

Governing Before Rules. Shadow AI among Legislative Advisors in Mexico City

Sergio A. Bárcena Juárez

Author ORCID Nr: 0000-0002-4860-6699

School of Humanities and Education, Tecnológico de Monterrey, sergio.barcena@tec.mx

Erik Cardona-Gómez

Author ORCID Nr: 0000-0002-7894-8622

Department of Politics and International Studies, SOAS, University of London, ec50@soas.ac.uk

Abstract: Generative AI enters legislative advisory work before parliaments adopt formal rules. This article examines this "adoption-rules gap" through an exploratory case study of the Mexico City Congress. Combining a diagnostic survey (N=42) and a co-design workshop, we investigate how frontline advisors manage "shadow AI" without stable institutional oversight. Findings reveal that while adoption is evenly divided, half of the surveyed advisors use AI for auxiliary tasks like document synthesis and preparatory drafting. Crucially, they articulate a shared practical boundary against the unsupervised delegation of final normative text and legal reasoning. Using a three-layer heuristic, we frame these informal rules-in-use as enacted internal governance. Ultimately, the study demonstrates that while informal safeguards help manage short-term cybersecurity and accountability risks, they remain fundamentally fragile until codified through formal training, secure infrastructure, and documentation procedures.

Keywords: Public-Sector AI Governance; Shadow AI; Legislative Advisors; Algorithmic Governance; Parliamentary Digitalization.

Statement on the use of AI: During the preparation of this work, the authors used Google Gemini solely to improve readability, format references, and proofread the manuscript. The tool was not used to generate the research design, data, analysis, findings, or conclusions. The authors reviewed and edited all AI-assisted output and take full responsibility for the final content of the publication.

1. Introduction

The rapid spread of generative AI in government offices has shifted part of AI governance from formal institutional design to everyday administrative practice. Public employees can now use general-purpose systems without procurement, institutional deployment, or specialized technical infrastructure. Through private accounts, personal devices, and external platforms, they may incorporate these tools before their organizations have defined rules for confidentiality, verification, authorship, or accountability. In Latin America, recent assessments describe AI adoption as a growing public-sector priority, while implementation remains constrained by uneven technological infrastructure, limited specialized expertise, and still-developing governance frameworks (ECLAC & CENIA, 2024; IDB, 2024). Under these conditions, frontline personnel begin to make operational decisions about AI before their institutions have established usable standards.

Legislatures intensify this problem because their work culminates in authoritative normative text. Parliamentary activity depends on legal precision, reviewable procedures, and identifiable chains of responsibility. Uses of generative AI for transcript summarization, background research, style editing, comparative review, or preliminary drafting therefore raise questions that go beyond administrative efficiency. They affect document integrity, confidentiality, political intention, and the traceability of human judgment. Crucially, when external tools are used for these tasks, internal legislative documents, which often contain sensitive political strategy, draft legal provisions, personal data, and confidential committee opinions, are uploaded to commercial platforms entirely outside institutional control, introducing severe data governance and cybersecurity vulnerabilities. Some legislative institutions have begun to develop formal responses, including Brazil's Ulysses platform, European experiments such as Archibot, initiatives in Finland's Eduskunta, and guidelines promoted by the Inter-Parliamentary Union (IPU, 2024; Gómez Macfarland, 2025; Palmirani et al., 2024). In settings that still lack internal procedures for everyday AI use, staff may carry out public work with tools that remain outside stable rules for authorization, documentation, verification, and accountability.

Public administration scholarship has long examined how technology reshapes discretion. One influential account describes a movement from street-level discretion to system-level bureaucracy, where technical infrastructures embed classifications, rules, and decisions into administrative processes (Bovens & Zouridis, 2002). Later work offers a more contingent view, showing that automation can relocate discretion and alter how officials justify decisions, without removing human judgment from administrative practice (de Boer & Raaphorst, 2023). Building on this, Technology Enactment Theory (Fountain, 2001) posits that organizational members do not passively receive technology, but actively incorporate it through established practices, professional expectations, and institutional constraints. The Mexico City Congress follows this trajectory. Generative AI has not entered through a centralized technical system, formal procurement, or an authorized drafting environment. Instead, discretion has moved into the daily practices of legislative advisors who decide when to rely on external tools, how to verify their outputs, and which tasks should remain under human control.

Informal adoption of this kind corresponds to shadow AI, understood as the use of generative AI systems inside institutional workflows without formal authorization, training, or oversight. Unlike traditional shadow IT, shadow AI does not only involve unauthorized software or informal communication channels; generative systems can synthesize documents, compare legal materials, propose language, organize arguments, and assist with initial drafting phases. Recent work has already examined informal generative AI use in public institutions at a larger scale, including shadow GenAI among public managers in EU member states (Tangi et al., 2025) and public-sector professionals in the UK (Bright et al., 2025). Consequently, this article does not frame shadow AI as a newly introduced public-sector governance problem. Rather, it presents an exploratory case study of how a targeted group of legislative actors in a subnational Latin American congress articulates boundaries around AI use before formal institutional governance is in place.

The Mexico City Congress provides a revealing case for observing this adoption-rules gap under conditions of institutional opacity and uneven administrative capacity. Within the broader universe of Mexican state legislatures documented by INEGI's National Census of State Legislative Powers, the case is useful because the advisory layer remains difficult to observe through public administrative records (INEGI, 2024). To contextualize this administrative opacity, Appendix C provides a brief comparison of structural reporting capacities between Mexico City and other state legislatures. The institution lacks a reliable public map of its advisory workforce and does not appear to maintain formal AI governance arrangements. Claims regarding the absence of formal AI governance in this setting are triangulated through a lack of public administrative records, participant self-reports indicating a lack of formal training, and corroborating contextual expert interviews. Rule-making around AI therefore begins, in practice, among the staff who prepare, review, and translate political decisions into legislative documents.

This article examines how internal AI governance takes shape when formal rules, secure tools, training programs, and documentation procedures remain underdeveloped. The empirical design combines a targeted diagnostic survey of active legislative advisors (N=42) with a three-hour in-person co-design workshop involving approximately 50 participants, organized into teams that produced 14 structured response forms. This multi-method approach is designed to uncover hidden practices in a setting where formal administrative records provide limited access. The recruitment for the survey was coordinated through the logistical channels created for the workshop; consequently, the survey respondents and workshop attendees represent partially overlapping groups rather than identical populations. The study focuses on a strategically relevant subgroup whose practices show how generative AI enters legislative workflows. This design makes it possible to observe rules-in-use, understood as practical norms that participants apply in daily institutional life when formal rules remain absent or incomplete (Ostrom, 2005).

Across both data sources, advisors describe AI as useful for auxiliary tasks, such as summarizing long documents, organizing background information, editing style, and preparing initial analytical inputs. However, they articulate a practical boundary against delegating final normative text to AI without human review. While the rejection of unsupervised delegation is what one might reasonably expect from professionals working with legally consequential text, the detailed documentation of how this boundary is actively negotiated clarifies what changes in our theoretical understanding:

informal safeguards are shown to be highly contingent. Advisors differentiate between preparatory assistance and the automated production of legal provisions, committee opinions, or legislative instruments. That boundary reflects professional risk aversion, concern for auditability, and a defense of human authorship in democratic lawmaking. It does not amount to formal policy; rather, it describes an operational consensus produced by frontline actors who must maintain productivity without surrendering responsibility for the final normative document.

The article contributes to public-sector AI governance by exploring how these informal boundaries take shape before institutional codification. Conceptually, the article uses a three-layer heuristic of internal governance. The operational layer contains everyday practices and informal rules-in-use. The collective-choice layer captures the institutional space where formal rules, secure infrastructure, and training should be produced, but remain underdeveloped in the case examined. The normative layer contains the professional commitments that justify limits on AI delegation. This framework demonstrates that while informal safeguards can reduce short-term risks, their protective capacity remains fundamentally fragile when they depend on tacit consensus and individual judgment rather than durable institutional rules.

