mean that existing debates have been entirely subsumed within data governance. A look at the presence of other terms in the popular literature highlights the fact that readers are much more likely to encounter work on “data protection” or “data management” in books and technical manuals than they are to find discussions on data governance. Even topics like open data – arguably just one particular approach to governing data –, have received significantly more direct attention in recent years than data governance has.

Chart 2 – COMPARATIVE MENTIONS OF THE TERMS “DATA PROTECTION,” “DATA MANAGEMENT,” “OPEN DATA,” “DATA GOVERNANCE,” AND “DATA RIGHTS” IN GOOGLE BOOKS NGRAM VIEWER CORPUS (1950-2019)



Source: Google (n.d.).

(...) the volume of data governance papers in academic literature has seen higher year-on-year percentage growth than those that focus on either data protection or data management.

These patterns in the popular literature are also broadly mirrored in scientific production, in which many more papers are published about data protection, data management, or open data each year than explicitly use the term “data governance” in their titles or abstracts. However, in recent years the volume of data governance papers in academic literature has seen higher year-on-year percentage growth than those that focus on either data protection or data management.

Ultimately, the continuous but incomplete increase in the use of data governance as a framing term in both research and policy is likely to pause. Authors are adopting the language of data governance from a range of starting points, and this will color what falls within the scope of their data governance definitions and prescriptions. For instance, as mentioned above, much of the literature on data governance within computing and management considers data governance only within the boundaries of an enterprise, whereas social studies and gray literature frequently explore data governance as a social issue. At the same time, the continuous production of work framed in terms of data protection, data management, and open data (to name just a few areas) may have much to contribute to the development of norms, policies, and practices of governing data and the Datasphere, even if a data governance language is not directly adopted.

Previous literature reviews reveal the diversity of the field

Many of the topics that increasingly fall within the broad frame of data governance were formerly discussed in terms of data protection (Greenleaf, 2012), data management (Ladley, 2019; Panian, 2009), or open data (Davies et al., 2019; Verhulst et al., 2020), each with their own particular agenda around privacy, the exploitation of enterprise data assets, and the public re-use of data, respectively. A shift towards framing these topics within the broader scope of data governance responds to recognition of the complexity and trade-offs involved in deciding when and how data should be collected, structured, shared, transferred, used, and deleted.

Efforts to resolve or reframe these trade-offs and tensions have also given rise to a range of new agendas around data sharing (Micheli et al., 2020) and new models of data ownership and stewardship (Delacroix & Lawrence, 2019; Lehtiniemi & Haapoja, 2020; Susha et al., 2017), which fall within the expanding field of data governance. In the gray literature on the topic, a strong normative element is increasingly evident, with the term being linked to wider agendas of good governance and global development. As Pisa et al. (2020, p. 2) state, the ideal of data governance incorporates “rules about how data is collected, analyzed, used, and shared in a way that protects citizens from abuse while supporting innovation, development, and inclusive growth.”

A review of eight past peer-reviewed data governance literature reviews, which were published between 2016 and early 2022, shows this shifting emphasis. While earlier work centered on data governance primarily in terms of data

and information management (Alhassan et al., 2016; Brous et al., 2016), recent work has increasingly addressed data governance as a broader public issue, requiring emphasis on inter-organizational data sharing (Abraham et al., 2019; Benfeldt Nielsen, 2017) and open data (Bozkurt et al., 2022). McCaig and Rezanian (2021, p. 5), however, argue that the literature ultimately remains “indicative of a sparse theoretical and empirical knowledge base” on data governance.

A broad working definition of data governance foregrounds both benefits and harms

Given the breadth of contexts in which data governance must be applied, it is unreasonable to expect a single unified definition that can tie together a single field of study. However, common aspects of data governance can still be distilled. For the purpose of this paper, the following working definition is offered:

- Data governance concerns the rules, processes, and behaviors related to the collection, management, analysis, use, sharing, and disposal of data – personal and/or non-personal.
- Good data governance should both promote benefits and minimize harms at each stage in the relevant data cycles.

At an organizational level this generally translates into: A focus on internal policies and their implementation; compliance with external regulation; and the creation of cross-functional frameworks and responsibilities for managing and extracting value from data as a business asset (Abraham et al., 2019). At the state level – be it national, regional, or international –, this may translate into a focus on the development and implementation of policies, standards, laws, regulations, agreements, and practices that cover the management of data within countries, and the transfer of data across jurisdictional boundaries (Aaronson, 2021). Organizational literature, however, often pays little attention to the state level, and vice-versa.

A number of authors also highlight that governance of data sits amongst a range of wider practical and governance concerns. Wendehorst (2020) describes data governance as one of a number of overlapping frameworks of governance concern in relation to artificial intelligence (AI), considering for instance how the same issue might be explored: Through the lens of data governance (considering how datasets are created, managed, and used); through a lens of AI systems’ design (using the language of bias or adequacy of methods); or by way of a focus on wider social governance (asking questions about the goals and governance of the wider policy areas to which datasets and AI systems relate).

Data governance concerns the rules, processes, and behaviors related to the collection, management, analysis, use, sharing, and disposal of data – personal and/or non-personal.

This highlights the importance of resisting the tendency to treat data as entirely in the abstract: Meaningful data is always about something, and this “something” is also frequently subject to its own governance regime (...).

This highlights the importance of resisting the tendency to treat data as entirely in the abstract: Meaningful data is always about something, and this “something” is also frequently subject to its own governance regime, with which any practical data governance will intersect. Many researchers have arrived at the topic of data governance because of challenges that are largely based on protecting, managing, or sharing data in relation to a particular field of action.

When turning to the academic literature it is important to have an understanding of the extent to which different projects and papers are part of a coherent research agenda, or – by contrast – to which each publication, using the language of data governance, may have developed in isolation from other works related to the topic.

The concept of data governance not only brings together academics previously working on distinct issues of data protection, management, and access, but it has also been invoked in disparate academic fields, from health research to work on international trade. Data governance in these fields can still appear more-or-less as a niche sub-field, rather than as cross-cutting field of inquiry in its own right.

Developing Datasphere narratives can offer a holistic perspective for future work on data governance

This section provides a brief overview of the concept of the Datasphere and explores what it may mean to look at the data governance literature through a Datasphere lens. The report *We Need To Talk About Data* (de La Chapelle & Porciuncula, 2021), draws on a law paper by Bergé et al. (2018a, p. 2) that offered a conceptually expansive, but digitally focused description of the Datasphere, which the paper describes as:

The notion of “Datasphere” proposes a holistic comprehension of all the “information” existing on earth, originating both in natural and socio-economic systems, which can be captured in digital form, flows through networks, and is stored, processed, and transformed by machines.

A desire to move outside of a narrow menu of policy options in part motivated the adoption of a refined Datasphere terminology, described as: “The complex system encompassing all types of data and their dynamic interactions with human groups and norms” (Porciuncula & de La Chapelle, 2022, p. 3).

This formula essentially draws attention to the mutual interactions between digital artifacts (datasets); constituencies and social relationships (human groups); and rules and social expectations (norms) – and to the multiplicity of each aspect. Notably, the model implies governance of one interconnected Datasphere, not many isolated instances, and does so with the purpose of providing a holistic lens into the evolving complexity of data governance and its impact on the creation of value and well-being for all. That is, the Datasphere

is seen as a single complex system (Siegenfeld & Bar-Yam, 2020). Or, going further, as per Porciuncula and de La Chapelle (2022), the Datasphere is a complex adaptive system with emergent dynamics.

Shifting from a discussion of “governing data” to “governing the Datasphere” involves identifying the particular regions of the Datasphere in focus and acknowledging the relationships between data governance in one region (for example, in relation to the individual or the firm), and data governance in other regions and at other levels (for example, organizational, industry, social, national or global). By offering the typology of datasets, human groups, and norms, the Datasphere framework then invites a clearer specification of the specific focus of any governance research and the factors being taken into account in proposing or evaluating particular governance regimes.

Conclusion

This paper offers a starting point for thinking about academic (and some gray literature) writings on data governance. It provides a high-level overview of research clusters and themes addressed in the literature, highlighting that there is, ultimately, not just one data governance field to speak of, but rather a range of distinct fields of work, each responding to thematic or sectoral challenges. While firm and society-level governance of data are broadly two sides of the same coin, relatively little work has explored issues of cross-boundary data governance, leaving a significant gap to be filled.

The paper also suggests that the conceptual framework of the Datasphere has a significant contribution to make to current data governance research and practice, in particular by putting forward the notion of “governance of the Datasphere” as a systems approach to data governance.

References

- Aaronson, S. A. (2021). *Data is disruptive: How data sovereignty is challenging data governance*. Hinrich Foundation. <https://www.wita.org/wp-content/uploads/2021/08/Data-is-disruptive-Hinrich-Foundation-white-paper-Susan-Aaronson-August-2021.pdf>
- Abraham, R., Schneider, J., & vom Brocke, J. (2019). Data governance: A conceptual framework, structured review, and research agenda. *International Journal of Information Management*, 49.
- Alhassan, I., Sammon, D., & Daly, M. (2016). Data governance activities: An analysis of the literature. *Journal of Decision Systems*, 25.
- Benfeldt Nielsen, O. (2017). A comprehensive review of data governance literature. *Selected Papers of the IRIS*, (8). <https://core.ac.uk/download/pdf/301373908.pdf>
- Bergé, J. S., Grumbach, S., & Zeno-Zencovich, V. (2018). The ‘datasphere’, data flows beyond control, and the challenges for law and governance. *European Journal of Comparative Law and Governance*, 5(2). <https://papers.ssrn.com/abstract=3185943>

(...) there is, ultimately, not just one data governance field to speak of, but rather a range of distinct fields of work, each responding to thematic or sectoral challenges.

