

THE RISE OF LLMS IN BUREAUCRACY AND MILITARY DECISION-MAKING AND THE CYBERSECURITY IMPERATIVE

Gloria Shkurti ÖZDEMİR*
ORCID ID: 0000-0001-8626-9761

Declaration*

Abstract

Large Language Models (LLMs) are rapidly transforming decision-making processes across bureaucratic and military institutions. Their ability to synthesize data, simulate complex scenarios, and generate real-time strategic insights is driving adoption in public sector settings, with initiatives like OpenAI’s “ChatGPT Gov” already deployed across U.S. federal agencies. However, the integration of LLMs into core governance and defense infrastructures introduces profound risks. Beyond technical concerns such as data poisoning, adversarial attacks, and insider misuse, these models also raise normative challenges, escalation bias in military applications, erosion of institutional accountability, and dependency on opaque corporate infrastructures. This article critically examines the operational use of LLMs in bureaucratic and military domains, analyzes the cybersecurity and geopolitical risks they pose, and frames their adoption within broader debates on technological sovereignty, corporate power, and data colonialism. Lastly, the article provides several recommendations that can offer some insight into how states, particularly middle and regional powers, can reclaim agency, enhance institutional resilience, and push for more effective regulatory frameworks in the face of accelerating LLM integration and corporate dominance.

Keywords: AI, Decision Making, Foreign Policy, Military, Threats, Cybersecurity

Introduction

Since the public release of ChatGPT in late 2022, not only has artificial intelligence undergone a pivotal transformation, but so too has the global landscape in which humans work, govern, and make decisions. The arrival of advanced large language models (LLMs)

* Director of Emerging Technologies and AI (ETAI) Center at Khazar University, Azerbaijan and Researcher at SETA Foundation, Türkiye

* The author acknowledges the assistance of ChatGPT model in language editing. All analysis and conclusions are solely the author’s responsibility.



marked a historic moment, fueling discussions around the “democratization of technology,” as once-exclusive computational capabilities became widely accessible to the public (Shkurti Özdemir, AB, Yapay Zekâ Düzenlemesinde Küresel Lider Olabilecek mi?, 2024).

Yet, as the initial excitement of open-access AI gave way to more critical reflection, the dual-use nature of these technologies became evident. While LLMs can empower individuals and increase productivity, they also hold strategic significance for governments and militaries. It was only a matter of time before their integration into the public sector and defense infrastructures began.

Today, the use of LLMs in governance is no longer speculative. Across the globe, bureaucratic agencies and defense institutions are actively experimenting with and deploying LLMs to automate routine functions, assist in policy analysis, and streamline administrative tasks. However, the most consequential shift lies not in automating clerical work, but in the gradual incorporation of LLMs into decision-making processes themselves, both in civil administration and in military contexts.

The appeal of LLMs stems from their capacity to scale cognitive labor and process vast amounts of information rapidly. Yet, their integration into core governance functions also introduces new vectors for cybersecurity threats, systemic vulnerabilities, and ethical concerns (Karaguezian, 2024, pp. 243-244). As these systems begin to shape high-stakes outcomes, the risks of bias, manipulation, and loss of institutional accountability grow accordingly.

This paper explores the dual-edged implications of LLM adoption in state systems. Specifically, it analyzes the ways in which LLMs are being operationalized within bureaucratic and military domains and assesses the emergent cybersecurity threats associated with their deployment.

Bureaucratic Adoption of Large Language Models

Bureaucracy, at its core, emerged as a response to the growing need for systematic information management. One of the earliest manifestations of this can be traced back to ancient Mesopotamia, where written records on clay tablets were used to document royal assets and economic transactions. However, as the volume of such records expanded, the challenge of organizing, storing, and retrieving critical information became increasingly apparent. Bureaucracy evolved as an institutional mechanism to address these problems,



structuring administrative functions and enabling information governance (Harari, 2024, pp. 45-48). Over time, bureaucracies adapted to successive waves of technological transformation, from paper-based filing systems to digital databases. Today, amid the exponential growth of data, we are witnessing another pivotal shift: the integration of advanced technologies such as large language models. These models are not merely tools for digitization, but catalysts for reimagining how bureaucratic systems process information, make decisions, and interact with the public.