The article is structured as follows. Section 2 reviews the literature on public-sector AI governance, digital discretion, and shadow AI to contextualize the adoption-rules gap. Section 3 outlines the conceptual framework, detailing the three-layer heuristic of enacted internal governance. Section 4 presents the multi-method research design, clarifying the sampling scope and overlapping populations of the diagnostic survey and the co-design workshop. Section 5 details the empirical findings, divided into the operational reality of the adoption-rules gap and the normative boundaries articulated by advisors through the traffic-light protocol. Finally, Section 6 discusses the implications of these informal rules-in-use for parliamentary AI governance and concludes with practical recommendations for institutionalizing these safeguards.

2. Literature Review

Public-sector AI governance research increasingly treats artificial intelligence as an institutional capacity problem rather than a simple question of adoption. Governments are expected to use algorithmic systems to improve efficiency, responsiveness, and evidence-based decision-making, while many public bodies still lack the rules, skills, infrastructures, and oversight routines needed to govern those systems. Recent work places organizational readiness at the center of public-sector AI governance, including procurement standards, verification routines, auditability, data stewardship, and accountability mechanisms (OECD, 2024). Generative AI gives this concern a sharper organizational form because many officials encounter these systems before their organizations have defined who may use them, for which tasks, with what safeguards, and under what forms of documentation.

Parliaments turn this governance problem into a problem of legal and political text. Legislative institutions depend on legal precision, procedural reviewability, and identifiable chains of responsibility. Parliamentary guidance recognizes that AI may support legislative research, administrative

work, citizen communication, and drafting-related tasks, while also creating risks that require human oversight, transparency, accountability, and clear internal rules (IPU, 2024; WFD, 2024). Work on AI in parliamentary contexts gives this problem a more precise technical and legal form. Once AI enters drafting environments, the concern extends beyond efficiency to traceability, legal consistency, reviewability, and human responsibility in the production of normative text (Palmirani et al., 2024).

Formal frameworks offer a necessary vocabulary for governing AI in parliaments, but they often assume that institutions can act through strategy, formal policy, pilot projects, data governance, risk management, and training. That assumption fits digitally mature legislatures better than institutions where staff already use generative tools before formal AI governance arrangements are in place. Guidelines for AI in parliaments address institutions at different levels of digital maturity (IPU, 2024), yet they still imagine parliaments as organizations able to adopt rules from the top. The problem examined here appears earlier in the institutional sequence. Generative AI enters through offices, advisors, personal accounts, informal routines, and unmonitored experimentation before the legislature has produced usable standards.

Legal drafting and legislative analysis create specific risks when AI use remains informal. Generative systems can assist with summarization, comparison, classification, translation, and preliminary writing, but those functions can also shape arguments and normative language. Crucially, in a legislative setting, these tasks often involve uploading internal documents to commercial platforms that operate entirely outside institutional control. Legislative documents routinely contain sensitive political strategy, draft legal provisions, confidential committee opinions, personal data, and information protected under applicable data protection frameworks. The informal use of third-party platforms therefore introduces severe data governance and cybersecurity vulnerabilities, an essential gap in contexts lacking formal oversight. Furthermore, concerns about technological due process are highly relevant because automated systems may embed errors, biases, or opaque classifications into public decisions while limiting the ability of affected actors to inspect or contest them (Citron, 2008). Black-box systems complicate public oversight and make institutional responsibility harder to assign (Fink, 2018; Busuioc, 2021), which directly affects the reviewability of the legislative drafting process.

The literature on discretion helps explain why generative AI does not simply automate public work. Earlier accounts of digital government described a movement from street-level discretion to system-level bureaucracy, where information systems organize administrative decisions through embedded rules and classifications (Bovens & Zouridis, 2002). Later studies offer a more contingent account: automation relocates discretion, mostly when frontline officials must interpret system outputs, correct errors, or decide when automated tools should be trusted (de Boer & Raaphorst, 2023; Giest & Klievink, 2024). This contingent view aligns closely with Technology Enactment Theory (Fountain, 2001), which proposes that organizational members do not passively receive technology, but actively incorporate and adapt it through established practices, routines, and institutional constraints. Legislative advisors do not merely receive a finished technological system; they enact it. They decide how to use external tools, how much weight to give their outputs, and where to draw boundaries around acceptable delegation.

Algorithmic mediation can enter public authority before any official automation system exists. The concept of artificial discretion captures this risk because algorithmic systems may structure, recommend, or partly substitute human judgment in public work (Young et al., 2019). Legislative advisors may use generative models to synthesize documents, compare legal frameworks, draft explanatory language, or prepare speeches. These tasks do not automatically transfer decision-making authority to AI, but they can shift part of the cognitive process into systems that the institution does not control. The unresolved question is how much analytical labor has already moved into unmonitored tools before a parliament formally adopts AI.

Unauthorized digital tools have usually been studied through the literature on shadow IT. Shadow AI extends that problem because generative systems produce language, classify inputs, generate summaries, and suggest argumentative structures. Recent studies have demonstrated that this phenomenon is widespread, which is seen in work that documents shadow generative AI use among public managers across EU member states (Tangi et al., 2025) and public-sector professionals in the UK (Bright et al., 2025). In public institutions, the risks of such informal adoption include confidentiality breaches, weak verification, inconsistent standards, and the unchecked delegation of analytical tasks to private systems (Silic et al., 2025; Socitm, 2025). Legislatures add a further difficulty because their outputs require identifiable human authorship and procedural traceability.

Policy advisory systems research helps locate the actors through whom this informal adoption occurs. Existing work has examined how governments receive expert knowledge, how advisory networks are organized, and how policy advice circulates within formal institutions (Halligan, 1995; Craft & Howlett, 2012). Much of this literature focuses on visible advisory arrangements, senior policy professionals, expert commissions, and institutionalized channels of advice. Legislative advisors occupy a less visible position. They prepare background materials, review legal language, compare policy alternatives, and translate political instructions into draftable text, often through partisan or office-based arrangements. That position makes them key to AI adoption in practice, even when formal AI governance remains underdeveloped.

Street-level bureaucracy offers an analogy for interpreting this internal advisory layer, although the fit is not exact. The concept usually refers to frontline workers who exercise discretion while interacting with citizens under resource constraints (Lipsky, 1980). Legislative advisors occupy a different institutional position. Their immediate clients are elected representatives, and the resource they manage is cognitive and procedural rather than distributive. The comparison remains analytically useful because advisors also work under heavy workload pressures, incomplete rules, and limited organizational support. Their daily decisions can become the institution's effective rules-in-use when formal rules are absent. In that sense, the advisory layer operates as an internal site where AI governance is enacted through practice rather than issued from above.

Latin American legislatures make the adoption-rules gap particularly visible because AI adoption often occurs in institutions with uneven administrative capacity, weak technological infrastructure, and limited internal professionalization (ECLAC & CENIA, 2024; IDB, 2024). In Mexico, the absence of comprehensive national AI legislation has coincided with limited parliamentary guidance for everyday AI use (Gómez Macfarland, 2025). Available studies offer limited empirical evidence on how

legislative advisors themselves use generative AI, how they define acceptable and unacceptable uses, and how they produce informal safeguards in the absence of formal governance.

The gap addressed in this article lies in the interval between informal adoption and formal institutionalization. Existing AI governance frameworks explain how parliaments should govern AI once institutional capacity can be mobilized. Legislative informatics identifies the risks of automated drafting and the need for traceable, reviewable systems. Public administration research explains how discretion changes when digital tools enter bureaucratic work. Less is known about the stage before formal governance, when legislative staff already use generative AI and must create practical boundaries without training, infrastructure, or enforceable rules.