- Bozkurt, Y., Rossmann, A., & Pervez, Z. (2022). *A literature review of data governance and its applicability to smart cities*. Hawaii International Conference on System Sciences. <http://hdl.handle.net/10125/79666>
- Brous, P., Janssen, M., & Vilminko-Heikkinen, R. (2016). Coordinating decision-making in data management activities: A systematic review of data governance principles. *Lecture Notes in Computer Science*, 9820.
- Datasphere Initiative. (2022). *Datasphere governance atlas*. <https://www.thedatasphere.org/wp-content/uploads/2022/04/Datasphere-Governance-Atlas-2022-Datasphere-Initiative.pdf>
- Davies, T., Walker, S. B., Rubinstein, M., & Perini, F. (Eds.). (2019). *The state of open data: Histories and horizons*. <https://www.idrc.ca/en/book/state-open-data-histories-and-horizons>
- de La Chapelle, B., & Porciuncula, L. (2021). We need to talk about data: Framing the debate around the free flow of data and data sovereignty. *Internet & Jurisdiction Policy Network*.
- de La Chapelle, B., & Porciuncula, L. (2022). *Hello datasphere: Towards a systems approach to data*. Datasphere Initiative. <https://www.thedatasphere.org/datasphere-publish/hello-datasphere/>
- Delacroix, S., & Lawrence, N. D. (2019). Bottom-up data trusts: Disturbing the “one size fits all” approach to data governance. *International Data Privacy Law*, 9(4).
- Google. (s.d.). *Google Books Ngram Viewer*. https://books.google.com/ngrams/graph?content=data+protection%2Cdata+management%2Copen+data%2Cdata+governance%2Cdata+rights&year_start=1950&year_end=2019&corpus=en-2019&smoothing=3
- Greenleaf, G. (2012). Global data privacy laws: 89 countries, and accelerating. *Privacy Laws & Business International Report*, (115). <https://papers.ssrn.com/abstract=2000034>
- Ladley, J. (2019). *Data governance: How to design, deploy, and sustain an effective data governance program*. Academic Press. <https://books.google.com?id=AkW9DwAAQBAJ>
- Lehtiniemi, T., & Haapoja, J. (2020). Data agency at stake: MyData activism and alternative frames of equal participation. *New Media & Society*, 22(1).
- McCaig, M., & Rezaia, D. (2021). *A scoping review on data governance*. International Conference on IoT Based Control Networks & Intelligent Systems.
- Micheli, M., Ponti, M., Craglia, M., & Suman, A. B. (2020). Emerging models of data governance in the age of datafication. *Big Data & Society*, 7(2).
- Panian, Z. (2009). Recent advances in data management. *WSEAS Transactions on Computers*, 9(7), 1061–1071.
- Pisa, M., Dixon, P., Ndulu, B., & Nwankwo, U. (2020). *Governing data for development: Trends, challenges, and opportunities*. Center for Global Development. <https://www.cgdev.org/sites/default/files/governing-data-development-trends-challenges-and-opportunities.pdf>
- Porciuncula, L., & de La Chapelle, B. (2022). Hello datasphere: Towards a systems approach to data governance. *The Datasphere Initiative*. <https://www.thedatasphere.org/news/hello-datasphere-towards-a-systems-approach-to-data-governance/>
- Siegenfeld, A. F., & Bar-Yam, Y. (2020). An introduction to complex systems science and its applications. *Complexity*.
- Susha, I., Janssen, M., & Verhulst, S. (2017). *Data collaboratives as a new frontier of cross-sector partnerships in the age of open data: Taxonomy development*. Hawaii International Conference on System Sciences. https://aisel.aisnet.org/hicss-50/eg/open_data_in_government/4
- Verhulst, S., Young, A., Zahuranec, A., Aaronson, S. A., Calderon, A., & Gee, M. (2020). *The emergence of a third wave of open data*. GovLab. <https://apo.org.au/node/311570>
- Wendehorst, C. (2020). *Data governance working group: A framework paper for GPAI's work on data governance*. Global Partnership on Artificial Intelligence. <https://gpai.ai/projects/data-governance/gpai-data-governance-work-framework-paper.pdf>

Interview I

Big Data and the production of public statistics

In this interview, Pedro Luis do Nascimento Silva, the secretary of the Society for the Development of Scientific Research (Sociedade para o Desenvolvimento da Pesquisa Científica [SCIENCE]), discusses the opportunities and challenges faced when adopting sources of Big Data for producing quality statistics, national statistical systems, and articulating data governance networks that involve multiple actors.

Internet Sectoral Overview (I.S.O.)_ What are the possibilities of adopting organic data sources or Big Data for producing quality statistics? Are there any essential precautions that need to be taken when using this type of data?

Pedro Luis do Nascimento Silva (P.S.)_ The generation and availability of the data we are seeing today is unprecedented. At the same time, there is a growing demand for more frequent and more detailed data on living conditions, the environment, etc. The opportunities for taking advantage of organic data sources or Big Data have led those organizations that produce public and official statistics to study, develop, and apply methods and systems for producing quality statistics from this information.

There are two main paths: using one or more new sources to generate statistics of interest directly; or combining data taken from one or more new sources with data taken from traditional sources, such as sample surveys, censuses, and administrative records. In both cases, the objectives may include covering gaps in unexplored topics, replacing statistics that were previously obtained from traditional sources, and expanding statistical production in terms of frequency or level of detail. In this sense, the use of organic data to generate more frequent or more disaggregated (small area) estimates from existing sample surveys is a field of great potential and interest.

In any case, new statistics must satisfy the quality requirements established in United Nations' (UN) Fundamental Principles of Official Statistics,³ so they meet the needs of interested parties and are suitable for use. There are several frameworks and codes of good practice that may guide the production of statistics that are derived from these new sources, with a particular emphasis on the proposal of the United Nations Economic Commission for Europe (UNECE).⁴ Among the main challenges surrounding this issue is the natural tension that exists between satisfying the demands for statistics and guaranteeing their quality.



Photo: Álvaro Vasconcellos

Pedro Luis do Nascimento Silva

Secretary of the Society for the Development of Scientific Research (SCIENCE).

³ Available at: <https://unstats.un.org/unsd/dnss/gp/fundprinciples.aspx>

⁴ Available at: <https://statswiki.unece.org/download/attachments/108102944/Big%20Data%20Quality%20Framework%20-%20final-%20Jan08-2015.pdf>

"(...) producing new statistics or replacing traditional statistics with others based on new sources while ensuring quality requirements are met is a delicate process. It involves applying or developing new methods and systems (...)."

There has been some progress, but there are still few examples of situations in which traditional sources have actually been replaced by new organic data sources. An interesting success case occurred when producing statistics on population mobility during the COVID-19 pandemic.⁵ This is a clear example of usage that could not be easily covered with data coming from traditional sources.

I.S.O._ What are the main challenges faced by official statistical institutes for using Big Data sources?

P.S._ There are three main challenges: Accessing data from many of the new organic sources; training staff to handle data from new sources and combining them with data from traditional sources; and the need to produce statistics that satisfy the demanding quality requirements imposed on “traditional” statistical production. Most of the data from new organic sources are produced and maintained by private organizations that view them as highly valuable assets and, for that reason, are unwilling to share them with third parties – not even with official statistical agencies that serve the public good. There are also issues of comparability over time, a lack of standards when capturing and harmonizing data, and even continuity in obtaining and storing information.

A striking example in the recent history of Brazil was the refusal of telephone companies to provide the Brazilian Institute for Geography and Statistics (IBGE) with information about their landline and mobile telephone customer records. The intention was to enable the IBGE to collect data for its main household sample survey via telephone, which was not possible to do with face-to-face interviews during the first year of the COVID-19 pandemic.

In order to increase and update their statistical production, both the IBGE and other agencies in Brazil and abroad have invested in training their teams so they can absorb, develop, and apply the methods and processes necessary for exploring new sources of organic data. However, the return on these efforts requires time to mature, and it is too early to say whether statistical agencies are ready to take full advantage of the new sources of organic data available to them.

Finally, producing new statistics or replacing traditional statistics with others based on new sources while ensuring quality requirements are met is a delicate process. It involves applying or developing new methods and systems, consulting specialist users, obtaining external validation, and undergoing various testing steps until the statistics are considered suitable for use and publication. In addition to the long time needed, the best tested proposals often fail to deliver data with the required quality.

I.S.O._ How can the Brazilian statistical system incorporate alternative sources of data that are produced and/or collected by other institutions (private, non-governmental, etc.)?

P.S._ An important step for accelerating the use of alternative data sources in Brazil would be to prepare a new legal framework for producing public and official

⁵ Find out more: https://www.gstatic.com/covid19/mobility/2022-10-15_BR_Mobility_Report_pt-BR.pdf

statistics. Brazilian law does not have the fundamental instruments that would allow statistical agencies to access data produced or maintained by private and non-governmental institutions. Even among public institutions there are access limitations. To mention an example that illustrates this point, the Brazilian federal revenue service never allowed the IBGE to access income tax microdata relating to either companies or individuals, not even when it was anonymized.

