
From Proposals to Enactment: The Procedural Bottleneck in AI Safety Regulation

Mansur Ali Khan
Skyline High School
khanman26@issaquah.wednet.edu

Mehmet Efe Akengin
Fatima Institute for Global AI Research
mefeakengin@alum.mit.edu

Ahmad Rushdi
Stanford Institute for Human-Centered AI
rushdi@stanford.edu

Abstract

1 While AI models advance at unprecedented rates, AI safety legislation remains
2 largely symbolic, stalled, or unrealized. Through a year-by-year analysis of AI
3 breakthroughs, U.S. congressional policy proposals, and international legislative
4 enactments, this study identifies a structural gap: the United States is not deficient
5 in AI safety bill proposals but in legislative action, with only 4.23% of U.S. AI
6 bills reaching any terminal outcome. We quantify enactment rates, map U.S. Con-
7 gressional AI bills across thematic domains, identify procedural bottlenecks, and
8 develop a logistic regression model to test which factors predict legislative stalling.
9 This study contributes five key advances: (1) a quantitative comparison of AI
10 legislation versus LLM breakthroughs, (2) a comprehensive taxonomy of proposed
11 and enacted policy subfields, (3) a dataset elucidating the structural causes of AI
12 legislation failure, (4) statistically significant evidence that number of sponsors neg-
13 atively affect bills' progress, and (5) policy recommendations grounded in planned
14 adaptation, preemptive enactment, and independent AI oversight. We demonstrate
15 that without enactment, AI safety regulation remains inert, highlighting the urgent
16 need for actionable, coalition-backed AI safety policies in the United States.

17 1 Introduction

18 Artificial Intelligence (AI) has shifted from a niche research domain into a widespread force shaping
19 economies, politics, and individual identities [1, 2]. This rapid expansion raises urgent questions
20 about AI's societal impact and whether governance structures are prepared to respond. The United
21 States has historically led technological innovation. However, in AI governance, legislative progress
22 lags far behind the pace of technological change [3, 4]. The U.S. remains active in global AI policy
23 discussions, but many AI-related bills fail or stall in Congress. This suggests that the barrier is not
24 awareness, but institutional inertia. The challenge is therefore not just to anticipate AI's influence,
25 but rather it is to overcome gridlock and enact substantive regulation [4]. Our analysis reveals this
26 gridlock operates primarily through procedural mechanisms rather than substantive disagreement
27 over policy content.

28 The European Union's AI Act has been praised as a landmark policy establishing formal obligations
29 for high-risk systems. While earlier drafts of the associated Code of Practice relied heavily on
30 provider self-assessment, more recent versions introduce elements of external oversight and third-
31 party evaluation. Nonetheless, scholars continue to note that the absence of individual redress and
32 the deferral of key enforcement mechanisms to future technical standards may limit its overall

regulatory force [6, 7, 8]. Across the Asia-Pacific region, most jurisdictions, including Japan, India, and Australia, have issued principle-based guidelines emphasizing fairness and transparency, but these remain voluntary and lack binding enforcement, while South Korea and Taiwan are still drafting national legislation [9]. China stands as a notable exception, having enacted enforceable measures such as the 2023 Interim Measures for GenAI Services, mandating licensing, security reviews, and real-name user registration. In contrast, the United States-despite being the world’s leading developer of frontier AI models, has yet to pass comprehensive federal AI safety legislation. High-profile proposals, including the AI Bill of Rights and bills to establish a federal AI Safety review office, have lapsed [10, 11, 12, 13]. However, this does not equate to arguing Congress is passive on AI safety. In 2025, the bill H.R. 1, known as "The One Big Beautiful Bill Act", was enacted on the 4th day of July. Its original version included a provision impacting AI safety by enacting a 10-year moratorium on states and localities from regulating AI models, systems, or automated decision systems. This moratorium was removed in the Senate by a 99 - 1 vote. Illustrating the recognition of regulated AI. State-level initiatives have emerged, for instance Washington State has convened an AI Task Force to advise legislators [14], California has pursued bills like SB 1047 and AB 2013 to mandate safety protocols and transparency [15], and California courts have begun shaping precedent in cases such as *Kadrey v. Meta*, which established that developers may face regulation when their systems disrupt economic markets [16]. In the 2024 election cycle, more than 151 state bills targeted deepfakes and deceptive media, demonstrating that in moments of acute perceived risk, U.S. states can act with speed and breadth [17]. Furthermore, outside of the legal jurisdiction, the 47th President of the United States declared the "Winning the Race AMERICA’S AI ACTION PLAN" [12] [13]. A component of the plan is to urge for more open source models to encourage further developments. A side benefit of this is that open source models increase transparency from frontier model developers. This method has proven to be effective, as a week after the announcement, OpenAI released two open-source models, gpt-oss-120b and gpt-oss-20b [18].

2 Related Work

2.1 Legislation Trackers

The 2025 Human-Centered AI (HAI) Artificial Intelligence Index Report from Stanford University is a comprehensive report detailing AI research and development, technical performance, and responsible AI. It also explores the role of AI in education, the economy, science, medicine, and public opinion. Lastly and most relevant to this paper, the index report tracked AI policy globally. The report showed the focus of U.S. states on deepfake legislation, a rise in policy proposals in the 2024 election season, and various other significant findings by using comprehensive data sets. It also highlighted the significant discrepancy between the U.S. State and Federal governments’ enactment and proposal of AI regulations. However, due to the nature of the report, the conclusive reasoning for why this gap exists is missing. Additionally, the Index Report shows through multiple data sets that the United States Congress is increasingly focusing on AI policies, at a near-exponential rate, by referencing mentions of AI in bills, committees, speeches, and agencies. However, it lacks a breakdown of what sectors of AI (ethical usage, LLMs, AGI, agentic-AI, etc.) are being referenced and focused upon [17].

The Brennan Center’s Artificial Intelligence Legislation Tracker compiles and maintains a comprehensive repository of AI-related bills introduced in the U.S. Congress, specifically, those referencing “artificial intelligence” across the 118th and 119th sessions only. This resource helps industry experts, advocates, and the public understand the legislative attention being paid to AI and the risks identified by lawmakers, such as election integrity, misinformation, and surveillance. It offers transparency into how federal lawmakers propose to regulate AI through bills and executive actions. The data set is limited in the aspect that it is not comprehensive, not providing a full overview of AI legislation [19].