To analyze this interval, this article utilizes a framework capable of connecting informal practice with institutional rule formation. Institutional analysis distinguishes between rules-in-form and rules-in-use, allowing researchers to examine how organizations are governed by practical norms even when formal rules remain absent or incomplete (Ostrom, 2005). Combined with participatory design principles (Ehn, 2008) and technology enactment, this framework allows informal AI adoption to be analyzed as a form of enacted internal governance rather than as a mere violation of organizational procedure.

Formal principles, parliamentary guidelines, and institutional designs leave the advisory layer less visible, where generative AI already enters legislative work. This article addresses that blind spot through an examination of how legislative advisors in the Mexico City Congress construct practical boundaries around shadow AI and how those boundaries can be understood as emerging rules-in-use.

3. Conceptual Framework: The Adoption-Rules Gap and Enacted Internal Governance

Generative AI creates an adoption-rules gap when its de facto use in parliamentary work advances before de jure rules, secure tools, training routines, or oversight procedures are available. Staff incorporate generative tools into daily workflows before bureaucratic procedures can be revised. Because adoption often occurs in informal networks while regulation depends on centralized coordination, formal authorities may lack detailed knowledge of daily AI practices. The gap becomes particularly consequential in legislative work because generative AI can enter the production of normative text. Once external tools become part of drafting-related work without institutional standards, the practical judgments of the advisory layer become the primary, albeit provisional, site of governance.

To analyze how this advisory layer manages the gap, we rely on the concept of enacted internal governance. This refers to the practical ordering of AI use by institutional actors who lack formal rule-making authority but still create operational boundaries in daily work. This concept builds upon Technology Enactment Theory, which posits that digital tools acquire institutional meaning

through situated use. Organizational members do not passively receive technology; rather, they actively incorporate and adapt it through established practices, professional routines, and institutional constraints (Fountain, 2001; Orlikowski, 2000). The distinction between rules-in-form and rules-in-use further explains how organizations can be governed by these enacted practical norms when formal rules remain absent or incomplete (Ostrom, 2005).

A further analogy comes from street-level bureaucracy, although legislative advisors do not interact with citizens in the classic sense. Their discretion is internal, cognitive, and procedural, exercised under workload pressure, incomplete guidance, and limited organizational support (Lipsky, 1980). Their daily decisions can become the institution’s effective rules-in-use when formal rules are absent.

To operationalize this framework for the empirical data, we utilize a three-layer heuristic that separates the sites where informal AI boundaries are practiced, justified, and left uncodified. 1) the operational layer contains everyday decisions about whether to use AI for summarization, comparison, style editing, research support, or preliminary drafting. These decisions operate as effective rules-in-use when formal instructions are missing. 2) the collective-choice layer captures the institutional capacity to convert practice into rules, secure infrastructure, training, procurement standards, documentation routines, and accountability procedures. In the Mexico City case, this layer remains underdeveloped in the domains examined here, leaving AI governance dependent on localized judgment. 3) the normative layer contains the professional commitments that shape operational boundaries. In legislative work, those commitments include authorship, auditability, document integrity, and responsibility for the final normative text. Table 1 illustrates this three-layer model of internal governance, outlining the institutional reference, technological logic, and specific function of each layer within the case study.

Table 1: The Three-Layer Model of Internal Governance

Layer	Institutional reference	Technological logic	Function in the case study
Operational	Rules-in-use and operational discretion (Lipsky, 1980; Ostrom, 2005)	Enacted technology and behavioral adaptation (Fountain, 2001; Orlikowski, 2000)	Advisors create informal boundaries around AI use, including a shared prohibition against automated drafting of normative text.
Collective-choice	Rule codification and organizational coordination (Ostrom, 2005)	Infrastructure, standards, and formal coordination (OECD, 2024; IPU, 2024; WFD, 2024)	Weaknesses in formal rules, secure infrastructure, training, and documentation prevent daily practices from becoming

			stable institutional standards.
Normative	Professional commitments and institutional responsibility (Ostrom, 2005)	Embedded norms, accountability, and reviewability expectations (Citron, 2008; Busuioc, 2021; Palmirani et al., 2024)	Professional risk-aversion, authorship, auditability, and responsibility for legal text justify operational safeguards.

Within this framework, shadow AI functions as a provisional bypass. Advisors enact external generative systems when secure tools, clear standards, or practical guidance are not yet available in the relevant institutional domains. This arrangement helps them manage workload and maintain legislative output, and shifts parts of the cognitive process into systems outside formal control. Through an analysis of the boundaries these advisors draw, which are most clearly observed in their refusal to allow AI to generate normative text without human authorship and verification, we can trace how frontline workers create practical rules before the organization can formalize them.

4. Methodology and Research Design

Informal AI use inside legislative advisory work is difficult to observe because relevant practices often occur through personal accounts, external platforms, and task-specific routines rather than through formal records. For that reason, this study uses an exploratory multi-method case study of the Mexico City Congress, which combines two contextual expert interviews, an online diagnostic survey, and an in-person co-design workshop. Claims regarding the absence or weakness of formal AI governance in this setting are based on a triangulation of three sources: the lack of public administrative records establishing such rules, participant self-reports that indicate a lack of formal training and guidance, and corroboration from the contextual expert interviews.

The Mexico City Congress was selected as a strategic high-opacity case. With 66 legislators and no centralized public roster of legislative advisors, the chamber offers limited visibility into the staff layer where much legislative preparation occurs. Based on contextual interviews and institutional estimates, the advisory workforce is approximately 180 staff members, a figure treated as an informed approximation rather than an official count. Administrative opacity matters for the research design because it limits the possibility of constructing a conventional sampling frame and makes formal top-down observation of advisory work difficult. The case is therefore useful to examine how internal boundaries around AI are articulated when AI-specific policy, secure institutional tools, documentation standards, and systematic training remain weak or unevenly developed.

Active legislative advisors were defined as currently contracted or formally attached staff members performing legislative support functions. This category included legal drafting, policy analysis, legislative research, committee work, speechwriting, communication, document review, agenda tracking, and strategic support to legislative decision-making. Rather than treating advisory work

as a single job title, the definition captures a functional field. It includes staff working in legislative offices, parliamentary groups, committees, and technical units, provided that their daily work involved the preparation, review, interpretation, or communication of legislative materials.

Before the workshop, two contextual expert interviews were conducted with actors familiar with legislative drafting, committee work, and institutional organization in the Mexico City Congress. Conducted in October and November 2025, these interviews helped refine the vocabulary of legislative work, identify likely points of AI adoption, and locate institutional constraints relevant to the survey and workshop protocol. Their role was contextual rather than evidentiary in the main empirical analysis. Interview accounts framed AI use as a problem involving infrastructure, confidentiality, legal precision, political intention, and human review in drafting-related tasks, not only as a technical matter.

An online diagnostic survey was administered through Google Forms and closed on November 14, 2025, the day of the in-person session. Recruitment was conducted via digital professional networks and coordinated through the logistical channel created for the workshop, with completion requested prior to attendance to collect baseline information. Consequently, the 42 survey respondents and the approximately 50 workshop participants represent partially overlapping groups, rather than identical populations. Respondents created anonymous pseudonyms, and after incomplete or invalid entries were removed, the final database contained 42 valid individual responses. Because this is a non-probability, purposive sample, all survey results, particularly subgroup breakdowns, are presented strictly as illustrative observations of this specific group and carry no statistical inferential weight.

