

Governing with Artificial Intelligence: Use, Ideology, and the Benefits and Risks of AI in State Government

Abstract

Artificial intelligence is increasingly entering government institutions, yet we know relatively little about how government professionals evaluate its potential benefits and risks or how these attitudes relate to organizational adoption. This study uses an original survey of individuals with experience in U.S. state government to examine the technological, political, and institutional correlates of attitudes toward governmental AI. The results reveal a striking asymmetry. Frequency of AI use strongly predicts perceived benefits: respondents who use AI more frequently report substantially greater confidence in its potential governmental utility, even after accounting for AI familiarity, political ideology, government experience, and office size. Political ideology, by contrast, is the dominant predictor of perceived risk. More conservative respondents perceive significantly lower AI-related risks, while the association between AI use and risk perceptions becomes nonsignificant after ideology is introduced. Nested models reinforce this distinction: AI familiarity and use explain substantial additional variation in perceived benefits, whereas ideology contributes substantial additional explanatory power for perceived risks. These relationships remain robust to HC3 standard errors and influential-observation sensitivity analyses. Individual attitudes, however, do not significantly predict reported organizational AI implementation, although the adoption analysis is limited by sample size. The findings demonstrate that attitudes toward governmental AI are multidimensional: technological experience is most strongly associated with perceived utility, while political ideology is most strongly associated with perceived risk.

Keywords: artificial intelligence; state government; public administration; technology adoption; political ideology; AI governance; risk perception; government innovation

Introduction

Artificial intelligence is moving from the private sector into the everyday work of government. Generative AI systems can summarize documents, process large volumes of information, draft correspondence, write press releases, inform legislative decision-making, assist with constituent communication, support policy research, and automate administrative tasks. For state governments confronting limited staff capacity, growing workloads, and demands for more responsive public services, these technologies offer potentially substantial benefits. Yet the same systems raise concerns about algorithmic bias, misinformation, transparency, accountability, unequal treatment, privacy, and the capacity of public institutions to govern technologies that are evolving faster than many administrative rules and procedures.

These competing possibilities have generated a rapidly expanding debate over AI governance. Much of that debate, however, approaches government as the regulator of artificial intelligence rather than as an institution whose employees increasingly encounter, evaluate, and potentially use AI themselves and must promulgate rules for how it uses AI internally. This distinction matters. Decisions about whether and how governments integrate emerging technologies are made within organizations populated by public officials and employees who may differ considerably in their familiarity with AI, frequency of use, professional experience, and assessments of its benefits and risks. Understanding governmental adoption, therefore, requires attention not only to formal AI policies but also to the individuals operating within public institutions.

Existing research on technology adoption suggests several possible sources of variation. One possibility is experiential: individuals who interact more frequently with a technology may develop different assessments of its usefulness and limitations than individuals with little direct experience. Another is institutional: officials working in larger offices or possessing greater governmental experience may encounter different administrative capacities, professional norms, or implementation constraints. A third possibility is political. Artificial intelligence has increasingly become embedded within broader political debates about misinformation, corporate power, discrimination, government regulation, technological innovation and views on the broader culture of Silicon Valley. Perceptions of AI may therefore reflect ideological orientations as well as technological experience.

Importantly, support for AI need not constitute a single attitude. An individual may simultaneously believe that AI can improve governmental efficiency and that it creates serious risks for democratic governance. Treating attitudes toward AI as a simple continuum from optimism to skepticism may consequently obscure meaningful differences between assessments of utility and assessments of risk. The factors that lead government professionals to recognize AI's administrative benefits may not be the same factors that shape concerns about bias, misinformation, transparency, or public trust.

Artificial Intelligence, Bureaucratic Politics, and Public Administration

Artificial intelligence is entering government through institutions that are simultaneously administrative and political. Political science and public administration have long demonstrated that bureaucratic behavior cannot be understood as the mechanical implementation of decisions made elsewhere: public officials exercise discretion within institutional environments, and their

experiences, values, and political orientations can shape administrative behavior. Generative AI introduces a new arena for these longstanding dynamics. Decisions about whether and how AI is used are therefore not simply technical choices; they are administrative and political judgments made by officials embedded within bureaucratic institutions.

This observation connects the emerging literature on governmental artificial intelligence to longstanding scholarship in political science and public administration concerning bureaucracy. Political scientists have long rejected a purely mechanical conception of administration in which bureaucrats simply execute decisions made elsewhere. Ambiguous statutes, incomplete political directives, specialized expertise, and the complexity of implementation leave administrative actors with substantial discretion (Lipsky 1980; Meier 1975). Indeed, Meier's early work on representative bureaucracy emphasized bureaucratic values as potentially important for understanding administrative behavior, while subsequent research has demonstrated that discretion creates opportunities for administrators' values and predispositions to influence implementation (Meier and Bohte 2001; Sowa and Selden 2003).

This literature provides an important starting point for studying AI in government. If public employees possess meaningful discretion over how administrative work is performed, then their evaluations of emerging technologies may matter for how those technologies enter governmental institutions. Bureaucrats are not simply passive recipients of technological change. They encounter new technologies through their existing professional experiences, political beliefs, organizational responsibilities, and understandings of appropriate administrative behavior. Scholarship on representative bureaucracy similarly emphasizes that bureaucrats' attitudes, values, and predispositions can matter for policy implementation when administrative actors possess sufficient discretion (Sowa and Selden 2003; Riccucci and Van Ryzin 2017).

These questions become particularly important with artificial intelligence because automation can redistribute rather than simply eliminate administrative discretion. Bullock (2019) argues that AI can alter the location and character of discretion within bureaucratic organizations, making the relationship between technological systems and administrative judgment central to understanding AI in government.

Artificial intelligence makes these questions particularly consequential. AI potentially expands administrative capacity while simultaneously introducing new forms of discretion and dependence on technological systems. Public organizations may use AI to process information, draft materials, communicate with citizens, or support decisions. These applications can improve efficiency and reduce demands on employees, but government use of AI also implicates core public values because administrative decisions distribute benefits, burdens, opportunities, and rights. Concerns about algorithmic bias, opacity, misinformation, privacy, accountability, and unequal treatment, therefore, extend beyond technical performance to questions of democratic governance. Young, Bullock, and Lecy (2019) similarly conceptualize artificial intelligence as a new form of administrative discretion, emphasizing that its consequences for public administration must be evaluated not only in terms of efficiency but also in terms of values such as equity, manageability, and political feasibility. AI adoption, therefore, implicates broader questions about the appropriate exercise of public authority. The consequences of AI may also depend on the type of administrative task and the relationship between human and artificial discretion. Bullock, Young, and Wang (2020) argue that AI can reshape bureaucratic form by changing how discretion is distributed between public employees and automated systems. Governmental AI should therefore be understood as potentially reorganizing administrative judgment rather than merely providing officials with another technical tool.

Public-administration scholarship increasingly conceptualizes these concerns through the language of public values. Schiff, Schiff, and Pierson (2022), for example, demonstrate experimentally that perceived failures of fairness and transparency in government automated decision systems reduce public support, highlighting the importance of values that distinguish public-sector technology adoption from private consumption. Government AI consequently presents administrators with a dual evaluation: what can the technology enable the government to accomplish, and what risks does its use create for legitimate administration?