The UK’s International AI Safety Report 2025, chaired by Yoshua Bengio and authored by a global panel, represents the first major synthesis of scientific knowledge about the capabilities and risks of advanced AI systems. Rather than offering policy prescriptions, it compiles scientific evidence around the potential harms from deepfakes and cyberattacks to job displacement and biological threats to inform policymakers and build a shared fact base across nations [20].

85 2.2 Analysis of policy and method of regulation

86 In Comparative Global AI Regulation: Policy Perspectives from the EU, China, and the U.S., Chun,
87 de Witt, and Elkins analyze how different jurisdictions are approaching AI governance. The study
88 examines regulatory philosophy across regimes, risk-based frameworks like the EU AI Act, the
89 market-oriented U.S. model, and China’s centralized approach, while probing the fragmentation
90 between federal and state-level efforts in America, notably through the lens of California’s pending
91 (at the time of the production of the report) SB 1047. This comparative perspective highlights how
92 cultural, political, and institutional differences shape AI policy direction [4].

93 Regarding effective methods of regulation, a 2009 study investigates the use of dynamic legislation
94 that morphs based on changing situations, termed as planned adaptation. It is concluded that planned
95 adaptation, at a minimum, improved policy and should be implemented thoroughly in the United
96 States and perhaps beyond. Planned adaptation in AI regulation can be used to make sure that an
97 ever-changing field doesn’t grow out of established legislation [5].

98 Multiple sources break down the AI bills into sub-fields and their endpoint. However, they are limited
99 to the years 2023-2025. The first bill referencing AI directly was in 2017. In this paper, we cover
100 the gap in these legislation trackers to provide the first comprehensive deduction of sub-fields and
101 endpoints for public access [19, 21, 22, 23].

102 3 Methodology

103 This paper employs quantitative legislative analysis and a qualitative policy evaluation to examine
104 the U.S. legislative bottlenecks in AI safety regulation, while also contrasting international devel-
105 opments. The study focuses on identifying not only the volume of proposed and enacted laws but
106 also the procedural bottlenecks that prevent legislative action, primarily within the United States. All
107 code, scripts, results, and data used in the methodology are made available at [www.github.com/
108 MansurAKhan/The-Procedural-Bottleneck-in-AI-Safety-Regulation](https://www.github.com/MansurAKhan/The-Procedural-Bottleneck-in-AI-Safety-Regulation)

109 3.1 Data Sources and Compilation

110 To construct a comparison between AI development and AI regulation, this paper aggregates data
111 from several verified public and governmental sources. Unlike prior trackers, we provide the first
112 comprehensive dataset spanning 2017– August 2025 with bill sub-fields, endpoints, and modeled
113 bottleneck factors:

- 114 • **AI Breakthroughs:** Major large language model (LLM) releases from 2017 to mid-2025
115 were compiled using the *LLM Timeline Project*, Wikipedia’s *List of Large Language Models*,
116 and official publication announcements from leading developers such as OpenAI, DeepSeek,
117 Mistral, DeepMind, and Anthropic.
- 118 • **U.S. Legislation:** AI-related bills proposed at the federal level were retrieved from the offi-
119 cial U.S. Congress database (www.congress.gov)¹, filtered using the keywords “Artificial
120 Intelligence” and “AI”. Additional records were sourced from the National Conference of
121 State Legislatures (NCSL) and the Brennan Center for Justice. Each bill was categorized
122 into sub-fields. This is the first comprehensive categorization of all AI bills into sub-fields.
123 The sub-fields were formulated by a human subject matter expert and are verified by a lawyer.
124 The classification of the sub-fields was done by providing the sub-fields and the URLs of
125 all 150 bills for GPT-4o to output any sub-fields that correspond to each bill, and provide
126 a confidence level (Low, Medium, High). Bills with a low confidence rate were audited
127 and corrected manually. The accuracy of the labeled data was determined by selecting 50
128 classifications at random and verifying the accuracy through manual auditing. The accuracy
129 was found to be 94% (47 out of 50).

¹The HAI Index Report 2025 contains a dataset covering proposed bills, amendments, and simi-
lar legislative action in the United States, similar to the "US AI Laws Proposed" dataset count (view
Appendix Table 3). The methods applied for analysis in this paper are only possible on bills; thus,
we created a separate analytical dataset of the bills, available at [www.github.com/MansurAKhan/
The-Procedural-Bottleneck-in-AI-Safety-Regulation](https://www.github.com/MansurAKhan/The-Procedural-Bottleneck-in-AI-Safety-Regulation).

Through comprehensive manual human annotation, each bill was assigned an endpoint and a systematic rationale for reaching this point. This is the first comprehensive analysis of endpoints of AI bills. The following is a breakdown of each ending category:

- **Expired without action:** This label is attributed to the bills that have only one action regarding them, be it a referral to a committee or a hearing. If the bill does not step through the introduction, it is marked as "Expired without action."
- **Stalled in Committee (House/Senate):** This label is attributed to the bills that have multiple actions regarding them, and their last action is to be referred to a Committee, which can happen in both the Senate and the House.
- **Senate/House Calendar Inaction:** This label is attributed to the bills that have multiple actions regarding them, referred to and passed a committee, and their last action is to be placed on Calendar, which can happen in both the Senate and the House. This process can occur multiple times, but none of the failed bills have reached this stage.
- **No Action After Introduction (note):** This label is attributed to the bills that are from the currently occurring 119th Session of Congress. In brackets, a note is provided to them to show their latest action, not a category.
- **Declined:** This label is attributed to the bills that reached an end decision not to be enacted.
- **Passed:** This label is attributed to the bills that reached an end decision to be enacted.
- **Amendment Passed:** This label is attributed to the amendments of bills that reached a decision to be enacted.
- **International Policy:** Global enactment and policy activity were gathered using the European Council's AI legislation tracker, the OECD's AI Policy Observatory, and Asia-Pacific legal reports (e.g., *InsideGlobalTech*). International summits, declarations, and non-legislative initiatives were cross-verified through news archives and official government publications.