Survey items measured the conditions under which generative AI enters legislative advisory work. Block one collected demographic and professional information, including current role, type of institutional placement, team size, and years of advisory experience. Block two measured technological resources provided by the chamber, including internet access, equipment, technical support, and cybersecurity. Workload pressure was captured through questions on urgent assignments, work outside office hours, work from home, planning capacity, time-consuming tasks, and perceived resource constraints. Block 3, which was a skills block, asked respondents to assess their capacities in key legislative tasks, while Block four examined AI adoption, learning pathways, tools used, perceived risks, and expected effects on job security, quality, and speed. Block five measured preferences regarding human supervision, documentation units, confidentiality protocols, valid-use rules, and technical training.

On November 14, 2025, a three-hour in-person co-design workshop brought together approximately 50 people, including legislators, legislative advisors, and staff members involved in legal drafting, committee work, communication, and political support. While legislators were present, the workshop's team-based design captured a collective consensus regarding legislative workflows, keeping the analytical focus anchored on the advisory and administrative realities of the institution. Designed as a structured normative-elicitation exercise, the workshop examined whether a heterogeneous legislative group could identify ethical, procedural, and professional risks associated with

AI use, define those risks in their own terms, and converge around shared boundaries despite occupying different institutional positions. Before participants began their debate, the research team provided a brief explanation that distinguished between auxiliary, conditional, and prohibited AI uses. We acknowledge that this framing structured the exercise and likely influenced the tripartite categories that participants subsequently produced.

Participants evaluated possible AI uses from a position-neutral standpoint. They were asked to define general criteria for responsible legislative work rather than to report or defend the practices of a specific office, party, or role, reducing pressure to disclose personal behavior and redirecting discussion toward institutional standards. Each team deliberated over four phases of legislative work, following a traffic-light protocol in which teams classified AI practices as desirable (green), neutral or doubtful (yellow), or undesirable (red). Fourteen analyzable team forms came out of the workshop. This figure refers to collective submissions, not to individual participants; the gap between approximately 50 attendees and fourteen forms therefore reflects the team-based design of the exercise rather than individual non-response.

Workshop responses were segmented into a coding matrix linking legislative phase, traffic-light classification, governance layer, and risk rationale. First, a deductive pass used the three-layer framework to classify units as operational practice, collective-choice absence, or normative boundary. A second inductive pass identified recurrent subthemes, including hallucination risk, confidentiality, verification, loss of political intention, authorship, overreliance, auxiliary use, and prohibited delegation. To construct the consolidated findings (Table 5), raw team forms were synthesized through the grouping of synonymous descriptions of tasks and risks under standardized labels, an approach that ensured the final output accurately reflected the team-based consensus while removing redundancies. Ambiguous cases were reviewed against the full team response and assigned to the category that best captured the stated rationale.

Individual diagnostic perceptions and collective normative classifications were compared without treating either source as a full institutional audit. Survey data cannot estimate the prevalence of AI use across the entire Congress because the sample is not probabilistic. The analysis also relies on self-reported practices and perceptions rather than direct observation of AI use, system logs, or document-level audits. These limits define the scope of inference. The study maps how a strategically relevant group of legislative actors articulates practical boundaries around generative AI before stable formal AI-specific rules are available to participants.

Finally, participant protection fundamentally shaped the design because informal AI use can expose staff to reputational or administrative risk. Survey and workshop instruments did not collect names, party affiliations, office identifiers, confidential legislative drafts, sensitive political strategies, or internal documents. Responses were anonymized and analyzed in aggregate. Reported examples are presented without information that could identify participants, offices, legislators, or specific legislative negotiations. This protocol allowed participants to discuss tacit workflows and risk perceptions without disclosing confidential materials or compromising their professional positions.

5. Empirical Findings

Survey and workshop evidence show how informal AI governance took shape before the institution had stabilized rules for use. Survey responses capture the operational side of the gap. Adoption was recent, uneven, concentrated in commercial tools, and located among respondents close to legal, political, and communicative tasks. Workshop responses capture the normative side. Participants accepted AI as support for synthesis, organization, comparison, style editing, and preparatory drafting, but rejected the delegation of final normative text, legal reasoning, political judgment, and public representation without human review. The findings therefore trace how utility and restraint developed together under weak AI-specific governance.

5.1. Uneven Adoption, Operational Pressure, and the Adoption-Rules Gap

Responses to the survey show the operational side of the gap through an evenly divided field. Among the 42 valid responses, exactly 21 respondents reported using AI tools in congressional tasks and 21 reported no such use. Generative AI had become relevant for a substantial subgroup of legislative staff without becoming a generalized feature of advisory work. Some respondents had integrated these tools into organization, synthesis, and writing support, while others continued to perform those tasks without AI. Everyday practice was already uneven before a stable institutional settlement had emerged.

Recent adoption reinforces that reading. 13 of the 21 active users had used AI tools for six months or less, and 17 had incorporated them within the previous year. These responses refer to experimentation inside congressional work, not to consolidated institutional use. AI entered advisory routines while participant-reported guidance, training, and secure internal infrastructure remained incomplete. The adoption-rules gap appears in that sequence, as practice advanced before institutional codification.

It is important to note that the sub-group breakdowns below (by age cohort and team size) are based on small cell sizes of fewer than 10 respondents. Consequently, these figures carry no inferential weight and are presented strictly as illustrative observations with no pattern-level implications. Among the 42 respondents, AI use was reported by 5 of the 6 participants aged 25–29, 8 of the 10 participants working in teams of 4–6, and none of the 4 participants working alone. Given the small number of observations in these categories, these figures describe only the composition of this sample and do not support conclusions about the effects of age or team size on AI adoption. Without formal training, everyday team interaction becomes a channel to observe examples, exchange prompts, compare outputs, and adapt routines.

AI use in this sample was not confined to peripheral administrative tasks. Many respondents occupied roles close to legal, political, and communicative outputs. Most reported high ideological proximity to the legislator, caucus, or office they support, and large shares reported strong expectations of alignment in the documents and responses they produce. A majority also described themselves as close to key political and legislative decisions. AI-supported tasks therefore appear in work linked to legal drafting, strategic communication, party alignment, and responsibility for legislative

outputs. Table 2 details these illustrative sub-group breakdowns, demonstrating the distribution of AI adoption across different age cohorts, team sizes, and levels of political embeddedness.

Table 2: AI adoption by cohort, team size, and political embeddedness

Dimension	Category	Base	Count	Share
Age cohort	20–24	6	3	50%
	25–29	6	5	83%
	30–34	4	2	50%
	35–39	5	3	60%
	40–44	9	5	56%
	45–49	6	2	33%
	50–54	6	1	17%
	55–59	6	1	17%
Team size	Works alone	4	0	0%
	2–3 people	19	7	37%
	4–6 people	10	8	80%
	7 or more people	9	6	67%
Political embeddedness	High ideological similarity with legislator, caucus, or office	42	35	83%
	High demand for alignment or discipline in documents and responses	42	29	69%
	High involvement in key political and legislative decisions	42	27	64%

Note. For age cohort and team size, counts refer to respondents who reported AI use within each category. For political embeddedness, counts refer to respondents selecting scores 3 or 4 on a four-point scale. Survey N=42. These sub-group figures are illustrative observations, not generalizable patterns.

Tool choice also places adoption outside a controlled congressional environment. 16 of the 21 respondents who used AI reported using ChatGPT, while Gemini appeared in five cases and Copilot did not appear among active users. The concentration around general-purpose commercial platforms supports the shadow AI interpretation. External tools entered legislative workflows through individual access, personal experimentation, and ordinary office routines. Learning occurred through the same informal channels: 64% reported that they learned through personal exploration, trial, and error, and 29% from colleagues. Among active AI users, no respondent identified official congressional training as the way they learned to use AI for work.