The legal framework should clearly define the roles, rights, and duties of institutions that seek access to individual data on people, companies, and/or transactions in order to produce official statistics that are in the public domain. One of the obligations should be to protect the confidentiality of individual information, as provided for in current legislation, but it would be necessary to provide for the possibility of using it for the legitimate purpose of producing statistics that are in the public interest.

Another area in which a new legal framework would play a decisive role would be in setting up governance authorities that deal with the production, storage, and use of data for statistical purposes. Some countries have an interesting data archive arrangement, that is, institutions dedicated to the storage, curation, discovery, and dissemination of datasets that are of public interest. Examples include the UK Data Archive,⁶ in the United Kingdom (UK), and the Inter-university Consortium for Political and Social Research (ICPSR),⁷ based in the United States. We still do not have a similar institution in Brazil with the legal mandate and institutional apparatus necessary for promoting activities of this nature.

Finally, but no less important, it would be essential to create and activate an effective coordination body for the national statistical system. Today, this role is delegated to the IBGE, but it does not exercise it for lack of effective instruments. A relevant model is that of the UK Statistics Authority,⁸ which was created in 2007 when the most recent legal framework for the production of official statistics in the UK was established.

I.S.O._ Given the private nature of most Big Data sources, what possible paths are there for linking the institutions that own the data and those that use it, with the aim of managing these sources and exercising governance over them? What aspects should be considered when building networks for data governance involving multiple actors?

P.S._ One possible path is the creation of a “national data archive” for storing, curating, discovering, and disseminating datasets that are of public interest. Such organization could play an intermediary role between the owners of the data and the users and would manage the data of which it is the depository and exercise governance over it. The advantages of this arrangement include guaranteeing the permanence or longevity of the deposited data, and explaining the rules and conditions of access for all those interested in using it. But there are also limitations, such as the likely delay between the moment the data is produced and its availability for access by third parties.

"An important step for accelerating the use of alternative data sources in Brazil would be to prepare a new legal framework for producing public and official statistics."

⁶ Available at: <https://www.data-archive.ac.uk/>

⁷ Available at: <https://www.icpsr.umich.edu/web/pages/>

⁸ Available at: <https://uksa.statisticsauthority.gov.uk/>

"(...) governance mechanisms need to consider models such as the Brazilian Internet Steering Committee (CGI.br), formed by representatives from the various segments involved in the issue (...)."

It is also possible to have “usage contracts” between the institutions that own the data and the statistical agencies interested in using it, in a manner that is similar to the way in which external audit firms work. They require unlimited access to the economic, financial, and accounting data of the companies they are auditing, but undertake to maintain their confidentiality and only use the data for the purpose of providing the services they were hired to carry out.

Following this model, statistical agencies could receive unrestricted access to organic data of interest to a specific statistical operation by undertaking to preserve its confidentiality and use it exclusively for the purposes authorized in the usage contract. This type of arrangement allows direct access to data at the source, without intermediaries or time lag. On the other hand, there is a risk that proprietary institutions will charge amounts that the statistical agencies are unable to pay since they are usually funded by public sources, have limited capacity to raise funds on their own initiative, and need to make their statistics available to the public free of charge.

In either of the cases, we must always seek to preserve: the requirements for protecting the confidentiality of the individual data of people and organizations; the legitimate business interests of the institutions that own the data; and the public interest in the statistics to be produced from such data. Such aspects suggest that governance mechanisms need to consider models such as the Brazilian Internet Steering Committee (CGI.br),⁹ formed by representatives from the various segments involved in the issue, and which serves as a successful Brazilian example.

Article II

The promises and challenges of data-centric digital transformation in the age of Artificial Intelligence

By Moinul Zaber¹⁰

The emergence of Artificial Intelligence (AI) that can harness different types of data for building efficient tools or gathering insights has shown the

⁹ Find out more: <https://cgi.br/about/>

¹⁰ Ph.D. from Carnegie Mellon University, he is computational social scientist focusing on applied machine learning and data science for technology policy. He is currently a Senior Academic fellow at the United Nations University (UNU), and has worked closely with telecommunications regulators and competition authorities in the Global South. As a professor and researcher, he has worked in Bangladesh, Japan, Sweden, Sri Lanka, and Portugal.

potential to change how humans and institutions traditionally make decisions. The use of AI and data in public agencies is enabling more proactive and automated public services. However, since AI as a science is at a nascent stage, the institutional application of various AI tools is challenging. The most significant obstacle comes from the data itself – the raw material of AI. To avoid the risk of failing to use AI, those institutions aspiring to implement data-centric decision-making need to adapt to the new ways of digital transformation. This article sheds light on various aspects of data-centric digital transformation to enable AI-focused automation by public agencies.

The promises of data and AI

A public agency is any self-governing entity (such as a department, a commission, or an authority) established by either local or national government. Its primary objective is to perform the necessary duties mandated by government and those these organizations serve. These duties may include: Safeguarding national security; providing civil protection; regulating markets; and ensuring access to necessities like food, shelter, energy, communication, and environmental protection. Most of these services, however, have two distinct features. On the demand side there is decision-making at various levels of service delivery, while on the supply side there is engaging with service receivers. As with any decision-making process, these processes require a vast amount of data, sometimes from internal sources, but many times from various external sources.

Data-centric digital transformation helps automate the data movement process. The target is to achieve citizen-centric decision making and service delivery. It is an ongoing effort that involves the integration of digital technologies and data to improve business processes, create new business models, and deliver better services to customers. For example, public services that focus on social security provide individuals with a certain degree of income security when faced with contingencies such as old age, survivorship, incapacity, disability, unemployment, or rearing children. It may also offer access to curative or preventive medical care. Throughout the world, digital transformation is enabling the implementation of increasingly comprehensive public service systems.

The emerging adoption by public institutions of AI tools that have various forms of data as their raw material is enabling more proactive and automated public services. The use of data and AI can help improve the efficiency, effectiveness, and responsiveness of public service. By using these technologies, government agencies can better meet the needs of citizens and provide more value to the communities. Many agencies have been using data analysis to identify areas where public service is lacking or needs improvement. The advent of AI fields – such as machine learning, pattern recognition, natural language processing, computer vision, and data visual-



Photo: Cristina Braga

Moinul Zaber
United Nations
University
(UNU).

Predictive analytics can be used to anticipate future needs and trends and enable government agencies to plan and allocate resources more effectively.

ization – is shaping the way data in their various forms can be used to make public service more effective and user-centric.

DATA AS THE RAW MATERIAL OF DIGITAL TRANSFORMATION IN THE AGE OF AI

By gathering and analyzing data on the supply side, such as response rates, customer satisfaction, and waiting times, it is possible to improve the level of engagement in relation to service recipients. AI-powered chatbots and virtual assistants can provide round-the-clock assistance to citizens seeking information or help with government services. On the demand side, government agencies can identify areas that require the allocation of more resources or the implementation of new policies. By analyzing demographic and socioeconomic data, among other, AI can help public officials identify disparities in service delivery and take action to address them.

When machine learning algorithms are used to identify patterns in data, they can help government agencies detect fraud, waste, and abuse, thus saving taxpayers money and improving the overall efficiency of the services on offer. Predictive analytics can be used to anticipate future needs and trends and enable government agencies to plan and allocate resources more effectively.

Using predictive analytics, for example, governments anticipate spikes in demand for emergency services during certain times of the year. Data and AI can be used to identify potential risks and threats in real time, thereby helping law enforcement agencies prevent crime and improve public safety. Data visualization tools can present complex data in such a way that is easy to understand and interpret. This can help government agencies communicate with citizens more effectively and make data-driven decisions.

It is important, however, to ensure that these technologies are implemented responsibly, with appropriate safeguards to protect privacy and prevent bias. Since AI relies heavily on data to train models and make predictions, organizations need to ensure that their data is of high quality, reliable, and accessible. There are also issues related to ethics and legality, which encompass personal privacy, transparency, and fairness, and to national security issues. AI and data-based interventions also need to be compatible with the legacy systems and practices of the agencies to ensure effectiveness. Organizations will need to invest in human capital to reduce the skills gap. As legacy automation is being replaced by AI-based data-centric automation, skills in data science, machine learning, and AI development are required.

Computing traditionally focuses on code, while AI focuses on data. The performance of an AI-based system depends on a continuous supply of good-quality data. The accuracy and reliability of the results generated by the algorithms are heavily related to the quality of the data used to train them. For example, machine learning algorithms rely on patterns and relationships in the data to make predictions or classify data. If the data is inaccurate, incomplete, or contains errors, the algorithm may learn incorrect patterns or make inaccurate predictions, leading to poor performance. Moreover, the absence of quality data generates biased patterns that can result in discrim-

inatory or unfair predictions. It is important to note that these systems are designed to be continuous learners. Hence there should be a permanent supply of good-quality data, otherwise the output of the algorithms will not be commensurate with the context.

The facts: Brief introduction to AI and its branches

The algorithmic processes of AI are unlike traditional processes and may be more efficient for many tasks. In the case of AI, instead of writing a program for each specific task, many examples are collected that specify the correct (or incorrect) output of a given input. AI algorithms then take these examples to produce a program that does the job and that is scalable for new cases. Programs adapt to the changes in data as the essence of AI programs is to re-train themselves using the new data. Massive amounts of computational power are now available for these tasks, which is why its use is cheaper than writing a task-specific program.