The broader literature likewise demonstrates that governmental AI presents both opportunities and governance challenges. In their systematic review, Zuiderwijk, Chen, and Salem (2021) identify potential benefits of AI for public governance alongside concerns involving transparency, accountability, privacy, bias, and institutional capacity. These competing implications reinforce the importance of distinguishing perceptions of AI's administrative utility from perceptions of its governance risks.

The distinction is theoretically important because assessments of technological utility and assessments of technological risk need not be opposite ends of a single attitude. Public officials may simultaneously believe that AI can make government more efficient and that its use creates serious problems involving bias, transparency, or accountability. The political and administrative literatures suggest, moreover, that these two assessments may have different origins. Practical technological experience may be especially important for evaluations of utility, while political orientations may be especially important for evaluations of risk.

Technological Experience and Perceived Administrative Benefits

Research also suggests that public officials do not necessarily interpret emerging technologies uniformly. Criado and de Zarate-Alcarazo (2022) demonstrate that public managers can hold distinct technological frames through which they understand artificial intelligence, highlighting the importance of officials' perceptions and interpretations in shaping governmental encounters with AI.

Research on technological change in public organizations suggests that employee attitudes are an important component of implementation. Ahn and Chen (2022), studying U.S. government employees, find that willingness to support governmental AI is related to perceptions of its benefits and harms and that familiarity or experience with AI applications is positively associated with willingness to support its use. Their findings emphasize employee buy-in as an important component of AI-augmented public administration.

This argument is consistent with broader technology-acceptance research emphasizing perceived usefulness in individuals' orientations toward emerging technologies. Yet the bureaucratic setting adds an important dimension. Public employees evaluate technologies not simply as consumers but as administrative actors confronting workloads, informational demands, procedural requirements, and expectations of responsiveness. AI may be valuable precisely because it changes administrators' capacity to perform these tasks.

More recent experimental research further demonstrates that public servants' acceptance of AI depends on how the technology is incorporated into administrative work. Haesevoets, Verschuere, and Roets (2025) find that public-sector employees are more accepting of AI when it performs an advisory rather than fully autonomous role and when it is applied to tasks more amenable to technological assistance. Evaluations of

governmental AI, therefore, depend not merely on abstract attitudes toward the technology but also on how officials understand its practical administrative role. These attitudes are also potentially responsive to information and institutional design. Haesevoets et al. (2026) find experimentally that public servants' attitudes toward AI in policymaking vary according to the role assigned to AI and that providing information about AI applications can affect those attitudes. Such findings reinforce the importance of examining officials' technological knowledge and engagement when explaining evaluations of governmental AI.

Direct use should therefore provide information about AI that abstract familiarity does not. An individual can become familiar with artificial intelligence through news coverage, professional conversations, training, or policy debates without actually using the technology. Direct users, by contrast, encounter concrete examples of what AI does well and poorly. They can observe whether it reduces the time necessary to draft correspondence, synthesizes information effectively, assists research, or generates outputs requiring extensive correction.

The direction of the relationship is not necessarily causal. Employees who initially view AI favorably may be more likely to use it, while using AI may itself reveal applications that increase perceptions of usefulness. Nevertheless, either mechanism implies an important relationship between technological engagement and perceived administrative benefits.

This study therefore distinguishes AI familiarity from frequency of AI use. Existing research suggests that familiarity and experience can increase support for governmental AI (Ahn and Chen 2022), but comparatively less attention has been paid to whether direct use carries explanatory power independent of generalized familiarity. If experiential engagement provides

distinctive information about what AI can accomplish in an administrative setting, actual use should be especially strongly associated with perceptions of governmental utility.

H1: More frequent use of artificial intelligence will be associated with greater perceived benefits of governmental AI.

General knowledge may nevertheless matter independently because individuals with greater familiarity with AI may possess greater awareness of possible governmental applications even without frequent direct use.

H2: Greater familiarity with artificial intelligence will be associated with greater perceived benefits of governmental AI.

Political Ideology and Perceptions of AI Risk

Technological experience, however, may provide only a partial account of how public officials evaluate artificial intelligence. Bureaucrats are not politically neutral in the sense of possessing no political values or orientations; like other individuals, they bring predispositions to judgments about administrative problems and governmental action.

Recent research has more directly considered administrators' political orientations. Uttermark (2024), for example, argues that administrators operate in contexts of substantial discretion and examines whether their political orientations structure how they evaluate federal actions. This work connects a longstanding political-science concern with bureaucratic discretion to the possibility that ideological beliefs influence administrators' evaluations of governmental action.

This expectation also connects governmental AI to a broader political-science literature demonstrating that political predispositions can structure perceptions of risk, threat, and public problems. Individuals do not necessarily evaluate risks solely on the basis of technical information; political values can shape which harms receive attention, how threatening those harms appear, and whether government intervention is considered appropriate (Douglas and Wildavsky 1982; Kahan et al. 2012; McCright and Dunlap 2011). Such dynamics have been documented across politically contested issues involving science, technology, regulation, and environmental risk. Governmental AI may increasingly constitute another domain in which ostensibly technical judgments become filtered through political orientations.

Artificial intelligence provides a particularly important context in which to examine this possibility because AI has become embedded in political disputes over regulation, inequality, discrimination, misinformation, corporate power, surveillance, labor displacement, and government intervention. Evaluating AI's risks, therefore, involves more than estimating the probability that a technology will malfunction. It can involve judgments about which social harms deserve attention, whether inequalities require governmental intervention, how much transparency institutions owe citizens, and how aggressively governments should regulate emerging technologies.

This suggests that ideology may be particularly consequential for risk perception. Consider two officials who agree that a generative AI system can summarize documents effectively. That technical agreement does not require them to agree about whether algorithmic bias represents a severe threat, whether the government is adequately prepared to oversee AI, how consequential misinformation is, or how much regulation is appropriate. Technological

experience can provide shared information about capabilities without eliminating political disagreement about acceptable risks.

This argument is also consistent with the public-values literature on algorithmic government. Fairness, transparency, responsiveness, and accountability are not ancillary concerns in public administration; they are criteria through which legitimate government action is evaluated. Schiff, Schiff, and Pierson (2022) demonstrate that failures involving fairness and transparency can meaningfully affect evaluations of government automated decision systems. The perceived seriousness of these risks, however, may vary across administrators because judgments about social harm and governmental responsibility can themselves be politically structured.

The study, therefore, departs from approaches that conceptualize attitudes toward AI along a single continuum ranging from technological optimism to technological pessimism. Someone may simultaneously perceive substantial administrative benefits and substantial governance risks. Conversely, perceiving relatively little risk does not necessarily imply believing that AI is technically useful.

This multidimensional conception generates a theoretically important possibility: experience may structure what administrators believe AI can do, while ideology structures what they believe AI might threaten. The distinction places governmental AI at the intersection of technology-adoption research and political-science scholarship on bureaucratic values and political orientations.

If concerns about algorithmic discrimination, inequality, misinformation, and governmental preparedness are partly structured by broader ideological orientations toward

social risk and governmental responsibility, more conservative officials should express lower levels of concern about these dimensions of governmental AI. Accordingly:

H3: More conservative ideological identification will be associated with lower perceived risks from governmental artificial intelligence.