Workload pressure explains part of the practical appeal. A total of 57% of respondents reported receiving urgent or last-minute requests at least three to four times per week. Last-minute writing or editing was selected by 60% of respondents as one of the most urgent tasks in their work, followed by the preparation of speeches at 50%. The same pressure extended beyond ordinary office hours. More than half of respondents (52%) reported working at least four hours outside regular working hours each week. Dispersed or hard-to-find information was selected by 48% as a frequent obstacle, and 45% identified research as consuming the most time.

Respondents described a gap between basic technical assistance and the conditions needed to govern AI use. A majority rejected the statement that internet access in their office was fast and stable. Cybersecurity also appeared as a weak point, with 45% disagreeing that the chamber's information-security infrastructure was sufficient to mitigate risks when using external tools. Contextual interview evidence provides background support for this interpretation. One institutional account described AI use in the Mexico City Congress as dependent on individual offices rather than on a chamber-wide policy, framing it as "a personal matter of each deputy, technician, or advisor." Table 3 summarizes these findings, outlining the specific indicators of operational pressure, informal learning pathways, and institutional constraints that characterize this adoption-rules gap.

Table 3: Operational pressure, informal learning, and institutional constraints

Dimension	Indicator	Share
Adoption and tools	Uses AI in congressional work	50%
	Uses ChatGPT among active AI users	76%
	Has used AI for six months or less among active AI users	62%
	Has used AI for one year or less among active AI users	81%
	Noticed increased AI use in the workplace in the last 12 months	57%
	Perceives advisors aged 35 or younger as the group adopting AI most visibly	76%
Learning pathways	Self-directed learning through exploration, trial, and error	64%
	Learning with colleagues	29%
	Official congressional or chamber training	7%
Operational pressure	Urgent requests at least three to four times per week	57%
	Last-minute writing or editing among most urgent tasks	60%
	Speeches or public interventions among most urgent tasks	50%

	Works at least four hours outside regular working hours weekly	52%
	Works from home or outside the office at least once or twice per week	60%
	Dispersed or hard-to-find information among most frequent obstacles	48%
	Research and source collection among most time-consuming tasks	45%
	Would like to devote more time to legislative drafting	50%
Institutional constraints	Disagrees that office internet is fast and stable	62%
	Disagrees that internet in other chamber spaces is fast and stable	62%
	Disagrees that cybersecurity infrastructure is sufficient	45%
	Agrees that technical support solves software or hardware problems	69%

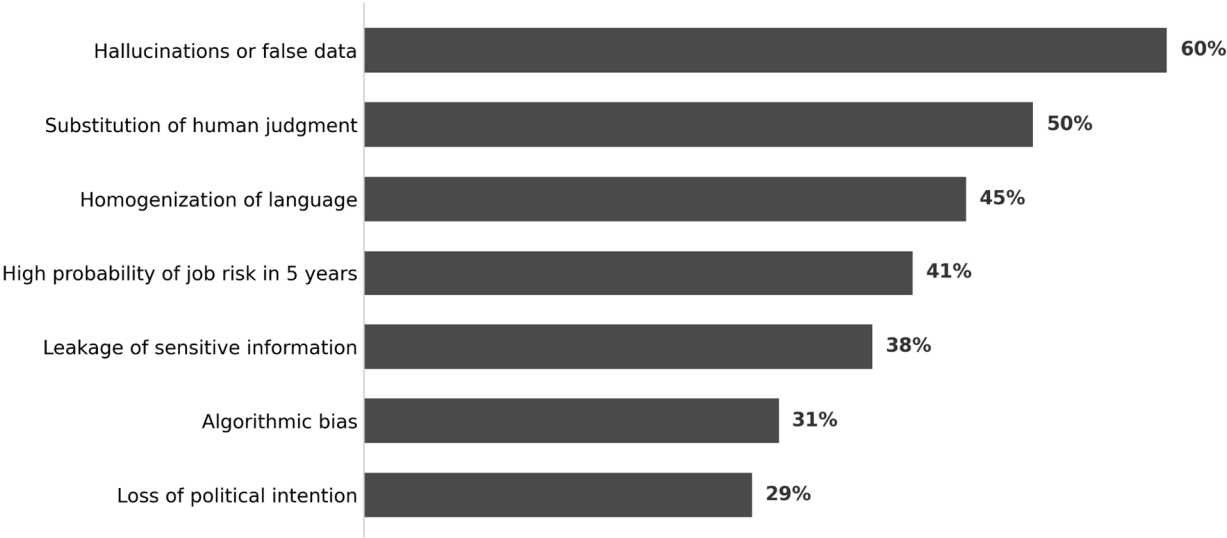
Note. Percentages calculated over the full sample (N=42) unless indicated otherwise (active users N=21). Multi-select indicators report the percentage of respondents who selected each option. Percentages are rounded to whole numbers.

5.2. Risk Perception, Red Lines, and the Traffic-Light Protocol

Respondents framed risk less as job displacement than as a problem of reliability, judgment, and authorship. Hallucinations or false data were the most frequently selected concern (60%). Half of the sample identified the substitution of human judgment as a major risk, and 45% pointed to the homogenization of legislative language. Leakage of confidential or sensitive information was selected by 38%. Only 41% assigned a high probability to AI placing their job at risk within five years, indicating that respondents expressed stronger concern about degradation in work quality, verification, and political meaning than about immediate displacement.

Political judgment marked the clearest limit to AI usefulness. Political operations and territorial management were selected by 62% of respondents, while relations with leadership or political negotiation were selected by 52%. AI was not treated as a substitute for forms of knowledge that depend on interpersonal judgment, organizational memory, and political timing. Consequently, respondents demanded governed use. Mandatory human supervision of final AI output was rated “important” or “very important” by 93% of respondents. Continuous technical training focused on ethics and strict information-security protocols both reached over 90% in importance. Figure 1 illustrates the specific perceived risks, while Table 4 quantifies the political limits of AI usefulness and the overwhelming demand for institutional governance protocols.

Figure 1: Perceived Risks of Artificial Intelligence in Legislative Work



Note. Data reflects the percentage of respondents selecting each risk. N=42.

Table 4: Political Limits of AI Usefulness and Governance Demands

Dimension	Indicator	Share
Tasks where AI is not considered useful	Political operations and territorial management	62%
	Relations with leadership and political negotiation	52%
	Coordination and mediation within the team or with other offices	31%
	Strategic design, contextual reading, and crisis anticipation	24%
Governance demands	Mandatory human supervision of final AI output, "very important"	64%
	Mandatory human supervision of final AI output, "important" or "very important"	93%
	Continuous technical training focused on ethics and verification, "important" or "very important"	95%
	Strict information-security and confidentiality protocols, "important" or "very important"	93%
	Rules defining valid uses in the chamber, "important" or "very important"	88%
	Internal caucus guidelines on permitted and prohibited tools, "important" or "very important"	88%

	Centralized documentation and AI analysis unit, “important” or “very important”	88%
--	--	-----

Note. Percentages are calculated over the full survey sample, N=42, and rounded to whole numbers. Risk and task indicators are multi-select items. Governance demand indicators are based on four-point importance scales.

In the co-design workshop, participants translated these concerns into practical classifications across four phases of legislative work. Green classifications referred to uses considered acceptable under ordinary conditions. Yellow classifications grouped uses requiring verification, contextual judgment, source control, or additional safeguards. Red classifications marked uses considered incompatible with responsible legislative work.

Bill drafting centered on the distinction between assistance and authorship. Participants accepted AI for synthesis, comparison, preliminary organization, and initial technical drafts subject to later review. Uses became conditional when AI generated policy ideas without a clear political position, searched for data without reliable sources, or supported drafting without subsequent review. Final legislative authorship drew the firmest boundary. Participants rejected the automated drafting of complete initiatives or generating legal foundations without updated sources. Several responses described the risk of “humanizing” the tool, which occurs when staff trust it as if it knew the social problem or allow it to replace human reasoning.