This capability of scalability and harnessing insights from data has made AI an essential and complementary tool for policymakers and service providers aiming at the social good. Various AI tools are being used for: Responding to a crisis; promoting economic empowerment; alleviating educational challenges; mitigating environmental challenges; ensuring equality and inclusion; promoting health; reducing hunger; information verification and validation; infrastructure management; public and social sector management; and even for security and justice.

MACHINE LEARNING AND THE NEED FOR BETTER-QUALITY DATA

AI is a broad field that encompasses many branches, each focusing on different aspects of intelligence and cognitive processing. A few of the main branches of AI include machine learning, natural language processing, computer vision, robotics, and expert systems. Among these, machine learning deals with the process of learning, reasoning, pattern finding, and decision-making. It is an umbrella of methods that help build practical tools for other branches of AI.

A learning problem can be defined as the problem of improving a measure of performance when executing tasks, using some type of training experience. For example, in learning to detect pension eligibility, the task is to determine “eligible” or “not eligible” for any given resident’s application. The performance metric may be to measure the accuracy of this eligibility classifier. The algorithm may be trained from a dataset containing historical eligibility information of applications, each of which is labeled in retrospect as being eligible or not. There may be many other alternate accuracy measures and training sets mixed with labeled and unlabeled data. Machine learning can be broadly categorized into three main branches: Supervised learning, unsupervised learning, and reinforcement learning.

This capability of scalability and harnessing insights from data has made AI an essential and complementary tool for policymakers and service providers aiming at the social good.

Machine learning can be broadly categorized into three main branches: Supervised learning, unsupervised learning, and reinforcement learning.

For many applications, it can be far easier to train a system by showing it examples of desired input-output behavior than to program it manually by anticipating the desired response for all possible inputs. Supervised learning is a type of machine learning in which the algorithm is trained on labeled data. Input here is paired with the desired output. The algorithm learns to map the input data to the output data, allowing it to make predictions on new, unseen data instance. Predicting whether an email is spam or not is an example of classification, while predicting the price of a house based on its features is an example of regression – two types of supervised learning algorithms.

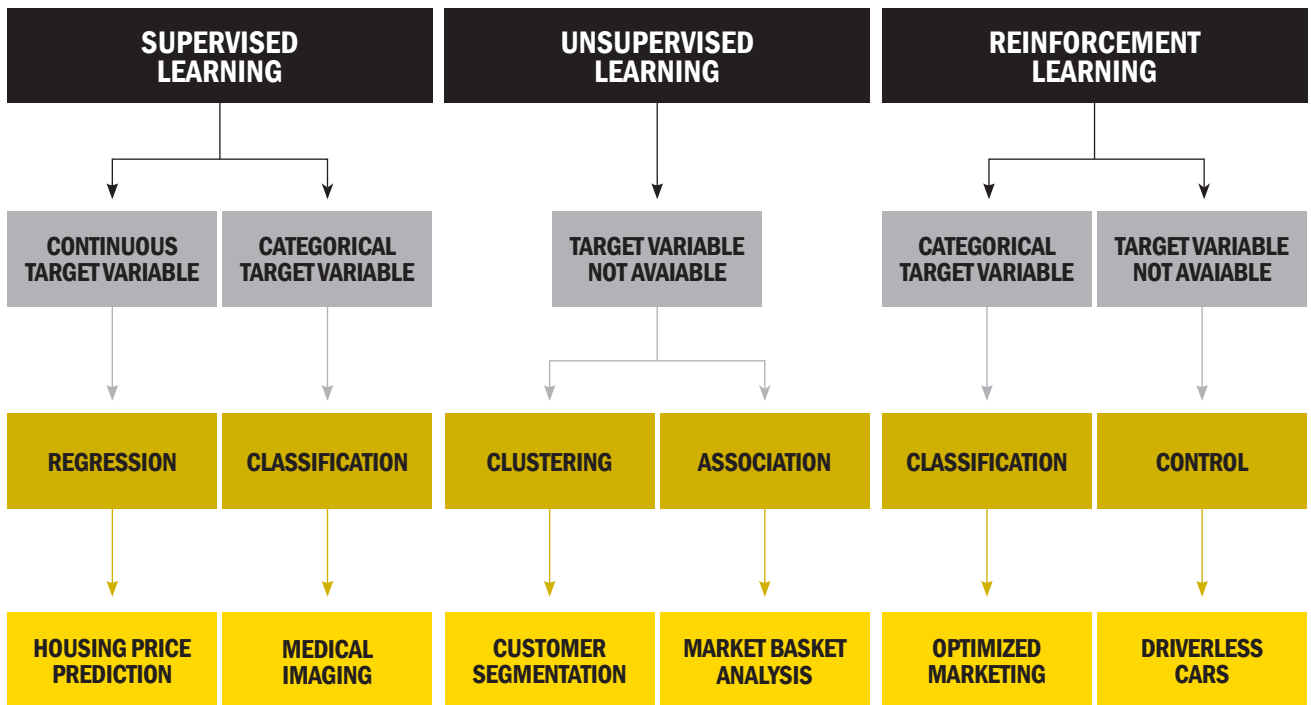
In unsupervised learning the algorithm is trained on unlabeled data; input data is not paired with any output variable. The algorithm learns to identify patterns and structure in the data without any prior knowledge of what the output should be. Clustering tasks – such as grouping customers into segments based on their purchase behavior – and dimensionality reduction tasks – such as reducing the number of variables in a dataset – are examples of unsupervised learning.

Reinforcement learning is a type of machine learning in which an agent learns to interact with an environment, learns by trial and error, and performs actions that maximize a reward signal. In trying to teach a pet a certain task, for example, we may give them a treat (reward) if it performs the task correctly. If it does not, then we may say “NO” indicating a penalty. Over time, the pet gets to associate the correct behavior with the reward, and it gets better at performing the trick. Reinforcement learning can be used in chatbots or conversational agents to improve their performance in understanding or responding to user queries leading to better user experience and increased satisfaction.

One of the high-impact areas of progress in supervised learning involves deep networks. Deep learning systems make use of gradient-based optimization algorithms to adjust parameters throughout multilayered networks based on errors at their output. Deep networks are based on an artificial neural network algorithm that is modeled after the structure and function of the human brain. Deep learning allows machines to learn from vast amounts of data and to recognize patterns that might be difficult or impossible for humans to identify.

Machine learning in general is heavily reliant on the availability of training data. The amount of labeled data required to train a machine learning model varies depending on the dataset and model adopted. The requirement increases with the complexity of the datasets and the “depth” of the models. Deep networks allow far more generalization than shallow neural networks and traditional machine learning approaches and therefore they achieve significantly better accuracy. When applying deep learning to a problem, the key challenge is the large amount of data required to train the models.

Figure 1 – A BROAD CLASSIFICATION OF DIFFERENT MACHINE LEARNING ALGORITHMS



Source: Prepared by the author, based on LITSLINK (2019).

LEARNING WITH LIMITED DATA

When the amount of data is limited, models are often provided with biased data, and sometimes the designed algorithms adjust so well to the training dataset that they fail to give the right solution in the real world. In many cases it is difficult to accumulate large amounts of data, or perhaps it does not even exist (Defence Science and Technnnology Laboratory [DSTL], 2020). For example, an institution wants to create a model to predict certain traits of the users, but they only have the historical data of 50 users. Traditional machine learning approaches to this model may have a biased output that goes against gender, age, and race due to a lack of variation. Machine learning problems associated with small datasets require a different set of techniques and approaches compared to those involving large data. These include: Feature engineering that helps create new features; regularization to avoid over-combining outputs from multiple models; and learning transfer using a model that has been pre-trained on a larger dataset. Table 1 shows different types of machine learning methods based on the amount of data available.

Table 1 – MACHINE LEARNING METHODS BASED ON THE AMOUNT OF DATA

DATA AMOUNT	LABELED OR UNLABELED	LEARNING METHOD USED	COMMENT
Small amount	Mostly unlabeled	Zero-shot learning	Uses description of concept to train the model, concept ontology, semantic word embedding
Small amount	Mostly unlabeled	Manual labeling	Manually labels the data
Small amount	Mostly labeled	Shallow machine learning, meta-learning, knowledge reasoning	Train a metamodel to be applied for unseen tasks, support vector machines, decision trees, multi-layer perceptrons, ontological approach leveraging a description of objects
Large amount	Mostly labeled	Deep learning	Convolutional neural network
Large amount	Mostly unlabeled	Active learning, semi-supervised learning, self-supervised learning, unsupervised learning	Utilizing both labeled and unlabeled data, query to select examples for a human operator to label, clustering, anomaly detection, latent variable, autonomous labelling

Source: Prepared by the author.

The challenges of the inefficient use of data for AI

AI tools imitate the way humans think and act. This means that algorithms can be inaccurate on many occasions. Such inaccuracies may cause risks to personal privacy, national security, fairness, transparency, and accountability. Inaccuracy may also engender data, algorithms, and human interaction with the design process.

If trained on biased datasets, AI systems can perpetuate and even amplify existing biases in data. If the training data is unrepresentative or lacks diversity, the AI system will learn to make biased predictions. A massive amount

of data is fed into the machine to recognize certain patterns. Unstructured data from the web, social media, mobile devices, sensors, and smart devices (i.e., the Internet of Things) make data absorption, linking, sorting, and manipulation difficult. In the absence of careful data curation, the dataset may be fraught with incomplete or missing data, as well as inaccurate or biased data.