The analysis also considers ideology as a predictor of perceived benefits. The multidimensional framework, however, provides no reason to assume that ideological differences in perceived risks must produce equivalent ideological differences in perceived technological utility. Emerging research on generative AI also suggests that public servants' orientations toward the technology cannot necessarily be reduced to simple acceptance or rejection. Zhu and Gao (2026) identify heterogeneous ways in which civil servants understand generative AI, illustrating that officials can hold distinct configurations of beliefs about its capabilities, appropriate uses, and consequences. This heterogeneity reinforces the value of treating perceived utility and perceived risk as analytically distinct dimensions.

Public Organizations, Institutions, and AI Adoption

Individual attitudes are only one component of technological change in government. Political science and public administration have long emphasized that organizational behavior is embedded within institutional environments that constrain individual choice. Public agencies operate through rules, routines, professional norms, budgets, hierarchies, legal mandates, and relationships with other political institutions. Consequently, even administrators possessing substantial discretion operate within organizational structures that shape which innovations can actually be implemented.

Institutional theory provides one framework for understanding this distinction. DiMaggio and Powell (1983) argue that organizations respond to coercive, normative, and mimetic pressures within their institutional environments. Public organizations may consequently adopt technologies because of legal or political requirements, professional expectations, funding incentives, or the behavior of peer organizations rather than solely because individual employees consider those technologies useful.

Recent research demonstrates the relevance of this framework for governmental AI. Rodriguez Müller and colleagues examine AI adoption among public managers through institutional theory and find that coercive, normative, and mimetic pressures can increase willingness to adopt AI. Importantly, managerial perceptions interact with these institutional pressures rather than completely determining adoption themselves. The implication is that AI adoption emerges from interactions between individuals and institutional environments rather than from individual attitudes alone.

The present analysis does not directly measure these institutional pressures. Rather, this literature provides a reason to expect that individual characteristics alone may offer an incomplete explanation of organizational adoption.

This distinction is particularly important in government. An employee may perceive substantial benefits from AI but work in an organization without procurement authority, cybersecurity approval, technical infrastructure, training, leadership support, or formal policies permitting its use. Conversely, an organization may implement AI because of leadership decisions, external mandates, professional pressures, or strategic considerations even when some employees remain skeptical.

Research on public-sector AI adoption similarly points to determinants operating at multiple levels. Madan and Ashok's (2023) systematic review demonstrates that AI adoption in public administration reflects technological, organizational, and environmental conditions rather than technological characteristics alone. This perspective suggests that individual enthusiasm for AI may be insufficient to produce organizational implementation when institutional capacity or organizational conditions are unfavorable.

Organizational choices concerning AI also involve decisions about how technological expertise is incorporated into existing administrative structures. Selten and Klievink (2024) show that public organizations adopting AI must navigate tensions between separating specialized AI capabilities and integrating them into established organizational processes. Adoption is therefore an organizational process rather than simply an aggregation of employees' attitudes toward the technology.

Organizational capacity further complicates this relationship. Larger organizations may possess greater financial and technical resources for evaluating new technologies, but they may also confront more complex procurement processes and bureaucratic procedures. Experienced administrators may possess greater knowledge of institutional processes while simultaneously being more embedded in established routines. Thus, characteristics that shape an individual's evaluation of AI need not predict organizational implementation in a straightforward manner.

This produces an important distinction between the micro-politics of bureaucratic attitudes and the meso-level politics of organizational adoption. Political science scholarship suggests that administrators' values matter when they possess discretion, while institutional scholarship demonstrates that organizational outcomes cannot be reduced to individual

preferences alone. Indeed, representative-bureaucracy research itself emphasizes that the translation of individual characteristics and values into administrative outcomes is conditional on discretion and organizational context (Sowa and Selden 2003).

The present study, therefore, treats organizational AI implementation as a separate outcome rather than assuming that favorable individual attitudes necessarily translate into institutional adoption. This allows the analysis to distinguish two questions that are sometimes conflated: what shapes how government professionals evaluate artificial intelligence, and what shapes whether governmental organizations actually implement it?

Contribution

Bringing these literatures together reveals an important gap. Research on governmental AI has increasingly examined technological acceptance, employee willingness to use AI, institutional pressures, and public values. Political science and public administration, meanwhile, have developed extensive literatures demonstrating that bureaucratic behavior is shaped by discretion, values, political orientations, and institutional context. Yet these perspectives have rarely been integrated to explain why government professionals may evaluate the benefits and risks of the same emerging technology differently.

This study bridges those literatures by conceptualizing governmental AI attitudes as multidimensional. Rather than asking whether government professionals are simply supportive or skeptical of AI, it distinguishes perceived administrative utility from perceived governance risk and examines whether those dimensions have different correlates. It further separates individual evaluations from reported organizational implementation.

This framework generates a broader theoretical expectation: technological experience and political ideology structure different components of administrative judgment. Experience informs judgments about what a technology can do; ideology structures judgments about what it might threaten. Organizational adoption, meanwhile, is embedded within institutional constraints that extend beyond either individual technological experience or political beliefs.

This approach places artificial intelligence within a longstanding political-science and public-administration question: how do the characteristics of the people who govern interact with the institutions in which they operate to shape the exercise of public authority? AI changes the technology available to government, but it does not remove the political and institutional processes through which government decides what technologies are useful, what risks are acceptable, and when innovations should become part of administrative practice.

This study examines these questions using an original survey of individuals with experience in U.S. state government. It asks three related questions: What predicts perceptions of the governmental benefits of artificial intelligence? What predicts perceptions of its risks? And do these individual attitudes correspond with reported organizational implementation of AI?

The findings reveal a striking asymmetry. Perceptions of AI's benefits are associated primarily with actual AI use. Respondents who use AI more frequently perceive substantially greater potential for the technology to improve governmental functions, even after accounting for familiarity, professional experience, ideology, and office size. Political ideology, by contrast, is the dominant correlate of perceived AI risks. More conservative respondents report substantially

lower levels of concern about AI-related risks, while the apparent association between AI use and lower risk perceptions disappears after ideology is taken into account.

Nested models reinforce this distinction. Adding AI familiarity and use to institutional characteristics increases explained variance in perceived benefits by approximately 48 percentage points, while subsequently adding ideology provides little additional explanatory power. For perceived risks, ideology adds approximately 24 percentage points of explained variance beyond technological and institutional characteristics. Yet neither pattern clearly translates into organizational adoption: none of these individual-level characteristics significantly predicts whether respondents report that their offices have implemented AI, although the relatively small number of adopting offices warrants caution.

Recent research further illustrates the organizational constraints surrounding public-sector AI. Robinson (2026) demonstrates that public agencies' choices concerning AI models are shaped by technological and organizational barriers, underscoring that adoption decisions depend on institutional capabilities and constraints beyond individual employees' assessments of AI.

These findings suggest that government professionals' orientations toward artificial intelligence are multidimensional. Experience appears to structure perceptions of what AI can do, while ideology structures perceptions of what AI might threaten. The distinction has implications for research on technology adoption, administrative capacity, and the governance of emerging technologies. Governments confronting AI are not merely deciding whether to adopt a new technical tool; they are doing so through institutions whose personnel hold systematically different understandings of both its promise and its dangers.