3.2 Policy Classification Criteria

Each policy item was classified into one of the following categories:

- **Proposed Legislation:** Any AI-related bill formally introduced into a national or supranational legislature.
- **Enacted Legislation:** Policies that successfully passed both legislative branches and came into legal effect.
- **Non-Legal Action:** Includes summits, executive orders, AI task forces, public safety frameworks, and ethical guidelines that lack binding authority.
- **Failed Bill:** A bill that did not reach a final decision but failed in the process of getting to one. If a bill did not pass, it is a failed bill.

U.S. legislative outcomes were further coded into paths: stalled in committee, calendar inaction, declined, passed, or no action after introduction. These were used to calculate the **Action Rate**:

$$\text{Action Rate} = \frac{\text{Passed Bills} + \text{Failed Bills}}{\text{Total Proposed Bills}}$$

This metric serves as a representative for congressional engagement and legislative momentum.

3.3 Analytical Framework

Trends were analyzed year-by-year from 2017 through July 2025. Visualizations (e.g., Sankey diagrams, bar charts) were created to represent the results through a visual medium. Policy bottlenecks were interpreted using committee records and external legislative studies, such as ProQuest Congressional Insight and the Stanford HAI AI Index Report (2025).

3.4 Deduction of Legislative Factors Attributed to Stalling

To deduce the factors that affect a bill's capacity to stall, the human-annotated data was expanded through an API key from <https://www.congress.gov/help/using-data-offsite>. The Congress.gov API specifically allows users to view, retrieve, and reuse the machine-readable data provided. The following parameters are utilized: Chamber, Sponsor Party, Bipartisan, number of Sponsors, and the sub-fields to predict whether the bill will stall, totaling 12 parameters. A penalized logistic regression with ridge penalty model [24] was trained and run on Google Cloud using a virtualized Intel Xeon CPU @ 2.20GHz. The results were further analyzed and interpreted using Python libraries.

Out of 150 bills, 124 have reached their end. Meaning the session in which the bill belongs has expired, and thus cannot have further action taken upon it. Since the other 26 bills in the 119th session of Congress have not reached an end, while their current placement of being stalled in Congress can be attributed to, the collective reasons and factors contributing to stallation cannot be deduced; thus, these bills are removed from the data used for the Logistic Regression model.

The Logistic Regression model provided each attribute with an associated coefficient representing its effect on the log-odds of a bill being stalled. A positive coefficient represents the likelihood of a bill passing, a negative coefficient represents the improbability of a bill stalling. To improve model stability, several low-frequency sub-fields (LLM, AGI, and Autonomous Driving) were combined into a broader Advanced AI category. Individually, these sub-fields had only one or two bills each, which produced highly volatile coefficients and p-values; grouping them allowed for cleaner, more interpretable results. To view accuracy reports, see Appendix Table 5.

The penalized Logistic Regression model was implemented using scikit-learn with ridge (L2) penalty and the following hyper-parameters: regularization parameter $C = 1.0$ (inverse regularization strength), maximum iterations = 100, solver = "lbfgs", and class_weight = "balanced" to handle class imbalance. Bootstrap resampling with a maximum of 100 iterations was employed while splitting the dataset into 80% train and 20% into test, and was used to estimate standard errors and p-values for coefficient significance testing. Feature scaling was performed using StandardScaler (z-score normalization). Statistical significance was assessed at $\alpha = 0.05$, and coefficients with $|z\text{-score}| > 1.96$ were considered meaningful predictors. Model performance was evaluated using accuracy, precision, recall, and a confusion matrix, with the model achieving 62.2% training accuracy and 76.0% test accuracy after 12 iterations at convergence.

4 Results

4.1 Gap in Proposed vs Enacted Legislation

The data in Figure 1 shows that since the release of ChatGPT in 2022, more than 55 new LLMs have been released. In 2023, 19 were released; in 2024, 20 were released, and by August 2025, 19. Which includes tools like GPT-5, Deep Research, and the ChatGPT Agent. Yet during the same time frame, only a handful of AI-related laws were enacted worldwide. In 2023, 13 policies were enacted globally; that number decreased to 6 in 2024, and in 2025, only one policy was enacted. Model development grows in a near-exponential pattern, and legislative enactment remains sub-linear. Two additional observations can be deduced from the data. First, in 2024, Non-legal AI Safety Actions decreased by more than 50% compared to previous years, and new AI bill proposals in the U.S. doubled in the same time frame. This may suggest policymakers prioritized substantive action over symbolic measures, though election-year dynamics could also explain the increase. This rise in law proposals, however, had minimal outcomes. In 2023, 13 policies were enacted globally alongside 28 U.S. proposals. The following year, proposals nearly doubled to 59, while enactments fell to just 6. Further exemplifying the claim that the true gap is not in making laws, but enacting them.

Lastly, of the 150 proposed bills in Congress, only three were enacted: the AI Training Act (2021), the James M. Inhofe National Defense Authorization Act for Fiscal Year 2023 (2022), and the One Big Beautiful Act (2025). None of them explicitly works to ensure AI safety, highlighting the gap in AI safety regulation despite the existence of AI legislation. This gap is further explored through the categorization of bills into sub-fields in the next subsection. To view a full breakdown of the data in Figure 1, see Appendix Table 3.

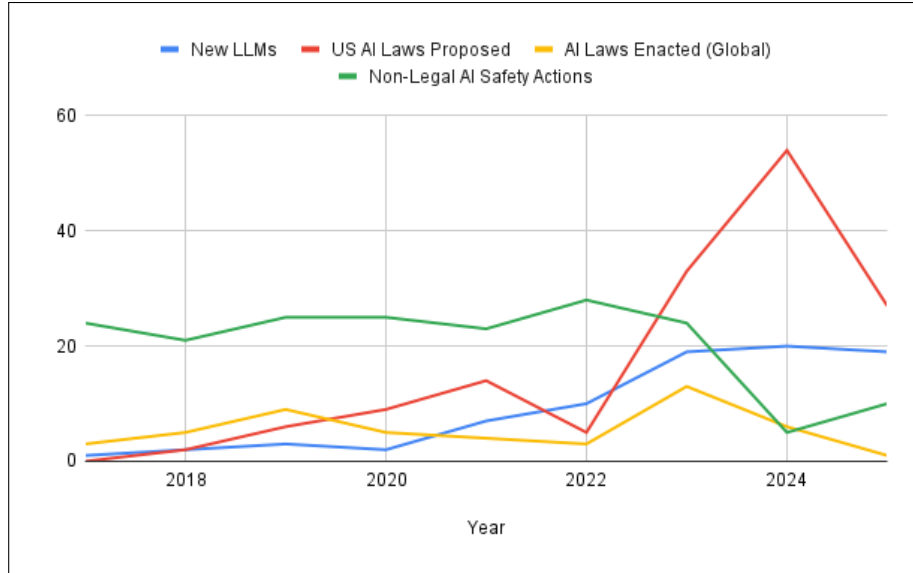


Figure 1: Trajectory of LLM Breakthroughs vs. Regulation for AI Safety

4.2 Sub-fields of Proposed Legislation in the U.S.

While laws are actively being proposed and discussed in the United States, and the immense low level of enactment of these laws is evident, we next break down the areas of their focus. Figure 2 visualizes the sub-fields in focus across the United States Congressional policy proposals.