In broad terms, there are four types of bias: Sample bias; measurement bias; algorithmic bias; and bias against groups or classes of objects and people. However, algorithmic bias seems to be the least discussed. Some algorithms are systematically biased toward a specific type of data. Several reasons make bias correction difficult. Firstly, the introduction of bias is not always obvious during a model's construction. Secondly, it is hard to retroactively identify where the bias originated from. Thirdly, machine learning, one of the AI fields largely used for data analytics, needs to train, test, and validate its algorithm with the dataset. To do so, in many cases data is divided randomly for training, testing, and validation which may keep the same biases. Fourthly, a lack of context due to the failure to understand the targeted users can create bias. A system designed in country A cannot be applied in country B as different communities have different ways of facing public policy problems. Lastly, context is not only affected by the communities, but it is also defined by the institutions. For example, "fairness" in the case of the "unemployment problem" may differ from "criminal justice."

Personal data may be removed from one dataset while another dataset may have it that the AI system may reveal. There is a risk that this may cause an inadvertent revelation of sensitive data unless care is taken to remove personal data from all datasets. Moreover, bias depends largely on the developers who curate the data or design the algorithms, and decide how these will be deployed and, ultimately, how they are used. A lot depends on how a problem is framed. While framing a problem, scientists decide what in fact they want to achieve when they create a learning model. Even the composition of the engineering teams can be biased. Problem framing depends on who designs it, who decides how it is deployed, what the acceptable level of accuracy is, and if the applications of AI are ethical. Failure to address these issues has proliferated the number of algorithms that dictate what political advertisements people see, how recruiters filter job seekers, and even how security agents are deployed in neighborhoods.

Finally, the interaction between humans and machine needs to be evaluated. If operators of AI tools do not recognize when systems need to be overruled, accidents and injuries are possible. For example, an Air France flight over the Atlantic Ocean in June 2009, crashed due in part to the over-reliance of the cockpit crew on the autopilot (the speed sensor confused the pilots). Human judgment can be faulty when it overrides systems. A lapse in data management, scripting error, and misjudgment in model training, can compromise fairness, privacy, security, and compliance. Data collectors may unintentionally induce bias if they access the data of people of a certain demographic over others.

Finally, the interaction between humans and machine needs to be evaluated. If operators of AI tools do not recognize when systems need to be overruled, accidents and injuries are possible.

EXPLAINABILITY AND DATA

Machine learning algorithms learn patterns and relationships from vast amounts of data, often without explicit programming of the rules or decision-making criteria. As a consequence, the algorithms can produce results that are accurate but may not be intuitive or easily understood by humans. Much of AI (particularly deep learning) is plagued by the “black box problem.” These models can be highly complex, with many layers and interconnected nodes. We often know the inputs and outputs of the model, but we do not know what happens in between. To ensure trust and accountability it is imperative to ascertain how an intelligent machine suggests certain decisions. Moreover, if AI systems become explainable, they may be able to significantly increase the profits of organizations, increase the accuracy of models by 15% to 30% and reduce monitoring efforts by up to 50%.¹¹

There are several reasons for the lack of explainability of machine learning algorithms. The foremost is related to data and their use. Machine learning algorithms can perpetuate biases present in the data resulting in outputs that reinforce existing societal biases. As a consequence, these models are also prone to adversarial attacks and biases. More significantly, and due to the black box nature of the algorithms, it is hard to pinpoint the features that cause such biases. Machine learning algorithms often operate in high-dimensional spaces. This results in non-linear relationships between features and output predictions making it difficult to explain.

Making machine learning models explainable is an active research field. Some of the notable work done to understand the feature-output relationship are: SHapley Additive exPlanations (SHAP); Local Interpretable Model-Agnostic Explanations (LIME); and Gradient-weighted Class Activation Mapping (Grad-CAM). One of the popular counters to the black box problem is Explainable AI (XAI) – a set of machine learning processes that allows human users to comprehend, trust and manage AI. The goal of XAI is to enable interactions between people and AI systems by providing information about how decisions and events come about (Tjoa & Guan, 2021). This has been so widely embraced that is mentioned by the General Data Protection Regulation (GDPR), of the European Union, and since 2016 the Defense Advanced Research Projects Agency (DARPA), of United States government, has made it its research focus.

Due to the data-driven nature of AI algorithms – as opposed to the program-driven nature of traditional algorithms –, traditional mechanisms for auditing software systems are quite inadequate. AI systems base their output on millions of data points. Changes in training samples can also induce

¹¹ Find out more: <https://www3.technologyevaluation.com/research/brochure/new-technology-the-projected-total-economic-impact-of-explainable-ai-and-model-monitoring-in-ibm-cloud-pak-for-data.html>

different learning. Hence, there is no expected result with many AI algorithms. Systems learn what the best prediction is, which makes it difficult to validate. Such unpredictability is a challenge. This means that auditing datasets and the output is not sufficient for evaluating AI tools.

CHALLENGES OF THE PROPER GOVERNANCE OF DATA

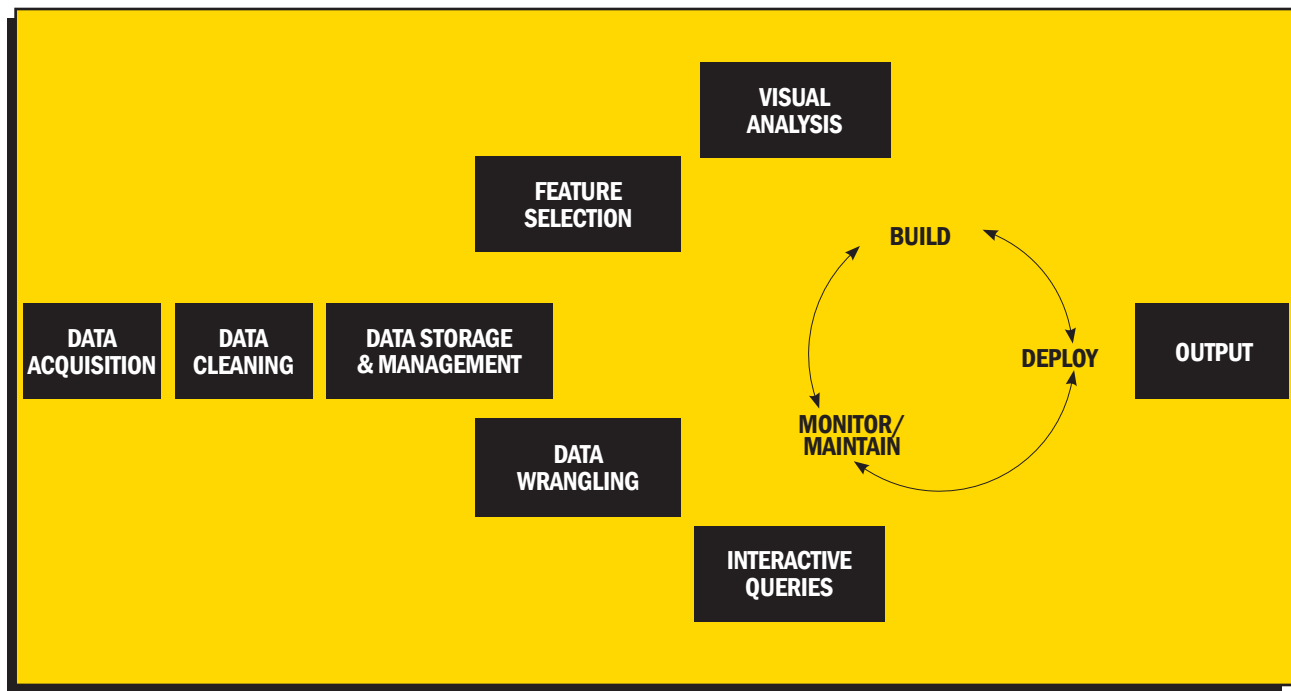
Traditionally public institutions work as silos that create barriers to data access and availability for other institutions. Lack of data access results in poor data analytics and AI tools. Implementing data-centric AI may require significant changes to existing IT systems and processes. This means that organizations need to ensure that AI solutions are integrated seamlessly with their systems and workflows without causing disruption.

The accuracy and reliability of AI models depend on the quality of the data used to train them. To achieve the most out of data, it is imperative to employ a data governance model that manages and ensures the quality, accuracy, completeness, and security of the data used to train and develop AI algorithms (Abraham et al., 2019). In aiming to convert data into information, data must go through a pipeline that consists of a series of steps, and the results of one step may influence the next. There is a specific order that may not be linear, as data processing may be an iterative process.

The steps start with data collection, in which raw data is gathered, and with preprocessing, when data is cleaned and transformed to ensure quality. Data is then stored in various forms in data warehouses, from where it moves to the analysis phase, in which various patterns are identified. It is then modeled, when various mathematical models are used to detect anomalies or predict outcomes, and finally moves on to visualization, during which the insights are visually summarized.

(...) it is imperative to employ a data governance model that manages and ensures the quality, accuracy, completeness, and security of the data used to train and develop AI algorithms.

Figure 2 – REPRESENTATION OF A SIMPLE MACHINE LEARNING PROCESS PIPELINE, IN WHICH DATA ARE TRANSFORMED INTO OUTPUT



Source: Prepared by the authors.