Data and Methods

The analysis uses an original survey of individuals with experience in U.S. state government. Seventy-nine completed survey records were collected. Four records contained no substantive responses across the principal variables, leaving approximately 75 respondents for most descriptive analyses. Sample size varies across multivariate models because of item nonresponse, particularly on political ideology.

The survey measures respondents' professional backgrounds, familiarity with and use of artificial intelligence, assessments of governmental AI, and whether their offices have implemented AI. The analysis focuses on three outcomes: perceived benefits, perceived risks, and reported organizational implementation.

Dependent Variables

The AI Benefits Index combines four five-point Likert items measuring whether respondents believe AI can improve governmental efficiency, information processing, constituent communication and responsiveness, and public-service delivery. Responses range from 1, strongly disagree, to 5, strongly agree. The items demonstrate excellent internal reliability (Cronbach's $\alpha = .90$, 95% CI [.86, .93]). Corrected item-total correlations range from .71 to .84, and removing any item reduces reliability. The four items are therefore averaged into an index ranging from 1 to 5 ($M = 3.11$, $SD = 1.15$).

The AI Risk Index combines seven items capturing concerns about bias and fairness, reproduction of existing inequalities, transparency, public trust, misinformation, rapid governmental adoption, and governmental preparedness. The seven items demonstrate strong

reliability ($\alpha = .84$; standardized $\alpha = .85$; 95% CI [.78, .89]). Corrected item-total correlations range from .48 to .69. The resulting index ranges from 1 to 5, with higher scores indicating greater perceived risk ($M = 4.07$, $SD = .77$).

The third outcome measures organizational AI implementation. Respondents reporting that their office had implemented AI are coded 1, and those reporting no implementation are coded 0. Respondents who selected "unsure" are excluded from the binary model. Eighteen respondents reported implementation and 52 reported no implementation.

Independent Variables

Frequency of personal AI use ranges from 1 ("never") to 5 ("daily"). AI familiarity is coded ordinally from lower to higher familiarity. Government experience ranges from 1 ("less than one year") to 5 ("16+ years"). Office size ranges from 1 ("1–5") to 5 ("51+"). Political ideology ranges from 1 ("very liberal") to 7 ("very conservative"). Respondents selecting "prefer not to answer" on ideology are treated as missing.

Estimation

OLS regression models estimate perceived benefits and perceived risks. The full specifications include AI familiarity, AI use, government experience, ideology, and office size. The complete-case sample for these specifications is $N = 66$.

Nested models evaluate the incremental explanatory contribution of different sets of predictors using the same complete-case sample. Baseline models contain government experience and office size. AI familiarity and use are subsequently introduced, followed by ideology.

The organizational-adoption outcome is analyzed using binary logistic regression. Because only 18 respondents report adoption and the complete-case model contains 61 observations, estimates from this analysis are interpreted cautiously.

Diagnostic tests indicate no meaningful multicollinearity; variance inflation factors range from approximately 1.09 to 1.46. Breusch-Pagan tests provide no evidence of heteroskedasticity for either the benefits (BP = 6.66, $p = .247$) or risks (BP = 2.58, $p = .765$) models. HC3 heteroskedasticity-robust standard errors are nevertheless used as a robustness check.

Results

The multivariate analyses reveal a clear distinction between the factors associated with perceptions of artificial intelligence's governmental benefits and those associated with perceptions of its risks. Whereas frequency of AI use is the strongest predictor of perceived benefits, political ideology is the strongest predictor of perceived risks. Figure 1 summarizes the coefficient estimates from the fully specified benefits and risks models. Across both models, government experience, office size, and general familiarity with AI are comparatively weak predictors.

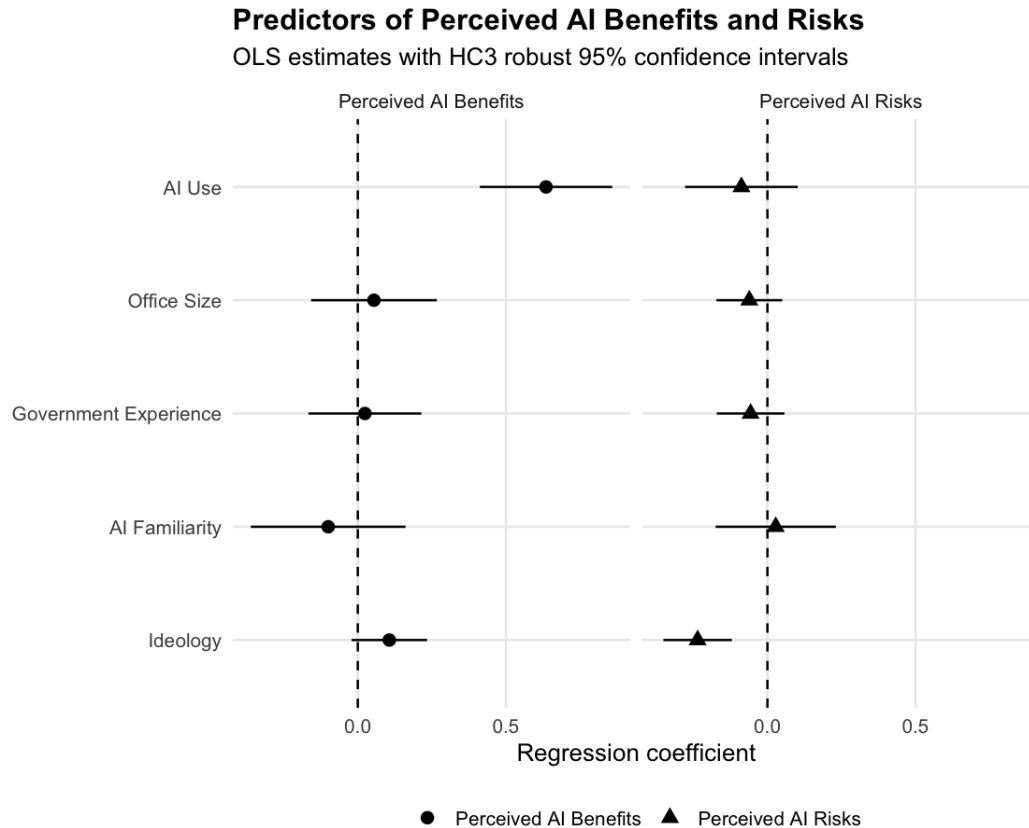


Figure 1. Predictors of Perceived AI Benefits and Risks.

Notes: Points represent unstandardized OLS regression coefficients from the fully specified models. Horizontal lines represent 95% confidence intervals calculated using HC3 heteroskedasticity-robust standard errors. Models include AI familiarity, frequency of AI use, government experience, political ideology, and office size. The dashed vertical line represents a coefficient of zero.

AI Use and Perceived Benefits

The first analysis examines perceptions of AI's potential governmental benefits. The fully specified model explains approximately half of the observed variation in the AI Benefits Index ($R^2 = .515$; adjusted $R^2 = .475$), $F(5, 60) = 12.75$, $p < .001$.