Committee analysis and dictamen work kept legal assessment under human control. Participants accepted AI for comparing versions, detecting duplications, and summarizing external opinions. Conditional uses involved impact estimates without verifiable data or summaries requiring legal review. Red classifications flatly rejected full dictamen generation, reliance on AI outputs as the sole basis for evaluation, or legal analysis without human reasoning. A contextual interview illustrated this boundary by noting that during a recent reform of the Organic Law, AI helped review over 400 articles to identify inconsistencies, representing a high-value auxiliary use. However, as the interviewee noted, a change in “a verb, a noun, or the construction of a sentence” alters legislative intent, reinforcing why normative drafting requires human review.

Speeches were treated as situated political interventions. Participants accepted AI when it helped organize ideas, prepare rough speech drafts, and improve style. Conditional uses included automatic counterarguments or speeches lacking source verification. Red classifications rejected speeches generated entirely by AI without review, real-time unverified AI responses, and outputs containing dignity-violating content. The red line protected ideological alignment, audience sensitivity, and public language responsibility.

Public communication raised sharp risks because outputs reach citizens directly. Participants accepted AI for citizen summaries, press releases based on verified inputs, and audience segmentation. Conditional uses included generic posts without institutional tone or messages lacking cultural sen-

sitivity. Red classifications strictly prohibited false notes, unverified information, automated simulation of citizen interaction, and synthetic media that could harm personal integrity. It is important to note that both Table 5 and the preceding text must be understood with the caveat that these 'green-light' uses, such as document comparison, synthesis, and preparatory drafting, are considered acceptable only under conditions of proper confidentiality, strict data-protection compliance, and the use of approved tools. Table 5 consolidates these workshop findings, detailing the specific acceptable, conditional, and prohibited AI uses across the four phases of legislative work.

Table 5: Co-designed traffic-light protocol for legislative AI use

Legislative phase	Green uses	Yellow uses	Red lines
Bill drafting	Synthesizing information; comparing initiatives; structuring preliminary ideas; delimiting reform scope; spelling/style support; preparing technical drafts subject to review.	Generating policy ideas without clear position; searching data without verification; using unreliable sources; outputs with possible bias or weak attention to gender.	Full automated drafting of initiatives; final legislative text without human review; AI-generated legal foundations without updated sources; humanizing AI as if it understood the problem.
Committee analysis / dictamen	Comparing versions; detecting duplications; summarizing opinions; identifying objective elements for procedural evaluation.	Impact/budget estimates without verifiable data; summaries requiring legal review; language with possible bias; modifications to a dictamen without reviewing opinions.	Full dictamen drafting without legal supervision; evaluation based only on AI output; omission of technical consultation; false/outdated data used to accelerate work.
Floor debate	Synthesizing debate points; preparing rough speech drafts for editing; retrieving supporting information; aligning draft with party ideology.	Automatic counter-arguments; prediction of other groups' positions; speeches without source verification; rhetorical appeals requiring political alignment.	Speeches fully generated by AI without review; real-time AI responses without verification; incorrect/sensitive data; racist, xenophobic, or dignity-violating content.
Communication and follow-up	Citizen summaries; press releases based on verified inputs; media monitoring; audience segmentation; adaptation of	Generic posts without institutional tone; content not aligned with group strategy; messages lacking gender/cultural sensitiv-	False notes; unverified information; automated simulation of citizen interaction; content published without review; deceptive synthetic media.

	technical material into accessible language.	ity; AI-generated images using legislator's likeness.	
--	--	---	--

Note. Table 5 synthesizes the 14 structured team forms generated collectively by the approximately 50 workshop participants. Categories were consolidated from repeated responses to reflect the team-based consensus regarding acceptable auxiliary uses (which assume adherence to proper confidentiality, data-protection, and approved-tool conditions), conditional uses requiring safeguards, and red lines involving authorship and representation.

Across the four phases, the same criterion recurred: assistance was acceptable when outputs could be checked, corrected, and subordinated to human judgment. Prohibitions appeared when AI was allowed to assume authorship, legal reasoning, parliamentary judgment, public representation, or sensitive communication outside institutional control.

The value of these rules-in-use lies in their practical clarity; their weakness lies in their lack of institutional force. Without rules, training, documentation procedures, or secure institutional tools, the boundary against prohibited delegation remains an operational consensus articulated by the advisory layer rather than an enforceable institutional standard.

6. Discussion: Informal Resilience and the Parliamentary Pacing Problem

The Mexico City evidence changes the usual sequence assumed by parliamentary AI governance. Guidance on AI in legislatures tends to imagine institutions moving systematically from strategy, rules, and training toward implementation (IPU, 2024; WFD, 2024). Legislative informatics makes a related demand by emphasizing traceability, legal consistency, reviewability, and responsibility in drafting environments (Palmirani et al., 2024). Here, practice came first. Advisors and offices encountered generative AI through ordinary workload pressures, personal access, and informal learning before the institution had stabilized rules for authorization, documentation, verification, confidentiality, or accountability. The pacing problem described by Marchant et al. (2011) therefore appears in a parliamentary form: the regulatory gap was already occupied by daily decisions about when AI could assist, when its output required verification, and when delegation should be refused.

Shadow AI alters the kind of work that moves outside institutional control. Unofficial software already raises problems of procedural control, data protection, interoperability, and institutional sovereignty. Generative AI adds a more difficult layer because it can summarize documents, classify information, compare legal frameworks, propose wording, and organize arguments. Crucially, as the findings show, even acceptable "green-light" practices like document comparison and synthesis involve the transfer of internal legislative documents to commercial platforms. Because these documents routinely contain sensitive political strategy, draft legal provisions, confidential committee

opinions, and personal data, shadow AI creates severe data-governance and cybersecurity risks. The risk therefore moves from the mere circulation of information to the production of legal and political meaning. Recent discussions of unapproved generative tools identify organizational risks around confidentiality, verification, and auditability (Silic et al., 2025; Socitm, 2025), but legislative work adds a democratic requirement because parliamentary outputs need identifiable authorship and a traceable connection between political intention and final text.

The evidence does not show formal delegation of parliamentary authority to an algorithmic system, but rather a slower relocation of artificial discretion (Young et al., 2019) through daily work. Earlier work on digital bureaucracy described a movement from street-level discretion to system-level bureaucracy (Bovens & Zouridis, 2002), while later accounts show that automation can redistribute discretion without removing human judgment (Giest & Klievink, 2024). In this legislature, discretion did not move into a centralized information system; it was distributed across advisors, offices, commercial platforms, and improvised review practices. Using the lens of Technology Enactment Theory, we see that legislative advisors do not simply accept or reject technology; they actively enact it by building practical boundaries around its use (Fountain, 2001). While it is predictable that professionals working with legally consequential text would reject unsupervised AI drafting, the systematic documentation of their methods carries theoretical significance. It demonstrates that the transition to algorithmic governance is mediated by professional risk aversion and a defense of human authorship, resulting in complex rules-in-use at the operational layer.

The analogy with street-level bureaucracy needs adjustment, but remains theoretically useful. Legislative advisors do not allocate benefits or sanctions directly to citizens (Lipsky, 1980). Their discretion is internal, cognitive, and procedural, operating between elected representatives, legal formats, committee procedures, and public communication. Under workload pressure and incomplete rules, their everyday decisions become the effective governance of institutional practice. This internal advisory layer remains difficult to observe because policy advisory systems research usually follows more visible actors and institutionalized knowledge flows (Halligan, 1995; Craft & Howlett, 2012). Generative AI enters through this less visible layer, and transforms the advisory field into a crucial site of digital governance even when formal AI strategies do not identify it as the locus of adoption.