To ensure the proper management of these steps, an appropriate data governance model is required, which involves defining policies and procedures for each of the above-mentioned steps. This includes identifying the sources of data, establishing data quality standards, defining data ownership and stewardship, as well as ensuring compliance with relevant regulations and industry standards. It is important to ensure that the data used to train AI models is consistent, accurate, and relevant.

Data governance for AI also involves establishing processes for data preparation and pre-processing, including data cleaning, normalization, and feature engineering. At the policy level data governance also addresses ethical and privacy concerns.

INSTITUTIONAL CHALLENGES OF IMPLEMENTING DATA-CENTRIC AI

The institutional challenges of implementing data-centric AI in public institutions can be divided into three categories: Legal, regulatory, and the availability of human resources. The legal challenges arise from the inherent issues of data and machine learning modeling such as ethics, data privacy, bias and discrimi-

nation, transparency and explainability, accountability, and intellectual property. Many countries have not yet succeeded in reformulating their policies on these issues, and it is becoming extremely hard for legal institutions to keep up with the pace of the rapid technological transformation that is taking place.

Several prominent regulations, such as the General Data Protection Regulation (GDPR), the California Consumer Privacy Act (CCPA), or the Asia-Pacific Economic Cooperation (APEC) privacy framework are now in place. Countries such as Brazil, India, Australia, and Canada have also introduced their own data protection acts. Besides data protection and personal privacy, several regulatory initiatives revolve around the ethical considerations of the use of AI; compliance with regulations, such as the Health Insurance Portability and Accountability Act (HIPAA) in the healthcare sector, or the Family Educational Rights and Privacy Act (FERPA) for educational institutions in the United States; issues related to who bears the liability for the decisions taken by AI models; and procurement regulations, like the Federal Acquisition Regulation (FAR) in the United States. Public institutions need to address these challenges so they can leverage data to make better decisions, improve services, and better serve their stakeholders, ensuring compliance with regulations while protecting privacy and security.

Apart from the legal and regulatory challenges at the policy level, there are several human resource challenges at the implementation level. Public agencies follow processes that are slow in comparison with private organizations. This means they are slow to respond to technological change. For civil servants to design the public policies they adopt, they must understand AI and data and be able to tap into their potential. They need to know the opportunities offered by AI while being fully aware of its risks and challenges. Working in AI and data requires specialized skills, such as data scientists, machine learning engineers, software developers, and engineering policy specialists. Organizations need to plan on hiring and retaining these skilled professionals. They also need to train and upskill the existing workforce. Implementing AI may result in the substitution of some jobs that can be automated. Organizations need to develop a plan for reskilling and redeploying employees whose jobs are affected by data-centric AI. Since implementing AI can involve significant changes to existing processes, workflows, and job roles, organizations need to develop a change management plan to help employees navigate these changes effectively.

For digital transformation based on AI and data to succeed, governments need to change the way they function. This is difficult. However, it can start if public employees at different levels acquire the competencies needed to understand the transformations that data-centric AI brings. It is therefore important to raise awareness of the required skills at different levels. This can be done by understanding the capacity-building needs at the individual, team, department, and government level. Increased collaboration and communication between institutions would help these departments share insights. Most significantly there should be continuous monitoring of the impact of capacity-building initiatives.

Since implementing AI can involve significant changes to existing processes, workflows, and job roles, organizations need to develop a change management plan to help employees navigate these changes effectively.

(...) the transparency and “explainability” of the AI application constitute an important issue, especially with regard to decisions that impact people.

Conclusions

Data is the most significant ingredient of progress in this age of Artificial Intelligence. AI is gradually becoming a key technology for public service organizations as it increases administrative efficiency. Its ability to take advantage of the enormous number of various types of data and find insights that help automate processes help with decision-making.

Although positive developments can be observed, however, there are various challenges. AI is data hungry. Therefore, a continuous stream of quality data must be ensured so AI tools do not make biased or incorrect decisions. Among the critical factors, data availability and quality are the most prominent need for training AI systems appropriately. Such “data needs” require establishing a data governance strategy to use internal data as well as potential data from other organizations and involves assessing compliance with data protection regulations.

While there are several AI algorithms that work well with small or big datasets, AI is a nascent science. These solutions should be scrutinized before mandating real-life use, especially in relation to the limitations and risks of AI, as well as the trade-off between process automation versus human control. The methodological differences between AI and traditional software development pose challenges to institutions when it comes to carrying out their projects. Particularly, the transparency and “explainability” of the AI application constitute an important issue, especially with regard to decisions that impact people.

Data-centric digital transformation can happen if institutions are prepared for the changes data and AI bring. Authorities contemplating such a transformation have to consider the legal, regulatory, and institutional challenges. Governments should assess the competencies of their civil servants and emphasize capacity building to ensure a smooth transition towards data-centric digital transformation.

References

- Abraham, R., Brocke, J., & Schneider, J. (2019). Data governance: A conceptual framework, structured review, and research agenda. *International Journal of Information Management*, 49(S1).
- Amodei, D., & Hernandez, D. (2018). *AI and compute*. OpenAI. <https://openai.com/research/ai-and-compute>
- Bienvenido-Huertas, D., Farinha, F., Oliveira, M., Silva, E., & Lança, R. (2020). Comparison of Artificial Intelligence algorithms to estimate sustainability indicators. *Sustainable Cities and Societies*, 63.
- Defence Science and Technology Laboratory . (2020). Machine Learning with Limited Data. *Future of AI for Defense project*. <https://www.gov.uk/government/publications/machine-learning-with-limited-data>
- Krizhevsky, A., Sutskever, I., & Hinton, G. (2012). ImageNet classification with deep convolutional neural networks. *Advances in Neural Information Processing Systems*, 1097-1105.
- LITSLINK. (2019, September 19). *An introduction to Machine Learning Algorithms*. <https://litslink.com/blog/an-introduction-to-machine-learning-algorithms>
- Minsky, M., & Papert, S. (1969). *Perceptrons: An introduction to computational geometry*. MIT Press.
- Stanford Institute for Human-Centered Artificial Intelligence. (2022). *The AI index report: Measuring trends in Artificial Intelligence*. <https://aiindex.stanford.edu/report/>
- Thompson, N. C., Greenewald, K., Lee, K., & Manso, G. F. (2020). The computational limits of deep learning. *MIT Initiative on the Digital Economy Research Brief*, 4. <https://ide.mit.edu/wp-content/uploads/2020/09/RBN.Thompson.pdf>
- Tjoa, E., & Guan, C. (2021). A survey on Explainable Artificial Intelligence (XAI): Toward medical XAI. *IEEE Transactions on Neural Networks and Learning Systems*, 32(11).

Interview II

Photo: Personal archive



Anita Lübbe

Judge of the
Regional Labor
Court of the 4th
Region (TRT-4).

Document management and memory in the Judicial Branch

In this interview, Anita Lübbe – judge of the Regional Labor Court of the 4th Region (TRT-4), president of the Permanent Brazilian Forum in Defense of the Memory of the Labor Court (Memojutra), and coordinator of the Digital Preservation Subcommittee of the National Program Committee for Document Management and Memory of the Judicial Branch (Proname) – discusses her experience in document and memory management in the Judicial Branch, as well as the policies implemented for preservation and access to their collection of documents.

Internet Sectoral Overview (I.S.O.)_ Considering the digitization of Judicial Branch processes and collections, and the storage possibilities facilitated by digital technologies, what is the importance of establishing policies for data governance and management of the judicial information ecosystem?

Anita Lübbe (A.L.)_ The importance lies precisely in the fact that this “storage” will only become an effective preservation activity when it is the result of a consistent document and memory management policy, in accordance with applicable legislation and technical guidelines. In short, data management encompasses actions that include: The generation of administrative and legal documents; classification, by applying the Brazilian Unified Procedural Tables to the processes; application of the Brazilian Classification Plan and Temporality Table of Documents of the Administration of the Judicial Branch (PCTTDA)¹²; and the safe collection of documents that are selected for permanent safe-keeping.

Article 16 of Resolution CNJ No. 324/2020¹³ (Brazilian Council of Justice [CNJ]) indicates the classification of Judicial Branch documents as being either current, intermediate, or permanent, combined with the consequent definition of periods of custody according to the respective Temporality Table. If they are discarded, the Representative Statistical Sample Plan also needs to be considered.¹⁴

In Brazil, we have a set of laws and regulations that need to be observed to establish and maintain policies for data governance and the preservation of documents and memory, starting with: Article 216 of the Federal Constitution; the Archives Law (Law No. 8159/1991); the Access to Information Law

¹² Available at: <https://www.cnj.jus.br/programas-e-acoes/gestao-documental-e-memoria-proname/gestao-documental/tabelas-de-temporalidade-da-area-administrativa/>

¹³ Available at: <https://atos.cnj.jus.br/atos/detalhar/3376>

¹⁴ Find out more: <https://www.cnj.jus.br/programas-e-acoes/gestao-documental-e-memoria-proname/gestao-documental/instrumentos-do-proname/>

(LAI – Law No. 12527/2011); the General Data Protection Law (LGPD – Law No. 13709/2018); and resolutions of the National Archives Council (Conarq) and the CNJ, particularly aforementioned Resolution No. 324/2020, which established parameters, definitions, and objective rules for document and memory management for the Judicial Branch. I draw attention to the Document Management and Memory Management Manuals of the Judiciary,¹⁵ published by the CNJ in February 2021, to enable the provisions of Resolution CNJ No. 324/2020 to be implemented and managed.