Figure 2 shows that since 2019, the bill proposals focusing on the sub-fields: Data Usage, Policy Advisory, and General Ethical Usage (GEU) have been majority. In 2022, these three sub-fields covered 66.7% of the proposed bills, highlighting that the bills proposed in Congress address AI safety.

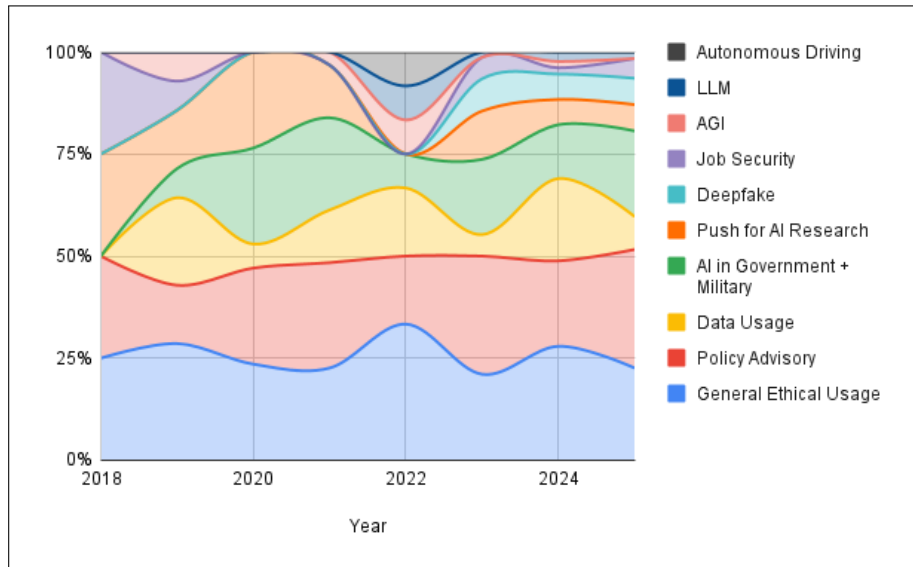


Figure 2: Sub-fields of Proposed Legislation in the U.S

In 2022, the only bills related to Large Language Models (LLMs) and Autonomous Driving were introduced. That same year, proposed bills concerning Artificial General Intelligence (AGI), LLMs, and Autonomous Driving collectively reached their all-time high at 24.9%. At the same time, bills

that pushed for AI research dropped to 0% in 2022. This marked a sharp decline, as in 2021 about 12.9% of bills supported AI research, and in 2020 the share was even higher at 23.5%. Since this decrease, the field of AI research has not regained its earlier recognition. The highest share it has seen since 2022 was only 11.8%.

Another important change appeared around deepfakes. Before 2023, no bills specifically addressed them. By 2025, however, H.R. 1 included deepfake legislation, and it was passed. Looking back before 2022, the sub-fields most focused on AI development (Push for AI Research and AI in Government + Military) made up a significant portion of proposed bills. After 2022, both categories declined, and lawmakers instead introduced more safety-related bills. Taken together, these patterns show that 2022 marked a turning point; Legislative priorities shifted away from AI development and towards creating safeguards.

Of all these bills, the following sub-fields have multiple enacted laws: AI in Government + Military, Policy Advisory, and Push for AI. While General Ethical Usage and Deepfake have only one bill. The bill that focuses on the General Ethical Usage of AI, only lightly ensures that AI is used ethically for military domains. To view a breakdown of the data in Figure 2, see Appendix Table 4

To find out why the bills being proposed focus on AI safety, yet none get enacted, we delved deeper into the paths of these bills to find out the reasons.

4.3 Reasons for failed legislation

We next analyze reasons for failed legislation. To quantify the effectiveness of Congress and its actions on AI policy, an action rate is necessary that accurately achieves said purpose. In this paper, the action rate was derived through the formula:

$$\text{Action Rate} = \frac{P + F}{T} \quad (1)$$

Table 1² displays that 89 of the 150 proposed legislation were stalled in a committee after some action. Only 4 bills reached an end, to pass or to fail. Being a new area, the action rate should have been higher; however, the action rate is 4.23%. 2.02% less than the national average from the same time frame, 6.25%. Based on the data from GovTrack US in the same timeframe as when AI policy legislation has been proposed, 2017-2024. Highlighting in greater detail that enactment is extremely low³.

Table 1: Final Destination of Proposed U.S. AI Legislation (2018–2025)

Year	Stalled in Comm.	No Action	Calendar Inact.	Failed Vote	Expired	Passed	Total	Action Rate	Action Rate %
2018	2	0	0	0	0	0	2	0	0%
2019	3	0	2	0	1	0	6	0	0%
2020	1 [#]	0	0	0	6	2 [*]	9	0	0%
2021	10 [#]	0	1	0	2 [#]	1	14	0.071	7.14%
2022	2	0	1	0	1	1	5	0.200	20%
2023	24 [#]	0	7	1	1	0	33	0.030	3.03%
2024	47	0	2	0	3	2 [*]	54	0	0%
2025	0	26	0	0	1	1	27	0.037	3.70%
Total	89	26	13	1	14	8	150	–	–

Notes: Average Action Rate: 0.0423 (4.23%).

In 2023 and 2024, when the largest number of proposals were submitted (33 and 54, respectively), none were enacted. While in the years 2021, 2022, and 2025, fewer proposals were made (14, 5, and 27, respectively), one bill was enacted each year. This finding suggests that in AI governance, fewer but more comprehensive and broadly supported bills have a higher chance of enactment than a large volume of symbolic proposals. In practice, this means legislative quality (coalition-building, clarity, and scope) matters more than quantity for advancing AI safety policy.