These practical boundaries give empirical content to the three-layer model. At the operational layer, advisors created limits around AI use, justified at the normative layer by professional commitments to authorship, auditability, and document integrity. However, because the institution has not yet converted these daily practices into standards, training, or secure tools at the collective-choice layer (Ostrom, 2005), this informal resilience remains fundamentally fragile. Relying on tacit consensus rather than institutional force leaves these provisional safeguards vulnerable to staff turnover, unequal technical skills, and workload escalation. The relevant institutional risk is therefore gradual normative erosion: repeated convenience can make unverified summaries or AI-generated public messages appear routine, expanding delegation through habit without any formal authorization.

The legislative setting changes the meaning of accountability because a bill, dictamen, or public communication carries a chain of responsibility that ordinary administrative writing does not. Work on technological due process and algorithmic opacity warns that automated systems can embed errors into public decisions while making responsibility harder to assign (Citron, 2008; Fink, 2018; Busuioc, 2021). Without records of AI use, verification procedures, or approved drafting environments, a legislature may be unable to reconstruct how a generative system shaped a document even when a human actor formally signs it. Contextual interviews supported this: while AI may support preliminary research, the legislator remains accountable for the political intention, justification, and final judgment attached to a legislative act. A generative system can support preparation, but it cannot bear public responsibility for representation.

The empirical scope of this study remains limited. Survey results do not estimate AI use across the entire Mexico City Congress, and workshop forms do not provide independent individual responses from all participants. Because the sample was exploratory, targeted, and linked to workshop participation, it likely overrepresents advisors willing to engage with digital tools and institutional reflection. Its contribution lies in showing how a strategically relevant group reasons about AI use under weak formal governance, without claiming to evaluate the legal quality of AI-assisted documents.

These limits point to a robust comparative agenda. Future work should examine whether similar red lines appear in subnational legislatures and national parliaments with different levels of digital maturity. Comparative research could test whether boundaries around normative drafting remain stable across varying party organizations, staff professionalization, legal traditions, and access to secure institutional tools. Document-level research should also compare AI-assisted drafts with human-only drafts to objectively assess effects on legal language, citation quality, and political positioning.

7. Conclusion

This article has advanced the argument that the rapid integration of generative AI into legislative advisory work has created an "adoption-rules gap," where everyday administrative practices serve as provisional governance in the absence of formal institutional rules. By examining the Mexico City Congress, we demonstrated how frontline advisors actively manage "shadow AI" by articulating informal boundaries against the unsupervised delegation of normative text, legal reasoning, and political judgment.

For legislatures, the institutional task is to convert frontline knowledge into durable governance. The adoption-rules gap in the Mexico City Congress is already filled with informal decisions about acceptable assistance, conditional use, and prohibited delegation. While these enacted rules-in-use mitigate short-term risks, they cannot substitute for formal organizational readiness. To stabilize these informal safeguards, legislatures must move from operational coping mechanisms to collective-choice codification by implementing several concrete recommendations.

First, institutions should procure and deploy secure, internal generative AI environments that prevent sensitive legislative data from entering public commercial models. Second, they must establish clear confidentiality rules defining what types of data (e.g., draft legal provisions, personal data, and confidential committee opinions) are strictly prohibited from being processed by unapproved AI systems. Third, transparent documentation routines for AI-assisted drafting must be required to ensure that the use of AI in legislative synthesis or comparison leaves a reviewable institutional trace. Additionally, legislatures should mandate human review and staff training by instituting continuous technical instruction focused on ethics, fact-checking, and the necessity of mandatory human supervision over all AI outputs. Finally, it is crucial to codify explicit red-line prohibitions that formally adopt the boundaries identified by advisors, explicitly forbidding the unsupervised delegation of final normative text, legal reasoning, political representation, and public communication to AI systems.

Legislatures must translate frontline decisions into secure tools and accountability procedures before convenience, staff turnover, and opaque systems permanently weaken human authorship, reviewability, and public responsibility in democratic lawmaking.

Author Contributions

S.A.B.J. and E.C.-G. contributed equally to all aspects of this manuscript, including conceptualization, methodology, formal analysis, data curation, writing, review, and editing. All authors have read and agreed to the published version of the manuscript.

About the Authors

Sergio A. Bárcena Juárez

Sergio A. Bárcena is a full-time professor and researcher at the School of Humanities and Education of Mexico's Tecnológico de Monterrey. He specializes in legislative politics, electoral behavior, political representation, and data-driven analysis of democratic institutions. His work combines political science, public policy, and computational approaches to study legislative performance, civic participation, and institutional change in Mexico and Latin America.

Erik Cardona-Gómez

Erik Cardona-Gómez is a Teaching Fellow in the Department of Politics and International Studies at SOAS, University of London. He earned his Ph.D. in Political Theory from the University of York, where his research focused on multiculturalism and Indigenous rights. His current scholarship explores a wide range of critical issues, including structural injustice, settler colonialism, and the algorithmic metamorphosis of political representation.

References

- Bovens, M., & Zouridis, S. (2002). From street-level to system-level bureaucracies: How information and communication technology is transforming administrative discretion and constitutional control. *Public Administration Review*, 62(2), 174–184. DOI: 10.1111/1467-9299.00300

- Bright, J., Enock, F. E., Esnaashari, S., Francis, J., Hashem, Y., & Morgan, D. (2025). Generative AI is already widespread in the public sector: evidence from a survey of UK public sector professionals. *Digital Government: Research and Practice*, 6(1), 1–13. DOI: 10.1145/3700140
- Busuioc, M. (2021). Accountable artificial intelligence: Holding algorithms to account. *Public Administration Review*, 81(5), 825–836. DOI: 10.1111/puar.13293
- Citron, D. K. (2008). Technological due process. *Washington University Law Review*, 85, 1249–1313.
- Craft, J., & Howlett, M. (2012). Policy formulation, governance shifts and policy influence: Location and content in policy advisory systems. *Journal of Public Policy*, 32(2), 79–98. DOI: 10.1017/S0143814X12000049
- de Boer, N., & Raaphorst, N. (2023). Automation and discretion: Explaining the effect of automation on how street-level bureaucrats enforce. *Public Management Review*, 25(1), 42–62. DOI: 10.1080/14719037.2021.1937684
- Economic Commission for Latin America and the Caribbean, & National Center for Artificial Intelligence. (2024). *Latin American Artificial Intelligence Index (ILIA) 2024*. ECLAC and CENIA.
- Ehn, P. (2008). Participation in design things. In *Proceedings of the Tenth Anniversary Conference on Participatory Design 2008* (pp. 92–101). DOI: 10.1145/1795234.1795248
- Fink, K. (2018). Opening the government's black boxes: Freedom of information and algorithmic accountability. *Information, Communication & Society*, 21(10), 1453–1471. DOI: 10.1080/1369118X.2017.1338142
- Fountain, J. E. (2001). *Building the virtual state: Information technology and institutional change*. Brookings Institution Press.
- Giest, S. N., & Klievink, B. (2024). More than a digital system: How AI is changing the role of bureaucrats in different organizational contexts. *Public Management Review*, 26(2), 379–398. DOI: 10.1080/14719037.2022.2095001
- Gómez Macfarland, C. A. (2025). La inteligencia artificial al servicio del Poder Legislativo en México: Acciones para la capacitación del personal de las cámaras del Congreso de la Unión en la materia. *RIESED - Revista Internacional de Estudios sobre Sistemas Educativos*, 3(16), 781–802. DOI: 10.5281/zenodo.15697666
- Halligan, J. (1995). Policy advice and the public service. In B. G. Peters & D. J. Savoie (Eds.), *Governance in a changing environment* (pp. 138–172). McGill-Queen's University Press.
- INEGI. (2024). *Censo Nacional de Poderes Legislativos Estatales (CNPLE) 2024*. Instituto Nacional de Estadística y Geografía.
- Inter-American Development Bank. (2024). *Accelerating AI adoption in Latin America and the Caribbean: Aligning regulations, deploying pilots, and conducting training (Project RG-T4439)*. Inter-American Development Bank.
- Inter-Parliamentary Union. (2024). *Guidelines for AI in parliaments*. Centre for Innovation in Parliament.
- Lipsky, M. (1980). *Street-level bureaucracy: Dilemmas of the individual in public services*. Russell Sage Foundation.