According to the above legislation and regulations, document and memory management policies must be instituted by all courts, thereby making the State's duty to safeguard documents and the citizen's right to access information. To this end it is fundamental that documents are correctly classified, Temporality Tables and sampling plans are duly applied, and that the Reliable Digital Archive Repositories (RDC-Arq) guidelines are implemented.¹⁶ Having an adequate document and memory management policy, various information ecosystems within the Judiciary can be identified and created.

I.S.O._ What are the main aspects to consider when setting up document and memory management systems and when properly storing legal data?

A.L._ I would highlight particularly forming teams that have qualified professionals so that right from the start of their action plans they can diagnose the extent of the collection to be considered, interpret this diagnosis, and select, classify, and properly archive the data and prepare a sample plan if documents are to be disposed of. The participation of professionals from the areas of archive studies, history, library science, museum studies, and information science is essential, as this strengthens a culture of preservation in its various aspects in each institution. It is extremely important to have these professionals included in the permanent staff of the courts.

With regard to the technical requirements of the information systems, it should be noted that the CNJ is in the final stages of reviewing the Requirements Model for Computerized Systems for the Management of Processes and Documents of the Judicial Branch (MoReq-Jus),¹⁷ which lists in detail all the mandatory or desirable requirements to be considered. It is also important and necessary to set up the RDC-Arq, repositories where documents are sent and archived, thereby guaranteeing their preservation, the custodial chain, and adequate access. In March 2023, the CNJ launched the Judiciary Document Digitization Manual,¹⁸ which should be used in conjunction with the document management and memory management manuals.

What here should be recognized is the fact that the Court of Justice of the Federal District and Territories (TJDFT) was the first in the country to start

¹⁵ Available at: <https://www.gov.br/arquivonacional/pt-br/cnj-lanca-manuais-para-gestao-de-documentos-e-da-memoria-do-judiciario>

¹⁶ Find out more at: https://www.gov.br/conarq/pt-br/centrais-de-conteudo/publicacoes/conarq_diretrizes_rdc_arq_resolucao_43.pdf

¹⁷ Find out more: <https://www.cnj.jus.br/programas-e-acoas/gestao-documental-e-memoria-proname/gestao-documental/moreq-jus-e-sistemas-informatizados/>

implementing an RDC-Arq in 2018, in partnership with the Brazilian Institute of Information in Science and Technology (IBICT), thus creating a paradigm to be followed. From the outset, the TJDFT sent a clear message that its solutions were and are available to other courts. But the Electronic Judicial Process (PJe) that is used by the five branches of the Judiciary – State Justice, Labor Justice, Federal Justice, Electoral Justice, and Military Justice – is not exactly the same, and each branch has specific matters and competences, which makes it very difficult to define a single RDC-Arq to suit everyone. In mid-2022, the Superior Council of the Labor Court (CSJT) and TRT-4 started their project to introduce the RDC-Arq, also in partnership with IBICT. The goal is to apply and install the RDC-Arq in all 24 regional labor courts; TRT-4 in Rio Grande do Sul state is the pilot court for this. In the case of the Labor Court, the initiative is strengthened by aiming to unify the RDC-Arq applied in these courts, since it is based on the same PJe model and branch of the Judiciary. The project is expected to take five years to implement, the first two being dedicated to creating the necessary solutions and another three years for monitoring it being set up in the courts.

I.S.O._ Could you tell us a little bit about your experiences in initiatives to articulate in different institutions and actors for managing data and memory in the Judiciary? What were the main lessons learned from these initiatives?

A.L._ The Memory preservation process began with the Written Records of TRT-4,¹⁸ which were created by TRT-4 Administrative Resolution No. 22/2003, which was chaired at the time by Justice Rosa Weber, the current President of the Federal Supreme Court (SFT) and CNJ. I subsequently joined the Written Records Coordinating Committee. It is important to register that already in the mid-1990s, memory preservation movements began in various courts of all branches of the Judiciary, with the establishment of spaces and initial projects for preserving documents.

In 2006, the Permanent National Forum in Defense of the Memory of the Labor Court (Memojutra) was created, which I have been a member of since the beginning and that I currently chair. It is made up of judges and civil servants from the Superior Labor Court (TST), the Superior Council of Labor Court (CSJT), and the 24 regional labor courts, with the permanent participation of all the coordinators of the respective Memory Centers.

In Memojutra, I remember pointing out how important it was to select and present the collections of the regional labor courts within the scope of the Memory of the World Program (MoW), of the United Nations Educational, Scientific and Cultural Organization (UNESCO). The initiative was based on the example of the Regional Labor Court of the 6th Region (TRT-6) in Pernambuco, the first Labor Court to receive the MoW seal as a documentary heritage of humanity in 2012.

¹⁸ Available at: <https://www.cnj.jus.br/wp-content/uploads/2023/03/proname-manual-digitalizacao-15-03-2023.pdf>

¹⁹ Available at: <https://www.trt4.jus.br/portais/memorial>

In subsequent years, several institutions in Brazil obtained this recognition from UNESCO, among them: TRT-4 in 2014; the Regional Labor Court of the 3rd Region (TRT-3) in Minas Gerais, in 2015; and the Regional Labor Appeals Court (TST) in 2016. Also, with regard to important actions for preserving memory, I would mention the now-extinct Sectoral Chamber on Judicial Archives (CSAJ), of Conarq, of which I was briefly a member in 2019.

I am currently a member of the Committee of the National Program for Document Management and Memory of the Judicial Branch (Proname), which I have participated in since 2019. In 2022 I started coordinating the Digital Preservation Sub-committee. In the Committee, we have been working to update and create regulations aimed at document and memory management, such as: Resolution No. 316/2020, which established the Day of Memory of the Judiciary and the annual National Memory Encounter of the Judicial Branch (ENAM); Resolution CNJ No. 429/2021, which established the CNJ Memory of the Judiciary Award; and Resolution No. 324/2020 and the Document Management and Memory Management Manuals of the Judiciary that I have already mentioned.

I think that the main lesson we have learned has been the dialogue established between all the members of these initiatives: Judges, civil servants, professionals in the areas of history, archive studies, museum studies, library science, information technology, information science, and others. With their expertise they teach us and make it possible to improve document and memory management, thus guaranteeing access and information security.

I.S.O._ In your opinion, what is the maturity level is the Brazilian Judiciary with regard to memory management and making its collections available to a wider audience?

A.L._ The Judiciary has carried out several actions for preserving and ensuring access to its collections in recent decades. We have gradually moved away from the physical process to the electronic process, with the implementation of the Electronic Judicial Process System based on Resolution CNJ No. 185/2013. Although we already have the electronic process in all branches of the Judiciary, the electronic proceeding is not fully implemented. We still have legacy physical processes that need to undergo classification and definition of their safekeeping and digitalization periods, so that they can then be preserved in reliable digital repositories. Today, throughout the Brazilian Judiciary, there is a significant number of digitally-born processes (that is, they were created digitally) and others in physical format (created on paper). Of the latter, a part is already digitalized, while another significant part is still in the digitalization phase, defining its temporality and, depending on the case, eliminating it, with the consequent formation of a statistical sampling plan.

In addition to the many initiatives to preserve its memory, the CNJ has held National Meetings on the Judicial Branch's Memory, as regulated in Ordinance CNJ No. 80/2022. The first meeting in May 2021 was held virtually due to the COVID-19 pandemic. The following year, in May 2022, the II ENAM took place in person at the Pernambuco Court of Justice (TJPE) in Recife,

with a significant number of participants including judges, civil servants, professionals from different areas – such as history, archive studies, library science, museum studies, information technology, and information science – and students. Also face-to-face, the third edition (III ENAM) will take place from May 10 to 12, 2023²⁰ in Porto Alegre (RS), hosted by the five courts located there: The Rio Grande do Sul Court of Justice (TJ-RS), TRT-4, Federal Regional Court of the 4th Region (TRF-4), Rio Grande do Sul Regional Electoral Court (TRE-RS), and Rio Grande do Sul Military Court of Justice (TJM-RS). Under the theme “Structuring memory,” the closing ceremony on May 12 will be attended by Justice Rosa Weber. The objective is to provide all courts with tools and alternatives to assist in their implementation or expansion of activities to preserve the memory of the national Judiciary, as well as improving the document and memory management of each institution. One of the hallmarks of the III ENAM is the online pre-meeting on April 13 and 14, 2023, the aim being to update those registered to attend on recent legislative concepts and changes (including CNJ regulations), thus motivating participants for the face-to-face lectures of the event in May. The theme of memory and document management in the Judiciary subject is, without a doubt, an inspiring path to follow.