^{2#}Includes a bill that passed one chamber. ^{*}Includes 2 amendments to bills that did pass.

³Bills that passed one chamber or passed as amendments are not reflected in the enactment totals.

The majority of the bills proposed, on average, 59.3% were stalled in committees, see Figure 3, commonly related to fields of technology, ethics, or the economy (breakdown of data can be accessed at www.github.com/MansurAKhan/The-Procedural-Bottleneck-in-AI-Safety-Regulation). AI regulation policies are being sent to subcommittees, where they stall, are placed on a calendar, and ignored. In the political realm, this phenomenon is called pigeonholing. A large reason for this is that in committees and subcommittees, “if the leadership decides the bill does not fit within its overall agenda, a decision not to act will ‘kill’ the bill just as effectively as a vote against it” [26]. This illustrates the need for an independent committee that deals with the issues of AI safety and its regulation.

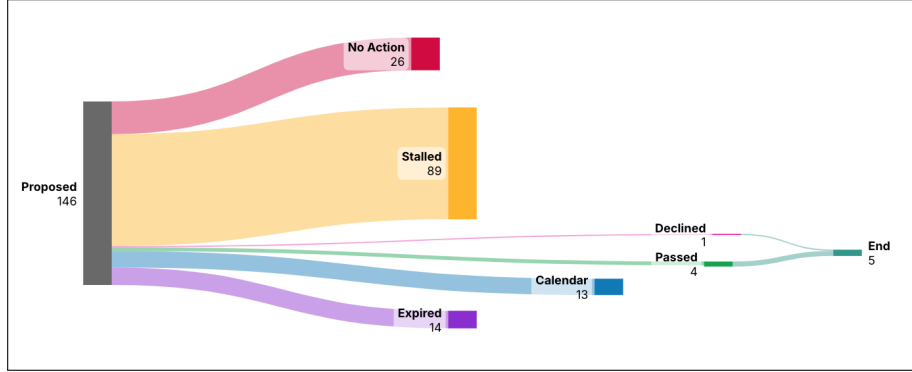


Figure 3: Path of Bills by Volume to Visualize the Last Stages of Each bill

4.4 Further factors contributing to stalling

Table 2: Logistic Regression Coefficients with Sub-fields Bolded

Feature	Coefficient	Std. Error	Z value	P-value	Coefficient
Advanced AI	-0.7662	0.7939	-0.9652	0.3345	0.7662
AI in Government + Military	-0.5247	0.5283	-0.9931	0.3207	0.5247
Job Security	-0.3767	0.5593	-0.6736	0.5006	0.3767
Push for AI Research	0.3405	0.5380	0.6329	0.5268	0.3405
Policy Advisory	0.2048	0.5763	0.3554	0.7223	0.2048
General Ethical Usage	0.2292	0.5264	0.4353	0.6633	0.2292
Data Usage	-0.1252	0.5774	-0.2168	0.8283	0.1252
Deepfake	0.2187	0.6908	0.3166	0.7516	0.2187
Bipartisan	-0.3282	0.5748	-0.5710	0.5680	0.3282
Chamber_Binary	0.7971	0.5817	1.3703	0.1706	0.7971
Num_Sponsors	0.8068	0.3649	2.2113	*0.0270	0.8068
Sponsor_Party_Binary	0.2350	0.5662	0.4151	0.6781	0.2350

Notes: * $p < 0.05$.

We next analyze the factors that contribute most to a bill being stalled beyond the nature of a committee to propose a solution to the enactment gap effectively.

Table 2 displays the coefficients that indicate the impact of each subfield on stalling. The bolded subfields all have p-values well above 0.05, meaning that with the available data, we cannot establish a statistically significant relationship between a bill’s subject matter and whether it stalls. This does not imply that subject matter plays no role; rather, we cannot confirm a causal link based on this analysis.

By contrast, structural and political factors appear to have a stronger impact. The most significant predictor is the number of sponsors of a bill, with a coefficient of 0.8068 and a p-value of 0.0270, which is below the 0.05 significance threshold. This suggests that bills with more cosponsors are substantially less likely to stall. This finding aligns with political science literature showing that

broader coalition support increases the likelihood of movement through committees [25]. Additionally, Chamber_Binary (0.7971) and Sponsor_Party_Binary (0.2350) show positive but non-significant coefficients, suggesting that where a bill originates and the party of its sponsor may influence outcomes, though our analysis does not confirm these effects.

Applying Bonferroni corrections helps ensure that the observed significance is not due to random chance from multiple comparisons, thereby reducing the likelihood of false positives. The number of sponsors remains significant even after correction, reinforcing the robustness of this factor as a predictor of whether a bill stalls as seen on Appendix Table 6.

5 Recommendations

To close the widening gap between AI development and safety, policymakers must shift their focus from simply drafting new legislation to ensuring that laws are enacted and enforced. This requires a new legislative mindset, one that emphasizes enforceability, speed, and coordination across sectors. The following recommendations are proposed to address this issue:

- 1. Establish Dedicated AI Policy Committees:** Because committee pigeonholing emerged as the dominant stalling factor, we recommend establishing dedicated AI committees to bypass this choke point. Both the Senate and the House of Representatives should form standing committees or subcommittees focused exclusively on AI Policy and Ethics. The current absence of such dedicated bodies is a major barrier to meaningful oversight and regulatory momentum in AI safety.
- 2. Implement Preemptive Enactment Models:** Modeled after pandemic and cybersecurity laws, these frameworks would activate automatically when specific risk thresholds are crossed. For example, any model exceeding a certain computational power or dataset size would be subject to immediate regulatory oversight.
- 3. Introduce Sunset Clauses:** Rather than wait years for political consensus, legislators should pass temporary AI laws that expire unless actively renewed. This approach creates urgency, enforces regular policy reviews, and ensures that regulation keeps pace with technological evolution.
- 4. Create Independent AI Safety Agencies:** Just as the FDA oversees pharmaceuticals and the FAA governs aviation, the U.S. needs independent, specialized agencies empowered to regulate AI systems, audit compliance, and intervene in development when necessary.