- Marchant, G. E., Allenby, B. R., & Herkert, J. R. (Eds.). (2011). The growing gap between emerging technologies and legal-ethical oversight: The pacing problem. Springer. DOI: 10.1007/978-94-007-1356-7
- OECD. (2024). Governing with artificial intelligence: Are governments ready? OECD Publishing. DOI: 10.1787/26324bc2-en
- Orlikowski, W. J. (2000). Using technology and constituting structures: A practice lens for studying technology in organizations. *Organization Science*, 11(4), 404–428. DOI: 10.1287/orsc.11.4.404.14600
- Ostrom, E. (2005). Understanding institutional diversity. Princeton University Press.
- Palmirani, M., Bresciani, P. F., Bomprezzi, C., & Gatti, F. (2024). Report on AI in parliamentary context: Final draft, December 2024. University of Bologna, HyperModeLex Project.
- Silic, M., Silic, D., & Kind-Trüller, K. (2025). From shadow IT to shadow AI: Threats, risks and opportunities for organizations. *Strategic Change*. DOI: 10.1002/jsc.2682
- Socitm. (2025). Shadow AI in the public sector: Innovation without oversight? Society for Innovation, Technology and Modernisation. <https://socitm.net/resource-hub/blog/shadow-ai-in-the-public-sector-innovation-without-oversight/>
- Tangi, L., Rodriguez Müller, P., & Combetto, M. (2025). A Silent Partner: The Shadow Presence of Generative Artificial Intelligence in Public Administrations. In *Electronic Participation. ePart 2025 (Lecture Notes in Computer Science, Vol. 15978, pp. 69–86)*. Springer. DOI: 10.1007/978-3-032-02515-9_5
- Westminster Foundation for Democracy. (2024). Guidelines for AI in parliaments. Westminster Foundation for Democracy.
- Young, M. M., Bullock, J. B., & Lecy, J. D. (2019). Artificial discretion as a tool of governance: A framework for understanding the impact of artificial intelligence on public administration. *Perspectives on Public Management and Governance*, 2(4), 301–313. DOI: 10.1093/ppmgov/gvz014

Appendix A: Methodological Instruments

This appendix contains the supplementary methodological tools. The narrative analysis of this data is located in Section 4 (Methodology and Research Design) and Section 5 (Empirical Findings) of the main manuscript.

A.1. Phase 1: Diagnostic Survey Instrument (N=42)

Target Population: Active legislative advisors.

Data Source: Original instrument deployed through digital professional networks and coordinated through the workshop's logistical channels.

Note on Translation: The following survey instrument was originally administered in Spanish. The items presented here have been translated into English by the authors for the purpose of this manuscript.

Part I: Demographics and Baseline Infrastructure

What is your age range?

What is your current position?

How many years of cumulative experience do you have in parliamentary advisory roles?

Indicate your level of agreement with the following statements regarding the technological resources provided by the Chamber:

The internet connection in my office is fast and stable for performing my tasks.

The IT security infrastructure is sufficient to mitigate risks when using external tools.

The computer equipment assigned to me is suitable for using advanced digital tools.

The Technical Support area provides the necessary assistance to resolve problems.

Part II: AI Adoption and Digital Competencies

Do you currently use Artificial Intelligence (AI) tools in your tasks within the congress?

Since when have you been using AI tools in your work?

Which artificial intelligence or automation tools have you used in your work?

How have you learned to use Artificial Intelligence tools for your work?

How skilled do you feel in using AI tools in your daily work?

Evaluate your level in each area of AI use:

Searching for and verifying information with AI support.

Drafting or summarizing documents with AI.

Writing speeches and interventions with AI (arguments, rhetoric).

Part III: Risk Perceptions and Epistemic Threats

Which risks do you consider most concerning when using Artificial Intelligence in your work?
(Select up to 3)

In which essential tasks of your work would you not consider AI useful?

How likely do you consider it that the use of AI will put your job at risk in the next 5 years?

Do you believe that the implementation of AI can raise the technical quality of legislative work?

Part IV: Governance and Institutional Preferences

How important is it that the following elements exist to guarantee a responsible and ethical use of AI in Congress?

Mandatory human supervision of the final AI output.

Creation of a centralized AI documentation and analysis unit.

Strict information security and confidentiality protocols.

Rules or agreements that define valid uses within the chamber.

Continuous technical training focused on ethics and verification (fact-checking).

Appendix B. Co-Design Workshop Coding Matrix (The 'Traffic Light' Protocol)

Sample: Approximately 50 workshop participants generating 14 consolidated team forms. (Note: This participant group partially overlaps with the survey respondents in Appendix A, as coordinated through the event's logistical channels).

Objective: Categorize acceptable ("rules-in-use") and unacceptable ("red lines") AI applications across the legislative lifecycle based on qualitative consensus.

(Note: Text in the matrix is derived verbatim from the raw participatory data).

Table B.1: Qualitative Coding Matrix

Legislative Phase	Desirable Practices (Green Light)	Neutral or Doubtful Practices (Yellow Light)	Undesirable Practices / Red Lines (Red Light)
1. Drafting of initiatives (proposals)	"Analyze and synthesize information"; "Knowledge of the initiative and pros/cons regarding the research."	"Structure, background, and ideas for initiatives"; risk of using "unreliable sources."	"Having the AI perform all the work without human supervision"; "Grounding and data collection" (without verification).
2. Committee analysis / legal opinion	"Searching for document duplication and structure analysis"; "Gaining greater scope on a topic."	"Budgetary impact, summaries, and synthesis."	"Grounding and preparation of the entire legal opinion"; "Ceasing to analyze the situation due to a generalized response."
3. Plenary / debate / speeches	"Fast drafting of speeches, rebuttal, and voting reasoning"; "Better tools to focus on the problem."	"Explanation of motives and extension of speeches"; "Supplanting human understanding."	"Expression of the speech and data updating" (hallucination risk); "Setting aside the feeling of the specific problem."
4. Communication and follow-up	"Better structure when writing"; "Fast communication and projected segment."	Risks of "gender bias and cultural representation" in generated content.	"Using the voice or photographs of deputies"

			ties for misuse" (deep-fakes); "Limited veracity."
--	--	--	--

Appendix C. Institutional Context: State Legislative Comparison

Table C.1: National Census of State Legislative Powers (INEGI, 2024) - Selected Raw Data

State Legisla- ture	Total Mem- bers (Legislators)	Staff per Leg- islator (Ratio)	Computers per Legislator (Ratio)	Reported Tech Budget
Mexico City	66	ND	ND	ND
Jalisco	38	> 25.0	~ 17.0	Reported
Michoacán	40	> 25.0	~ 8.0	Reported
State of Mex- ico	75	~ 15.0	~ 17.0	Reported
Puebla	41	< 10.0	~ 5.0	Reported
National Av- erage	N/A	10.5	7.8	N/A

Source: 2024 National Census of State Legislative Powers (INEGI). Note: This table highlights the structural reporting anomaly of the Mexico City Congress compared to both high-capacity and low-capacity state legislatures regarding administrative traceability.