Domain Report

Domain registration dynamics in Brazil and around the world

The Regional Center for Studies on the Development of the Information Society (Cetic.br), department of the Brazilian Network Information Center (NIC.br), carries out monthly monitoring of the number of country code top-level domains (ccTLD) registered in countries that are part of the Organisation for Economic Co-operation and Development (OECD) and the G20.²¹ Considering members from both blocs, the 20 nations with highest activity sum more than 89.80 million registrations. In March 2023, domains registered under .de (Germany) reached 17.49 million, followed by China (.cn), the United Kingdom (.uk) and Netherlands (.nl), with 9.65 million, 7.19 million and 6.29 million registrations, respectively. Brazil had 5.09 million registrations under .br, occupying 5th place on the list, as shown in Table 1.²²

²⁰ Find out more: <https://sites.google.com/trt4.jus.br/enam>

²¹ Group composed by the 19 largest economies in the world and the European Union. More information available at: <https://g20.org/>

²² The table presents the number of ccTLD domains according to the indicated sources. The figures correspond to the record published by each country, considering members from the OECD and G20. For countries that do not provide official statistics supplied by the domain name registration authority, the figures were obtained from: <https://research.domaintools.com/statistics/tld-counts>. It is important to note that there are variations among the date of reference, although the most up-to-date data for each country is compiled. The comparative analysis for domain name performance should also consider the different management models for ccTLD registration. In addition, when observing rankings, it is important to consider the diversity of existing business models.

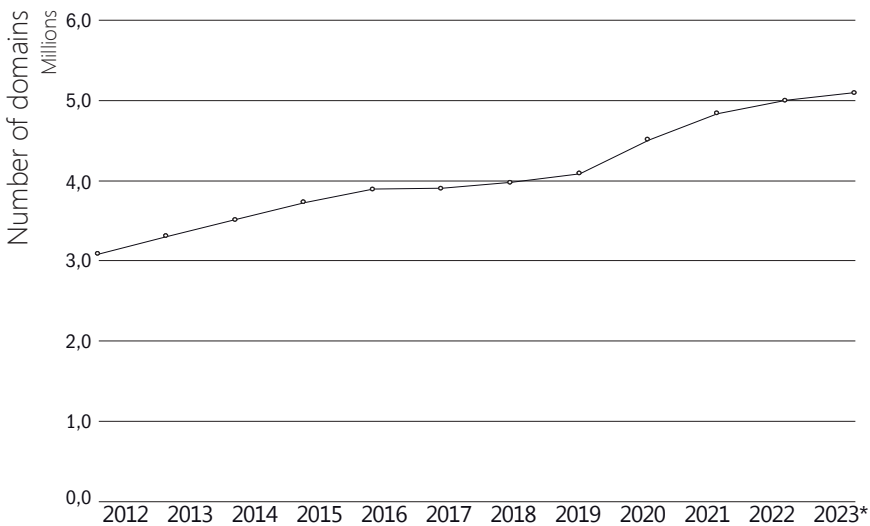
Table 1 – TOTAL REGISTRATION OF DOMAIN NAMES AMONG OECD AND G20 COUNTRIES

Position	Country	Number of domains	Date of reference	Source (website)
1	Germany (.de)	17,498,904	31/03/2023	https://www.denic.de
2	United Kingdom (.uk)	9,659,204	28/02/2023	https://www.nominet.uk/news/reports-statistics/uk-register-statistics-2023/
3	China (.cn)	7,193,640	31/03/2023	https://research.domaintools.com/statistics/tld-counts/
4	Netherlands (.nl)	6,295,609	31/03/2023	https://stats.sidnlabs.nl/en/registration.html
5	Brazil (.br)	5,094,470	31/03/2023	https://registro.br/dominio/estatisticas/
6	Russia (.ru)	4,935,204	31/03/2023	https://cctld.ru
7	Australia (.au)	4,214,524	31/03/2023	https://www.auda.org.au/
8	France (.fr)	3,975,180	31/03/2023	https://research.domaintools.com/statistics/tld-counts/
9	European Union (.eu)	3,669,172	31/03/2023	https://research.domaintools.com/statistics/tld-counts/
10	Italy (.it)	3,493,029	31/03/2023	http://nic.it
11	Colombia (.co)	3,365,252	31/03/2023	https://research.domaintools.com/statistics/tld-counts/
12	Canada (.ca)	3,361,681	31/03/2023	https://www.cira.ca
13	India (.in)	2,893,157	31/03/2023	https://research.domaintools.com/statistics/tld-counts/
14	Switzerland (.ch)	2,535,872	15/03/2023	https://www.nic.ch/statistics/domains/
15	Poland (.pl)	2,509,765	31/03/2023	https://www.dns.pl/en/
16	Spain (.es)	2,024,766	20/03/2023	https://www.dominios.es/dominios/en
17	United States (.us)	1,932,390	31/03/2023	https://research.domaintools.com/statistics/tld-counts/
18	Belgium (.be)	1,746,750	31/03/2023	https://www.dnsbelgium.be/en
19	Japan (.jp)	1,728,299	01/03/2023	https://jprs.co.jp/en/stat/
20	Portugal (.pt)	1,678,130	31/03/2023	https://www.dns.pt/en/statistics/

Collection date: March 31, 2023.

Chart 1 shows the performance of .br since 2012.

Chart 1 – TOTAL NUMBER OF DOMAIN REGISTRATIONS FOR .BR – 2012 to 2023*



* Collection date: March 31, 2023.
Source: Registro.br
Retrieved from: <https://registro.br/dominio/estatisticas/>

In March 2023, the five generic Top-Level Domains (gTLD) totaled more than 190.08 million registrations. With 159.71 million registrations, .com ranked first, as shown in Table 2.

Table 2 – TOTAL NUMBER OF DOMAINS AMONG MAIN gTLD

Position	gTLD	Number of domains
1	.com	159.717.469
2	.net	13.030.989
3	.org	10.754.641
4	.info	3.738.226
5	.xyz	3.649.013

Collection date: March 31, 2023.
Source: DomainTools.com
Retrieved from: research.domaintools.com/statistics/tld-counts

ANALYSIS OF BIG DATA IN THE PUBLIC SECTOR

With the increasing adoption of digital technologies, individuals, machines, systems, and sensors generate a large amount of data. In this context, a new data ecosystem has been establishing itself, in which Big Data sources are making unprecedented analyses and gathering information

for improving decision-making processes, including in the public sector.

In 2021, 25% of all federal and state government organizations performed Big Data analyses. The following indicators²³ show how the Brazilian public sector uses or does not use Big Data.

Federal and state government organizations that performed Big Data analyses by branch *Federal and state government organizations (2021)*

Legislative Branch



Judiciary Branch



Executive Branch



²³ Data from the ICT Electronic Government 2021 survey, from Cetic.br|NIC.br. Available at: <https://cetic.br/en/pesquisa/governo-eletronico>

/Answers to your questions

Federal and state government organizations by branch and reasons for not performing Big Data²⁴ analyses

Federal and state government organizations (2021)

	Lack of qualified personnel in the government organization to perform Big Data analyses	It is not a priority for the government organization	Lack of need or interest
Total	41%	31%	30%
 Legislative Branch	33%	20%	30%
 Judiciary Branch	64%	54%	43%
 Executive Branch	40%	30%	30%

²⁴ Other reasons for not analyzing big data that were identified in the ICT Electronic Government 2021 survey can be found at: <https://cetic.br/en/tics/governo/2021/orgaos/H1B/>

/Credits

TEXT

DOMAIN REPORT

Thiago Meireles (Cetic.br | NIC.br)

GRAPHIC DESIGN

Giuliano Galves, Larissa Paschoal, and Maricy Rabelo
(Comunicação | NIC.br)

PUBLISHING

Grappa Marketing Editorial

ENGLISH REVISION AND TRANSLATION

Ana Zuleika Pinheiro Machado and Robert Dinham

EDITORIAL COORDINATION

Alexandre F. Barbosa, Graziela Castello, Javiera
F. M. Macaya, and Mariana Galhardo Oliveira
(Cetic.br | NIC.br)

ACKNOWLEDGMENTS

Anita Job Lübbe (TRT-4)

Carolina Rossini (The Datasphere)

Datasphere Initiative

Delfina Soares (UNU EGOV)

Marlova Jovchelovitch Noleto and Aduino Candido
Soares (UNESCO – Brazilian Office)

Moinul Zaber (UNU EGOV)

Pedro Luis Nascimento da Silva (SCIENCE)

ABOUT CETIC.br

The Regional Center for Studies on the Development of the Information Society – Cetic.br (<https://www.cetic.br/en/>), a department of NIC.br, is responsible for producing studies and statistics on the access and use of the Internet in Brazil, disseminating analyzes and periodic information on the Internet development in the country. Cetic.br acts under the auspices of UNESCO.

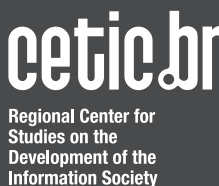
ABOUT NIC.br

The Brazilian Network Information Center – NIC.br (<http://www.nic.br/about-nic-br/>) is a non-profit civil Entity in charge of operating the .br domain, distributing IP numbers, and registering Autonomous Systems in the country. It conducts initiatives and projects that bring benefits to the Internet infrastructure in Brazil.

ABOUT CGI.br

The Brazilian Internet Steering Committee – CGI.br (<https://cgi.br/about/>), responsible for establishing strategic guidelines related to the use and development of the Internet in Brazil, coordinates and integrates all Internet service initiatives in the country, promoting technical quality, innovation, and dissemination of the services offered.

*The ideas and opinions expressed in the texts of this publication are those of the respective authors and do not necessarily reflect those of NIC.br and CGI.br.



CREATIVE COMMONS
Attribution
NonCommercial
(by-nc)





STRIVING FOR A BETTER INTERNET IN BRAZIL

CGI.BR, MODEL OF MULTISTAKEHOLDER GOVERNANCE

<https://cgi.br>

nic.br cgi.br