As LLMs increasingly move from the periphery to the center of technological ecosystems, their adoption within public administration has accelerated. What began as experiments in automating low-level clerical tasks has evolved into a much deeper transformation of the bureaucratic imagination. LLMs, at the beginning, were used as conversational agents, i.e. chatbots or virtual assistants, for different public-facing services (Lund & Ting, 2023) or as tools for the summarization and translation of documents (Council of the European Union, 2023, p. 9). However, currently they are being considered, and in some cases even actively integrated, into different tasks that can inform or impact administrative decision-making.

However, this intensifying integration of LLMs within the administrative decision-making brings several uncertainties with it. Specifically, when the cognitive labor previously done by human administrators is delegated to opaque and probabilistic systems such as LLMs this erodes the discretionary space that was reserved just for the human administrator. Even more importantly, such a delegation challenges directly the well-established normative foundations within the public sector, including here the fact that decisions need to be transparent, justifiable, and aligned with the public interest. Within this context, the concern becomes higher when we acknowledge that the decision-making in bureaucracy includes matters of great national importance, such as foreign policy, military interventions, and in some cases even decisions relations to the nuclear command. These risks augment further when we consider not only the threat coming from the models themselves but also their growing exposure to possible cybersecurity threats and from a global affairs perspective, the geopolitical dependence of the states that cannot develop these models on the foreign-owned AI systems. Within this framework, when we consider the fact that LLMs are transitioning from simple tools of administrative convenience towards important actors within the decision-making chain, it can be said that this marks a very important turning point requiring great oversight.



Real-World Deployments

While integrating LLMs within the public sector was considered to happen maybe later in the future, now their deployment, even for decision making purposes, is no longer speculative. Today we can speak about the integration of LLMs with prominent variations in scope, ambition, and institutional design across different national domains. Several governments have begun experimenting with or formally deploying LLMs within their administrative systems. A particularly significant case is that of the United States. In October 2024, the Biden administration released a policy directive urging U.S. national security institutions to prioritize the adoption of artificial intelligence technologies. The memo emphasized the importance of leveraging AI models and related tools across federal agencies, particularly within national security operations (The White House, 2024). Within this framework, soon after Trump assumed presidency a strategic partnership between OpenAI and public institutions has given rise to ChatGPT Gov, a customized version of ChatGPT designed specifically for governmental use. Launched in early 2025, ChatGPT Gov allows U.S. agencies to access OpenAI's frontier models within secure, self-managed cloud environments that adhere to federal cybersecurity standards (OpenAI, 2025).

38

The initiative marks a qualitative shift in how public bureaucracies conceptualize AI integration, not merely as an efficiency tool but as a structural component of digital governance. According to OpenAI, since 2024, more than 90,000 users across over 3,500 federal, state, and local government entities have exchanged upwards of 18 million messages using ChatGPT Enterprise to assist with their daily workflows (OpenAI, 2025). These use cases span a wide spectrum, from document drafting and administrative support to data analysis and internal communication.

Unlike commercial versions, ChatGPT Gov is deployed within government-controlled Microsoft Azure infrastructures, including both commercial and government community cloud environments. This architecture allows agencies to retain sovereignty over key aspects such as data privacy, security protocols, and compliance frameworks, offering a model of AI adoption that seeks to balance innovation with institutional risk management.

The U.S. model reflects not only technological ambition but also a growing recognition that future governance may hinge on the controlled, context-specific deployment of advanced



language models. Yet, at the same time, as it will be discussed below, it raises critical questions about long-term dependence on private sector actors for the core infrastructure of public administration.

Another notable example of LLM integration in the bureaucratic sphere, though currently in the research and pilot phase, is the Indonesian Ministry of Finance's development of KemenkeuGPT. This domain-specific language model has been trained on a substantial corpus of national economic data, fiscal policy frameworks, and regulatory documents, enriched by iterative expert feedback from within the Ministry itself. While not yet deployed for operational use, KemenkeuGPT is envisioned as a strategic decision-support system, designed to facilitate policy simulations, generate tailored financial reports, and enhance internal modeling and forecasting capacities (Febrian & Figueredo, 2024). Its development reflects a deliberate effort to build sovereign AI capabilities tailored to the unique informational demands of a specific governmental domain. As such, KemenkeuGPT offers an important contrast to off-the-shelf LLM deployments, representing a model of targeted, context-sensitive AI integration that seeks to retain institutional control over core knowledge infrastructures.