6 Limitations

Our analysis has two main limitations. First, the analysis is limited to federal legislation, excluding state-level nuances and their perspectives. State-level legislation can add additional variable insights. This would require collecting data for all 50 state legislatures. Second, the precision of data classification is 94%. We find that suitable for analysis, though further efforts can improve upon it.

7 Conclusion

The study has multiple positive contributions to help in the process of establishing AI safety legislation by clarifying the factors that are strong predictors of whether or not a bill is delayed. It also provides the first comprehensive subfield and end goal datasets. Artificial Intelligence is advancing at a pace that outstrips the capacity of U.S. legislative mechanisms. Although policymakers increasingly recognize AI's risks, recognition without implementation results in procedural delay rather than meaningful mitigation. This study shows that Congress remains largely confined to the proposal stage, with most bills stalling before enactment. The core challenge is not a lack of awareness, but systemic inaction. Unless Congress shifts from drafting to implementing enforceable measures, AI will continue to evolve without adequate oversight. These findings highlight the urgent need for adaptive, enforceable, and forward-looking legislative frameworks. Symbolic or stalled policies are insufficient; only actionable legislation can ensure AI develops safely and accountably.

References

- [1] Charlotte Alter. "Google's Veo and the Coming Deepfake Crisis." *TIME*, 2025. <https://time.com/7290050/veo-3-google-misinformation-deepfake/>
- [2] Dan Milmo. "DeepSeek AI model sparks safety concerns." *The Guardian*, Jan 29, 2025. <https://www.theguardian.com/technology/2025/jan/29/deepseek-artificial-intelligence-ai-safety-risk-yoshua-bengio>
- [3] LLM Timeline Project. <https://llmtimeline.web.app/>
- [4] Jon Chun, Christian Schroeder de Witt, and Katherine Elkins. *Comparative Global AI Regulation: Policy Perspectives from the EU, China, and the US*. arXiv preprint arXiv:2410.21279, October 5, 2024.
- [5] McCray, Lawrence E., Oye, Kenneth A., and Petersen, Arthur C. "Planned Adaptation in Risk Regulation." *Technological Forecasting and Social Change*, 2009. <https://www.sciencedirect.com/science/article/pii/S0040162509001942>
- [6] Ruschmeier, J. "Limitations of the EU AI Act." *PubMed Central*, 2023. <https://pmc.ncbi.nlm.nih.gov/articles/PMC9827441/>
- [7] AI and Ethics. "Civil Society and Accountability in the AI Act." *Springer*, 2024. <https://link.springer.com/article/10.1007/s43681-024-00595-3>
- [8] John Villasenor. "Soft Law as a Complement to AI Regulation." *Brookings*, 2024. <https://www.brookings.edu/articles/soft-law-as-a-complement-to-ai-regulation/>
- [9] Inside Global Tech. "Overview of AI Regulatory Landscape in APAC." *InsideGlobalTech*, 2024. <https://www.insideglobaltech.com/2024/04/26/overview-of-ai-regulatory-landscape-in-apac/>
- [10] U.S. Congress. "S.5616 - AI Safety Review Act of 2024." <https://www.congress.gov/bill/118th-congress/senate-bill/5616/text>
- [11] U.S. Congress. "The One Big Beautiful Bill Act" <https://www.congress.gov/bill/119th-congress/house-bill/1>
- [12] The White House. *America's AI Action Plan*, July 2025. Available as PDF via White House website: <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>. Accessed August 2025.
- [13] Donald J. Trump. "Trump lays out AI Action Plan during keynote address at summit in Washington, D.C." YouTube video, published one week ago. Available at: https://www.youtube.com/watch?v=JJ7mi0py_bU. Accessed August 2025.
- [14] Washington State Attorney General. "AI Task Force." <https://www.atg.wa.gov/aitaskforce>
- [15] California Senate Bill 1047. https://leginfo.legislature.ca.gov/faces/billNavClient.xhtml?bill_id=202320240SB1047
- [16] United States District Court, Northern District of California. *Order Denying the Plaintiffs' Motion for Partial Summary Judgment and Granting Meta's Cross-Motion for Partial Summary Judgment*, Case No. 3:23-cv-03417-VC (2025). Available at <https://media.npr.org/assets/artslife/arts/2025/order1.pdf>. Accessed July 2025.
- [17] Stanford HAI. "AI Index Report 2025." http://hai.stanford.edu/assets/files/hai_ai_index_report_2025.pdf
- [18] OpenAI. "Introducing GPT-OSS." OpenAI Blog, August 2025. Available at: <https://openai.com/index/introducing-gpt-oss/>. Accessed August 2025.
- [19] Brennan Center. "AI Legislation Tracker." <https://www.brennancenter.org/our-work/research-reports/artificial-intelligence-legislation-tracker>
- [20] UK Government. *International AI Safety Report 2025*. https://assets.publishing.service.gov.uk/media/679a0c48a77d250007d313ee/International_AI_Safety_Report_2025_accessible_f.pdf
- [21] National Conference of State Legislatures. "AI Legislation Tracker 2025." <https://www.ncsl.org/technology-and-communication/artificial-intelligence-2025-legislation>