Consistent with H1, frequency of AI use is a strong positive predictor of perceived benefits. Holding AI familiarity, government experience, political ideology, and office size constant, each one-category increase in frequency of AI use is associated with a .635-point increase on the five-point benefits index ($b = .635$, $SE = .107$, $\beta = .64$, $p < .001$). The association remains highly statistically significant when HC3 robust standard errors are employed ($b = .635$, HC3 $SE = .114$, $p < .001$).

The magnitude of this relationship is substantively considerable. Because frequency of AI use ranges from 1 (“never”) to 5 (“daily”), the coefficient implies substantial differences in predicted assessments of AI’s governmental benefits across the range of AI use. Figure 2 illustrates this relationship using adjusted predictions from the fully specified model. Predicted perceptions of AI benefits increase steadily as respondents move from never using AI to using it daily.

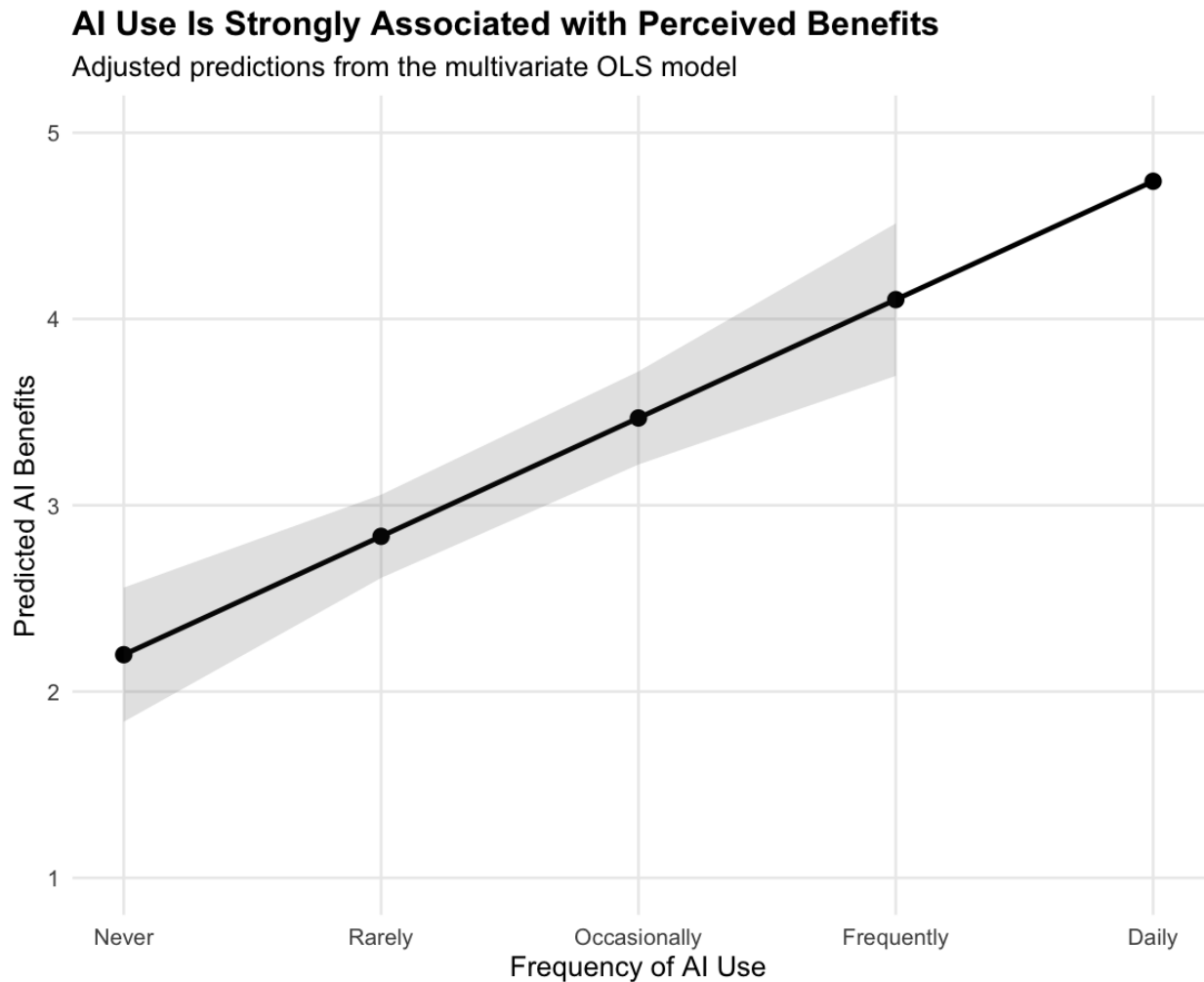


Figure 2. Predicted Perceptions of AI Benefits by Frequency of AI Use.

Notes: Predicted values are derived from the multivariate OLS model controlling for AI familiarity, government experience, political ideology, and office size. Shaded areas represent 95% confidence intervals. Higher values indicate greater perceived benefits of artificial intelligence in government.

General familiarity with AI, however, does not independently predict perceived benefits ($b = -.100$, $p = .452$). H2 is therefore not supported. This distinction between familiarity and use

is notable: direct engagement with AI is strongly associated with perceptions of its governmental utility, whereas simply reporting greater familiarity with the technology is not.

Neither government experience ($b = .024$, $p = .797$) nor office size ($b = .054$, $p = .561$) significantly predicts perceived benefits. Political ideology is positively associated with perceived benefits at the bivariate level, but its relationship is substantially reduced after accounting for AI use and the other covariates. In the fully specified model, ideology does not reach conventional levels of statistical significance ($b = .106$, $\beta = .18$, $p = .103$; HC3 $p = .107$).

The nested specifications further demonstrate the explanatory importance of technological engagement. Using the common complete-case sample ($N = 66$), a baseline model containing government experience and office size explains only 1.2% of the variance in perceived benefits ($R^2 = .012$). Adding AI familiarity and frequency of AI use increases explained variance to 49.3% ($R^2 = .493$), representing a 48.1-percentage-point increase in explained variance. This improvement is statistically significant, $F\text{-change} = 28.96$, $p < .001$. Adding political ideology increases R^2 only modestly, from .493 to .515 ($\Delta R^2 = .022$), and does not significantly improve model fit, $F(1, 60) = 2.75$, $p = .103$.

Taken together, these findings indicate that perceptions of AI's governmental utility are associated much more strongly with respondents' actual engagement with AI technology than with their ideological orientation, professional experience, or organizational context.

Ideology and Perceived AI Risks

A substantially different pattern emerges for perceptions of AI-related risks. The fully specified model explains 43.1% of the variance in the AI Risk Index ($R^2 = .431$; adjusted $R^2 = .383$), $F(5, 60) = 9.08$, $p < .001$.

Consistent with H3, political ideology is the strongest predictor of perceived AI risks. Because ideology is coded from very liberal to very conservative, the negative coefficient indicates that respondents become less likely to perceive AI as risky as they move toward the conservative end of the ideological scale. Holding AI familiarity, AI use, government experience, and office size constant, each one-category movement toward the conservative end of the scale is associated with a -.236-point decrease in perceived AI risks ($b = -.236$, $SE = .047$, $\beta = -.59$, $p < .001$). The relationship remains highly statistically significant using HC3 robust standard errors ($b = -.236$, HC3 $SE = .059$, $p < .001$).

Figure 3 illustrates the magnitude of this ideological gradient. Predicted perceptions of AI risk decline steadily across the ideological spectrum, with liberal respondents exhibiting the highest predicted levels of concern and conservative respondents exhibiting the lowest.

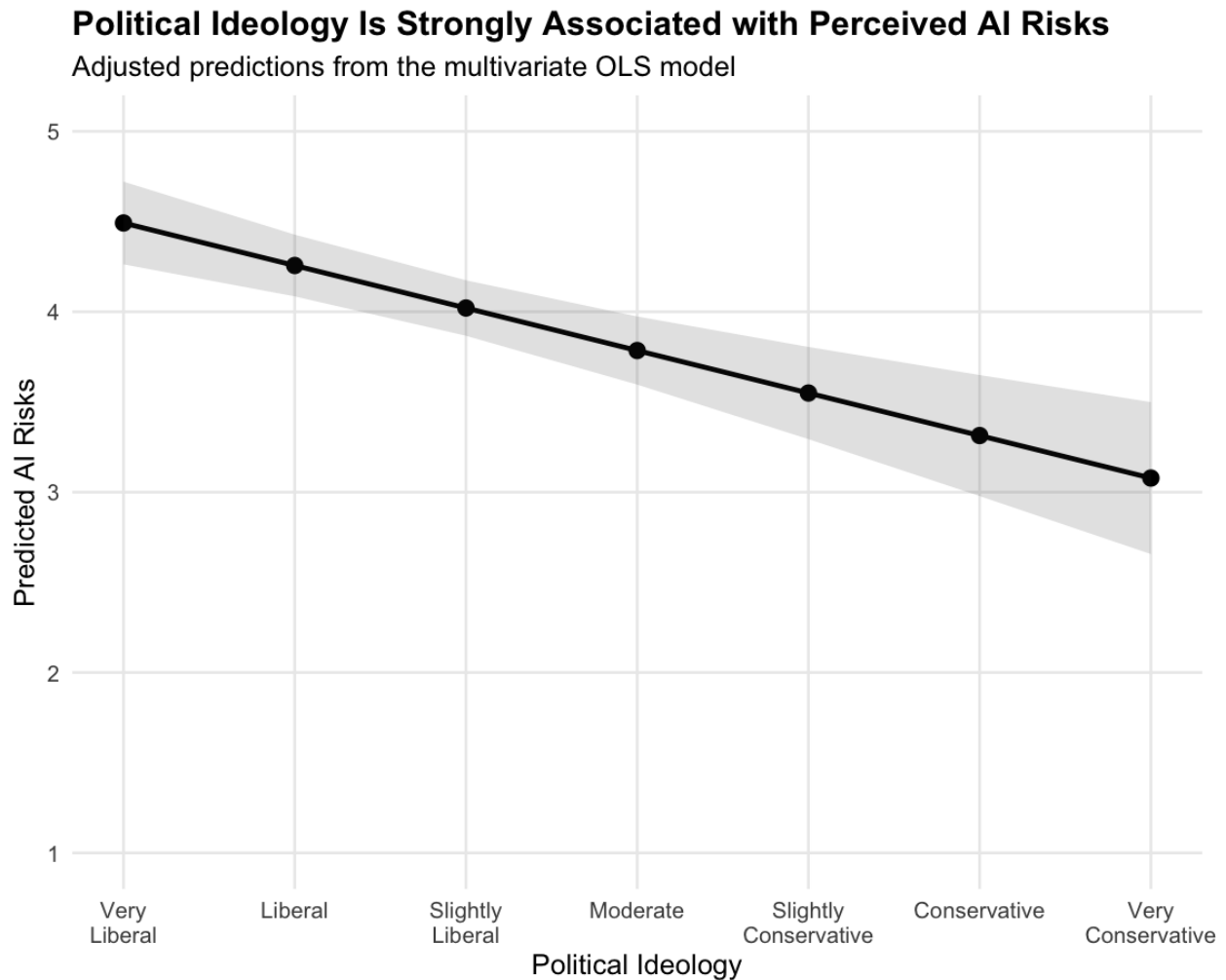


Figure 3. Predicted Perceptions of AI Risks by Political Ideology.

Notes: Predicted values are derived from the multivariate OLS model controlling for AI familiarity, frequency of AI use, government experience, and office size. Shaded areas represent 95% confidence intervals. Political ideology ranges from very liberal to very conservative. Higher outcome values indicate greater perceived risks associated with artificial intelligence in government.

Frequency of AI use exhibits an especially revealing pattern across specifications. Before political ideology is introduced, more frequent AI use is significantly associated with lower

perceived risks ($b = -.277, p < .001$). Once ideology is included, however, the magnitude of the AI-use coefficient declines considerably and becomes statistically nonsignificant ($b = -.088, \beta = -.13, p = .267$; HC3 $p = .369$). Thus, the association between AI use and perceived risk is substantially attenuated after accounting for respondents' ideological orientations. Because these data are cross-sectional, this attenuation should not be interpreted as evidence of causal mediation.

AI familiarity ($b = .028, p = .773$), government experience ($b = -.057, p = .405$), and office size ($b = -.061, p = .373$) do not independently predict perceived AI risks.

The nested models reinforce the importance of political ideology. In the common $N = 66$ sample, government experience and office size explain only 2.6% of the variance in perceived risks ($R^2 = .026$). Adding AI familiarity and use increases explained variance to 19.1% ($R^2 = .191$), an increase of approximately 16.5 percentage points, $F\text{-change} = 6.23, p = .003$. Adding ideology subsequently raises R^2 from .191 to .431, representing an additional 24.0 percentage points of explained variance, $F\text{-change} = 25.25, p < .001$.

The contrast between the benefits and risks models is therefore pronounced. Technological engagement provides substantial explanatory leverage for understanding perceived benefits, whereas political ideology provides substantial explanatory leverage for understanding perceived risks. These results support treating perceptions of AI's utility and perceptions of AI's risks as related but analytically distinct dimensions rather than opposite ends of a single attitudinal continuum.

Organizational Implementation of AI

The final analysis examines whether the individual characteristics associated with perceptions of AI also predict reported organizational implementation. Among respondents providing a definitive answer, 18 reported that their office had implemented AI, while 52 reported that their office had not. After accounting for missing predictor values, the multivariate logistic regression includes 61 respondents.

Unlike the attitudinal models, none of the individual-level or organizational predictors reaches conventional levels of statistical significance. The estimated odds ratio for AI familiarity is 1.72 (95% CI [.78, 3.80]), while the odds ratio for frequency of AI use is 1.48 (95% CI [.76, 2.88]). Government experience is associated with an odds ratio of 1.04 (95% CI [.58, 1.87]), political ideology with an odds ratio of 1.33 (95% CI [.92, 1.92]), and office size with an odds ratio of 1.35 (95% CI [.81, 2.27]).

These null results should be interpreted cautiously. Only 18 respondents report organizational AI implementation, resulting in limited statistical power and relatively wide confidence intervals. The findings therefore should not be interpreted as evidence that individual or organizational characteristics are unrelated to implementation. Rather, the analysis provides no statistically distinguishable evidence within this sample that the same characteristics strongly associated with individual perceptions of AI also predict whether an office has implemented the technology.