- [22] American Action Forum. *AI Legislation Tracker*. Available at: <https://www.americanactionforum.org/list-of-proposed-ai-bills-table/>. Accessed August 2025.
- [23] Mintz. *AI Legislation Tracker*. Available at: <https://www.mintz.com/ai-legislation?page=0>. Accessed August 2025.
- [24] le Cessie, S. and van Houwelingen, J. C. "Ridge Estimators in Logistic Regression." *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, vol. 41, no. 1, 1992, pp. 191–201. <https://doi.org/10.2307/2347628>.
- [25] Sotoudeh, Sarah, Porter, Mason A., and Krishnagopal, Sanjukta. "A Network-Based Measure of Cosponsorship Influence on Bill Passing in the United States House of Representatives." *arXiv*, 27 Jun. 2024. <https://arxiv.org/abs/2406.19554>.
- [26] ProQuest. "How a Bill Becomes Law: The Congressional Process." <https://proquest.libguides.com/congressionalhelp/process>
- [27] Brennan Center. "States Take the Lead Regulating AI in Elections—Within Limits." <https://www.brennancenter.org/our-work/research-reports/states-take-lead-regulating-ai-elections-within-limits>
- [28] Anderljung et al. "Frontier AI Regulation: Managing Emerging Risks to Public Safety" *arXiv*, 2023. <https://arxiv.org/abs/2307.03718>
- [29] Appel, Ruth E. "Preemptive AI Regulation: Lessons from Social Media." *arXiv*, 2024. <https://arxiv.org/abs/2412.11335>
- [30] GovTrack. "Statistics and Historical Comparison of Bill Progress." <https://www.govtrack.us/congress/bills/statistics>
- [31] Deirdre Ahern. "The New Anticipatory Governance Culture for Innovation: Regulatory Foresight, Regulatory Experimentation and Regulatory Learning" *arXiv*, 2025. <https://arxiv.org/abs/2501.05921>
- [32] Issa Rice. "Timeline of AI Policy." https://timelines.issarice.com/wiki/Timeline_of_AI_policy
- [33] Issa Rice. "Timeline of AI Safety." https://timelines.issarice.com/wiki/Timeline_of_AI_safety
- [34] Wiseman, Alan E. "The Bipartisan Path to Effective Lawmaking." *The Journal of Politics*, vol. 85, no. 3, May 2023, pp. 1048–1063. <https://doi.org/10.1086/723805>.
- [35] Craig Volden, Alan E. Wiseman, and Laurel Harbridge-Yong. *The Bipartisan Path to Effective Lawmaking. The Journal of Politics*, vol. (to be confirmed), 2024. Available at: <https://www.journals.uchicago.edu/doi/10.1086/723805>. Accessed 2025.
- [36] U.S. Congress. *H.R. 7776 - James M. Inhofe National Defense Authorization Act for Fiscal Year 2023*, 117th Cong. (2022). Became Public Law No. 117-263 on December 23, 2022. Available at: <https://www.congress.gov/bill/117th-congress/house-bill/7776>. Accessed August 2025.
- [37] U.S. Congress. *S. 2551 — AI Training Act*, 117th Cong. (2021–2022). Introduced in the Senate to establish grant programs and AI training for the federal acquisitions workforce. Available at: <https://www.congress.gov/bill/117th-congress/senate-bill/2551>. Accessed August 2025.
- [38] U.S. Congress. "AI Legislation Search Query." <https://www.congress.gov/search?q=%7B%22search%22%3A%22AI%22%2C%22source%22%3A%22legislation%22%2C%22congress%22%3A%5B%22119%22%2C%22115%22%2C%22116%22%2C%22117%22%2C%22118%22%5D%7D>
- [39] Wikipedia. "List of Large Language Models." https://en.wikipedia.org/wiki/List_of_large_language_models

438 A Appendix

439 A.1 Acknowledgments

440 The authors would like to thank Attorney Khalil Khan for his valuable review of the paper’s legal
 441 and legislative dimensions, particularly his guidance on the categorization of AI bills across relevant
 442 sub-fields.

443 A.2 Data Tables

Table 3: Comparison of LLM Breakthroughs and AI Safety Activity (2017–2025)

Year	LLM Breakthroughs	US AI Laws Proposed	AI Laws Enacted (Global)	Non-Legal AI Safety Actions
2017	1	0	3	24
2018	2	2	5	21
2019	3	6	9	25
2020	2	9	5	25
2021	7	14	4	23
2022	10	5	3	28
2023	19	33	13	24
2024	20	54	6	5
Jan–Aug 2025	19	27	1	10

Note: Data from 2025 is not comprehensive, as it was compiled in Aug 2025.

Table 4: Sub-fields of Proposed AI Legislation in the U.S. (2018–Jul 2025)

Year	General Ethical Usage	Policy Advisory	Data Usage	AI in Gov/Military	Push for AI Research	Deepfake	Job Security	AGI	LLM	Autonomous Driving
2018	1	1	0	0	1	0	1	0	0	0
2019	4	2	3	1	2	0	1	1	0	0
2020	4	4	1	4	4	0	0	0	0	0
2021	7	8	4	7	4	0	0	1	0	0
2022	4	2	2	1	0	0	0	1	1	1
2023	16	22	4	14	9	6	4	0	1	0
2024	36	27	26	17	8	8	2	2	3	0
Jan–Aug 2025	14	18	5	13	4	4	3	0	1	0
Total	86	84	45	57	32	18	11	5	6	1

Table 5: Logistic Regression Model Performance Metrics

Metric	Value
Training Accuracy	0.622
Test Accuracy	0.760
Training Error	0.371
Test Error	0.240
Number of Iterations	12
Max Number of Iterations	100

Table 6: Bonferroni-Corrected Logistic Regression Summary: “Advanced AI” = (LLM, AGI, Autonomous Driving Combined)

Feature	Coefficient	Std. Error	Z value	P-value	Coefficient
Advanced AI	-0.7662	0.7780	-0.9848	0.3247	0.7662
AI in Government + Military	-0.5247	0.5309	-0.9882	0.3231	0.5247
Job Security	-0.3767	0.5638	-0.6682	0.5040	0.3767
Push for AI Research	0.3405	0.5332	0.6386	0.5231	0.3405
Policy Advisory	0.2048	0.5834	0.3510	0.7256	0.2048
General Ethical Usage	0.2292	0.5226	0.4385	0.6610	0.2292
Data Usage	-0.1252	0.5574	-0.2246	0.8223	0.1252
Deepfake	0.2187	0.6934	0.3154	0.7525	0.2187
Bipartisan	-0.3282	0.5992	-0.5477	0.5839	0.3282
Chamber_Binary	0.7971	0.5703	1.3977	0.1622	0.7971
Num_Sponsors	0.8068	0.3599	2.2416	*0.0250	0.8068
Sponsor_Party_Binary	0.2350	0.5635	0.4171	0.6766	0.2350

Notes: * p<0.05 (Bonferroni-adjusted).

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The abstract clearly states the four key contributions: (1) quantitative comparison of AI legislation vs LLM breakthroughs, (2) comprehensive taxonomy of policy sub-fields, (3) dataset explaining causes of AI legislation failure, and (4) policy recommendations. These match the results presented in the paper, including the 4.23% enactment rate finding and the logistic regression analysis showing structural/political factors matter more than bill content.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We have added a limitations section.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA] .