The distinction between attitudes and implementation nevertheless raises an important possibility: organizational AI adoption may depend on institutional conditions beyond individual attitudes, including administrative capacity, procurement procedures, leadership decisions, legal

restrictions, technical infrastructure, cybersecurity requirements, and formal organizational policies.

Robustness and Diagnostics

Several diagnostic and sensitivity analyses were conducted to assess the stability of the multivariate findings. First, variance inflation factors are low across all predictors, ranging from approximately 1.09 to 1.46. These values provide no indication of problematic multicollinearity.

Second, Breusch–Pagan tests provide no evidence of heteroskedasticity in either primary OLS specification. The test is nonsignificant for the perceived-benefits model ($BP = 6.66$, $df = 5$, $p = .247$) and for the perceived-risks model ($BP = 2.58$, $df = 5$, $p = .765$). Nevertheless, HC3 heteroskedasticity-robust standard errors were estimated as an additional robustness check.

The central results remain substantively and statistically unchanged under robust inference. Frequency of AI use remains strongly associated with perceived benefits ($b = .635$, $HC3\ SE = .114$, $p < .001$), while political ideology remains strongly associated with perceived risks ($b = -.236$, $HC3\ SE = .059$, $p < .001$).

Cook's-distance diagnostics were also examined to determine whether the principal results were driven by a small number of influential observations. Using the conventional $4/N$ screening threshold, two observations exceeded the threshold in the benefits model and four exceeded it in the risks model. None approached a Cook's distance of 1, although the largest value in the risks model was .312. Because statistical influence alone does not provide a substantive justification for deleting otherwise valid observations, these cases were retained in the primary analyses.

As an additional sensitivity test, the fully specified models were re-estimated after excluding observations exceeding the 4/N threshold. The results remain highly consistent with the primary models. In the benefits model, AI use remains a strong positive predictor ($b = .603$, $\text{HC3 SE} = .112$, $p < .001$). The sensitivity model explains 52.7% of the variance in perceived benefits ($R^2 = .527$; adjusted $R^2 = .486$).

Similarly, political ideology remains a strong negative predictor of perceived risks after influential observations are excluded ($b = -.202$, $\text{HC3 SE} = .040$, $p < .001$). The corresponding sensitivity model explains 55.3% of the variance in perceived risks ($R^2 = .553$; adjusted $R^2 = .513$).

Taken together, these diagnostic and sensitivity analyses indicate that the paper's central findings are not attributable to heteroskedasticity, multicollinearity, or a small number of influential observations. Across specifications, the same substantive distinction emerges: frequency of AI use is most strongly associated with perceived governmental benefits, whereas political ideology is most strongly associated with perceived governmental risks.

Discussion

This study began with a simple question: How do people working within government understand artificial intelligence as it enters the institutions in which they work? The results suggest that there is no single answer because assessments of AI's benefits and its risks are structured differently.

The clearest finding concerns technological experience. Respondents who use AI more frequently perceive considerably greater potential for AI to improve government. The

relationship is large, survives the inclusion of ideology and institutional characteristics, and remains robust across alternative diagnostic specifications. General familiarity with AI, however, has no comparable relationship.

This distinction matters. Exposure to discussion about artificial intelligence is not equivalent to actually working with it. Direct users encounter concrete applications that can demonstrate the technology's utility: summarizing information, producing drafts, organizing material, assisting research, or accelerating routine tasks. The results are therefore consistent with an experiential account of technological utility in which perceptions of benefits are closely associated with practical use.

The direction of causality remains uncertain. Individuals who already believe AI is useful may choose to use it more frequently. Alternatively, repeated use may lead individuals to discover applications that increase their assessments of its value. Most plausibly, the relationship may be reciprocal. Longitudinal or experimental research is needed to distinguish among these possibilities.

Perceptions of risk tell a different story. Once ideology is considered, frequency of AI use no longer independently explains whether respondents perceive AI as threatening bias, inequality, transparency, public trust, or institutional preparedness. Political ideology instead becomes the dominant correlate.

This finding illustrates why attitudes toward governmental AI should not be represented simply as technological optimism versus technological pessimism. A respondent can perceive substantial administrative utility while simultaneously believing AI requires extensive

safeguards. Likewise, skepticism about particular risks does not necessarily imply skepticism about AI's technical capabilities.

The results instead suggest that assessments of utility and risk operate as related but distinct dimensions. The two indices are negatively correlated ($r = -.39$), but the relationship is far from sufficiently strong to treat them as opposites. Their predictors further demonstrate their distinctiveness.

The political character of risk perception also has implications for AI governance. Public institutions attempting to construct AI policies may encounter disagreements that are not primarily disagreements about technical performance. Officials may agree that AI can summarize documents or accelerate information processing while disagreeing substantially about the severity of algorithmic bias, misinformation, transparency problems, or institutional unpreparedness.

This distinction may complicate efforts to develop shared governance frameworks. Technical demonstrations of AI's usefulness may alter assessments of utility without resolving political disagreement over acceptable risk. Conversely, rules designed to mitigate risks may be evaluated through ideological frameworks that extend beyond the technical characteristics of the systems themselves.

The organizational-adoption results introduce another important distinction. Individual attitudes do not necessarily translate directly into institutional action. Although AI use strongly predicts perceived benefits and ideology strongly predicts perceived risks, neither clearly predicts whether respondents report organizational AI implementation.

Government adoption may therefore operate at a different level of analysis. Procurement systems, budgets, cybersecurity requirements, leadership decisions, legal restrictions, administrative capacity, vendor relationships, workforce rules, and formal AI policies may determine institutional adoption even when individual employees perceive considerable utility.

Future research should consequently distinguish between individual technological orientation and organizational technological capacity rather than treating the former as a proxy for the latter.

Limitations

First, the survey is relatively small. Although the attitudinal models identify large and statistically robust relationships, the sample limits the number of predictors that can reasonably be incorporated. The organizational-adoption analysis is particularly constrained because only 18 respondents report implementation.

Second, the data are cross-sectional. The models identify associations rather than causal effects. The strong relationship between AI use and perceived benefits, for example, cannot establish whether AI use changes attitudes or whether favorable attitudes encourage use.

Third, the analysis relies on respondents' reports of organizational AI implementation. Future work could supplement survey responses with organizational policies, procurement records, administrative directives, or other measures of actual implementation.

Fourth, the survey captures perceptions of AI rather than evaluations of specific AI systems. Attitudes may differ depending on whether AI is used for internal administrative work,

constituent communication, policy analysis, benefits administration, enforcement, or decisions affecting individual rights.

Finally, the respondents represent individuals with state-government experience rather than a probability sample of all state-government employees. The findings should therefore not be interpreted as population estimates of attitudes across American state government.

These limitations are important, but they do not diminish the central contribution of the analysis: identifying a theoretically meaningful difference in the correlates of perceived technological benefits and perceived technological risks.

Conclusion

Artificial intelligence is entering government institutions before those institutions have fully resolved what the technology should be permitted to do, what safeguards should govern its use, or how its consequences should be evaluated. Understanding this transition requires examining not only formal regulation and organizational adoption but also the people operating within government.

This study demonstrates that government professionals' assessments of AI are multidimensional. Experience structures perceived utility; ideology structures perceived risk.