Justification: This paper does not include theoretical results requiring formal proofs. The statistical analysis uses standard logistic regression without novel theoretical contributions.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: The paper provides complete reproducibility information including: specific hyperparameters ($C=1.0$, $\text{max_iter}=100$, $\text{solver}=\text{'lbfgs'}$), exact train/test split (80/20 with $\text{random_state}=42$), bootstrap procedure (500 iterations), feature scaling details (Standard-Scaler), and performance metrics. Combined with the available code and data links, this enables full reproduction.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).

550 (d) We recognize that reproducibility may be tricky in some cases, in which case
551 authors are welcome to describe the particular way they provide for reproducibility.
552 In the case of closed-source models, it may be that access to the model is limited in
553 some way (e.g., to registered users), but it should be possible for other researchers
554 to have some path to reproducing or verifying the results.

555 5. Open access to data and code

556 Question: Does the paper provide open access to the data and code, with sufficient instruc-
557 tions to faithfully reproduce the main experimental results, as described in supplemental
558 material?

559 Answer: [Yes]

560 Justification: Code, data, and any other files used in the methodology will be made available
561 at www.anonymized-link.com. Links are anonymized for review.

562 Guidelines:

- 563 • The answer NA means that paper does not include experiments requiring code.
- 564 • Please see the NeurIPS code and data submission guidelines ([https://nips.cc/](https://nips.cc/public/guides/CodeSubmissionPolicy)
565 [public/guides/CodeSubmissionPolicy](https://nips.cc/public/guides/CodeSubmissionPolicy)) for more details.
- 566 • While we encourage the release of code and data, we understand that this might not be
567 possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not
568 including code, unless this is central to the contribution (e.g., for a new open-source
569 benchmark).
- 570 • The instructions should contain the exact command and environment needed to run to
571 reproduce the results. See the NeurIPS code and data submission guidelines ([https://](https://nips.cc/public/guides/CodeSubmissionPolicy)
572 nips.cc/public/guides/CodeSubmissionPolicy) for more details.
- 573 • The authors should provide instructions on data access and preparation, including how
574 to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- 575 • The authors should provide scripts to reproduce all experimental results for the new
576 proposed method and baselines. If only a subset of experiments are reproducible, they
577 should state which ones are omitted from the script and why.
- 578 • At submission time, to preserve anonymity, the authors should release anonymized
579 versions (if applicable).
- 580 • Providing as much information as possible in supplemental material (appended to the
581 paper) is recommended, but including URLs to data and code is permitted.

582 6. Experimental setting/details

583 Question: Does the paper specify all the training and test details (e.g., data splits, hyper-
584 parameters, how they were chosen, type of optimizer, etc.) necessary to understand the
585 results?

586 Answer: [Yes]

587 Justification: We provide comprehensive experimental details, including hyperparameters,
588 data splits, feature preprocessing, convergence criteria, performance evaluation methods, as
589 well as sources of data. All necessary details for understanding and reproducing the results
590 are present.

591 Guidelines:

- 592 • The answer NA means that the paper does not include experiments.
- 593 • The experimental setting should be presented in the core of the paper to a level of detail
594 that is necessary to appreciate the results and make sense of them.
- 595 • The full details can be provided either with the code, in appendix, or as supplemental
596 material.

597 7. Experiment statistical significance

598 Question: Does the paper report error bars suitably and correctly defined or other appropriate
599 information about the statistical significance of the experiments?

600 Answer: [Yes]

Justification: The paper includes p-values and standard errors from bootstrap resampling (500 iterations) for the logistic regression coefficients, and reports z-scores for significance testing. The methodology paragraph describes the statistical procedures used.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We mention the Logistic Regression model was "trained and run on Google Cloud" using "a virtualized Intel Xeon CPU @ 2.20GHz".

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: This research analyzes publicly available legislative data and does not involve human subjects, privacy violations, or potential harmful applications. The work aims to improve AI governance and safety.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.

- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: The paper clearly discusses the positive impacts: improving AI safety governance and policy recommendations. The paper proposes recommendations to safeguard against the negative societal impacts of foundation AI development. To the best of our knowledge, there is no negative societal impact of our study.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA] .

Justification: This paper analyzes publicly available legislative data and does not release models or datasets with high misuse potential. The legislative analysis data poses minimal safety risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [\[Yes\]](#)

Justification: The paper used data sources: Congress.gov, HAI Index Report, Brennan Center etc. Declared references are provided.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: The paper creates new datasets (comprehensive AI bill classification, endpoint analysis) for public usage. They will be released as part of camera-ready submissions.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: This research does not involve crowdsourcing or human subjects research. The human annotation mentioned was performed by the authors.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.

- 755 • According to the NeurIPS Code of Ethics, workers involved in data collection, curation,
756 or other labor should be paid at least the minimum wage in the country of the data
757 collector.

758 **15. Institutional review board (IRB) approvals or equivalent for research with human**
759 **subjects**

760 Question: Does the paper describe potential risks incurred by study participants, whether
761 such risks were disclosed to the subjects, and whether Institutional Review Board (IRB)
762 approvals (or an equivalent approval/review based on the requirements of your country or
763 institution) were obtained?

764 Answer: [NA]

765 Justification: This research does not involve human subjects and therefore does not require
766 IRB approval.

767 Guidelines:

- 768 • The answer NA means that the paper does not involve crowdsourcing nor research with
769 human subjects.
- 770 • Depending on the country in which research is conducted, IRB approval (or equivalent)
771 may be required for any human subjects research. If you obtained IRB approval, you
772 should clearly state this in the paper.
- 773 • We recognize that the procedures for this may vary significantly between institutions
774 and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the
775 guidelines for their institution.
- 776 • For initial submissions, do not include any information that would break anonymity (if
777 applicable), such as the institution conducting the review.

778 **16. Declaration of LLM usage**

779 Question: Does the paper describe the usage of LLMs if it is an important, original, or
780 non-standard component of the core methods in this research? Note that if the LLM is used
781 only for writing, editing, or formatting purposes and does not impact the core methodology,
782 scientific rigorousness, or originality of the research, declaration is not required.

783 Answer: [Yes]

784 Justification: The paper clearly describes using GPT-4o for bill classification, including
785 confidence levels, accuracy validation (94%), and human oversight for low-confidence
786 classifications. This is properly documented as a core methodological component.

787 Guidelines:

- 788 • The answer NA means that the core method development in this research does not
789 involve LLMs as any important, original, or non-standard components.
- 790 • Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>)
791 for what should or should not be described.