Respondents who use artificial intelligence more frequently perceive substantially greater governmental benefits, independent of their ideological orientation, professional experience, familiarity with AI, and office size. Political ideology, meanwhile, is strongly associated with perceptions of AI-related risks: more conservative respondents express substantially lower levels of concern, even after technological experience is considered.

These relationships remain robust to alternative standard errors, nested specifications, and influential-observation sensitivity tests. At the same time, neither technological engagement nor ideology clearly predicts reported organizational adoption, suggesting that institutional implementation cannot be understood simply by aggregating individual attitudes.

The distinction has consequences for the governance of emerging technologies. Increasing governmental exposure to AI may alter perceptions of what the technology can accomplish without producing consensus about what risks it poses or which safeguards are necessary. Governments may therefore confront a future in which agreement over AI's practical usefulness coexists with persistent political disagreement over its consequences.

As artificial intelligence becomes increasingly embedded in public institutions, the central governance challenge will not simply be whether governments adopt AI. It will be how institutions reconcile technological experience, political disagreement, and administrative responsibility when determining the conditions under which AI should govern—and the conditions under which it should not.

References

Ahn, Michael J., and Yu-Che Chen. 2022. “Digital Transformation toward AI-Augmented Public Administration: The Perception of Government Employees and the Willingness to Use AI in Government.” *Government Information Quarterly* 39 (2): 101664.

<https://doi.org/10.1016/j.giq.2021.101664>.

Bullock, Justin B. 2019. “Artificial Intelligence, Discretion, and Bureaucracy.” *The American Review of Public Administration* 49 (7): 751–761. <https://doi.org/10.1177/0275074019856123>

Criado, J. Ignacio, and Lucia O. de Zarate-Alcarazo. 2022. “Technological Frames, CIOs, and Artificial Intelligence in Public Administration: A Socio-Cognitive Exploratory Study in Spanish Local Governments.” *Government Information Quarterly* 39 (3): 101688.

<https://doi.org/10.1016/j.giq.2022.101688>

DiMaggio, Paul J., and Walter W. Powell. 1983. “The Iron Cage Revisited: Institutional Isomorphism and Collective Rationality in Organizational Fields.” *American Sociological Review* 48 (2): 147–160. <https://doi.org/10.2307/2095101>.

Douglas, Mary, and Aaron Wildavsky. 1982. *Risk and Culture: An Essay on the Selection of Technological and Environmental Dangers*. Berkeley: University of California Press.

Haesevoets, Tessa, Bram Verschuere, Alain Van Hiel, and Arne Roets. 2026. “Public Servants' Attitudes toward the Use of Artificial Intelligence in Policymaking: An Experimental Study.” *Government Information Quarterly* 43 (2): 102134. <https://doi.org/10.1016/j.giq.2026.102134>.

Haesevoets, Tessa, Bram Verschuere, and Arne Roets. 2025. "AI Adoption in Public Administration: Perspectives of Public Sector Managers and Public Sector Non-Managerial Employees." *Government Information Quarterly* 42 (2): 102029.

<https://doi.org/10.1016/j.giq.2025.102029>

Kahan, Dan M., Ellen Peters, Maggie Wittlin, Paul Slovic, Lisa Larrimore Ouellette, Donald Braman, and Gregory Mandel. 2012. "The Polarizing Impact of Science Literacy and Numeracy on Perceived Climate Change Risks." *Nature Climate Change* 2 (10): 732–735.

<https://doi.org/10.1038/nclimate1547>.

Lipsky, Michael. 1980. *Street-Level Bureaucracy: Dilemmas of the Individual in Public Services*. New York: Russell Sage Foundation.

Madan, Rohit, and Mona Ashok. 2023. "AI Adoption and Diffusion in Public Administration: A Systematic Literature Review and Future Research Agenda." *Government Information Quarterly* 40 (1): 101774. <https://doi.org/10.1016/j.giq.2022.101774>

McCright, Aaron M., and Riley E. Dunlap. 2011. "The Politicization of Climate Change and Polarization in the American Public's Views of Global Warming, 2001–2010." *Sociological Quarterly* 52 (2): 155–194. <https://doi.org/10.1111/j.1533-8525.2011.01198.x>.

Meier, Kenneth John. 1975. "Representative Bureaucracy: An Empirical Analysis." *American Political Science Review* 69 (2): 526–542. <https://doi.org/10.2307/1959084>.

Meier, Kenneth J., and John Bohte. 2001. "Structure and Discretion: Missing Links in Representative Bureaucracy." *Journal of Public Administration Research and Theory* 11 (4): 455–470. <https://doi.org/10.1093/oxfordjournals.jpart.a003511>.

Riccucci, Norma M., and Gregg G. Van Ryzin. 2017. “Representative Bureaucracy: A Lever to Enhance Social Equity, Coproduction, and Democracy.” *Public Administration Review* 77 (1): 21–30. <https://doi.org/10.1111/puar.12649>.

Robinson, Nicholas. 2026. “Open to Open-Source AI? Navigating AI Model Choice in Public Sector Agencies.” *Government Information Quarterly* 43 (2): 102133. <https://doi.org/10.1016/j.giq.2026.102133>

Rodriguez Müller, A. Paula, Luca Tangi, and Amandine Lerusse. 2025. “Understanding the Adoption of Artificial Intelligence in Local Government Decision-Making: The Influence of Institutional Pressures and Managerial Perceptions.” *Public Administration*. <https://doi.org/10.1111/padm.70033>.

Schiff, Daniel S., Kaylyn Jackson Schiff, and Patrick Pierson. 2022. “Assessing Public Value Failure in Government Adoption of Artificial Intelligence.” *Public Administration* 100 (3): 653–673. <https://doi.org/10.1111/padm.12742>.

Selten, Friso, and Bram Klievink. 2024. “Organizing Public Sector AI Adoption: Navigating between Separation and Integration.” *Government Information Quarterly* 41 (1): 101885. <https://doi.org/10.1016/j.giq.2023.101885>

Sowa, Jessica E., and Sally Coleman Selden. 2003. “Administrative Discretion and Active Representation: An Expansion of the Theory of Representative Bureaucracy.” *Public Administration Review* 63 (6): 700–710. <https://doi.org/10.1111/1540-6210.00333>.

Uttermark, Matthew J. 2024. “Administrator Political Orientations and Views on Federal Encroachment.” *Administration & Society* 56 (1): 18–49.

<https://doi.org/10.1177/00953997231199230>.

Young, Matthew M., Justin B. Bullock, and Jesse D. Leczy. 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.

<https://doi.org/10.1093/ppmgov/gvz014>.

Zhu, Liang, and Tianyi Gao. 2026. “Hybrid Images of Generative AI: A Q Methodological Study of Civil Servants’ Perceptions.” *Government Information Quarterly* 43 (1): 102113.

<https://doi.org/10.1016/j.giq.2026.102113>

Zuiderwijk, Anneke, Yu-Che Chen, and Fadi Salem. 2021. “Implications of the Use of Artificial Intelligence in Public Governance: A Systematic Literature Review and a Research Agenda.”

Government Information Quarterly 38 (3): 101577. <https://doi.org/10.1016/j.giq.2021.101577>