A third example regarding the integration of LLMs in public administration is that of "Pubbie," a project developed by Canada's National Research Council (NRC). Pubbie, which was started as a part of a broader AI program launched by NRC in May 2024, is currently in the experimental phase and is designed to support government operations, especially in the area of research and innovation policy. Specifically, by searching vast academic and technical databases, spotting new fields with scientific value, and matching national research funding with strategic priorities, the model is intended to support the civil servants. Furthermore, Pubbie's main function is to improve the evidence-based decision-making within NRC by offering timely and contextualized insights, this way showing how LLMs can be effective when used for high-level policy coordination (Liu, Geng, & Hart, 2025). It is also important to state the fact that this model is part of a larger initiative in Canada, namely the Artificial Intelligence Strategy for the Federal Public Service, launched in March 2025. At some extend similar to the above-mentioned initiative by the U.S., the strategy in Canada establishes the main frameworks for the responsible integration of AI into federal agencies, placing a focus on openness, responsibility, and creativity in service provision (Government of Canada, 2025).



Lastly, another example of the integration of LLM within the bureaucratic domain is that of LLaMandement which is used in France. This model was designed to automate the analysis and summarization of parliamentary documents. LLaMandement improves the effectiveness and transparency of the parliamentary workflows and at the same time it reduces the administrative load on legislative staff. This way, by speeding up the processing and accessibility of legislative texts, the model helps to create a more responsive lawmaking process (Gesnouin et al., 2024). Concurrently, it can be stated that the adoption of this model within the French bureaucracy is a reflection of France's broader strategic objective that aims to achieve digital sovereignty. As a result, the LLaMandement represents how these models can be used not only to help the bureaucratic processed but when seen from the global perspective they are also seen as instruments of national autonomy.

The Military Turn: LLMs and the Rise of Agentic Warfare

Focusing on the military domain, the adoption of AI and LLMs especially within the military operations reflects a shift and change in the character of the warfare (Shkurti Özdemir, 2024). Considering the fact that LLMs can process large amount of data at a much faster rate than the human operators can, these models can then make decisions faster, can allocate resources more efficiently, and at the same time can improve the communication within the military hierarchies (Rivera, et al., 2024, p. 1). According to Puscas, these models can be used for several purposes including strategic simulations, wargaming scenarios, operational planning, the creation of multiple courses of action, and real-time threat identification (Puscas, 2024, p. 15). Their capacity to automate scenario development and streamline decision support systems makes them increasingly indispensable in high-tempo, complex conflict environments.

While traditionally framed as tools for textual generation and summarization, LLMs are now being embedded within agentic AI systems, autonomous frameworks capable of perception, decision-making, and dynamic interaction with real-world data (Jensen, Tadross, & Strohmeyer, 2025). This shift signals the emergence of what is increasingly referred to as *agentic warfare*, a new paradigm in which AI agents actively shape the tempo and direction of conflict across all domains.

States, now aware of the accelerating pace of the modern warfare, where the responses within military operations need to occur within seconds, are highly investing in AI adoption in military domain in order to avoid being strategically outmaneuvered. Considering also its



technological superiority, the U.S. stands out as a leader in terms of its efforts to incorporate LLMs and agentic AI systems into its defense infrastructure. The U.S. Department of Defense (DoD), in particular, is trying to take advantage of this transformation by integrating LLMs into different critical military infrastructure. The Pentagon's 2023 Data, Analytics, and Artificial Intelligence Strategy envisions AI-enabled systems as vital to accelerating decision-making and enhancing the precision of command structures (Farnell & Coffey, 2024). In practical terms, LLMs are now tested for operational roles ranging from scenario planning and intelligence analysis to cyber-operations and even command-and-control functions. Experiments within the DoD have shown that LLMs can digest vast troves of classified data and return actionable insights within minutes, a process that previously took human staff days to accomplish. As one military officer put it after a successful trial, "We are learning that this is possible for us to do" (Manson, 2023).

These developments have been catalyzed also by OpenAI's controversial January 2024 decision to lift restrictions on the military use of its models, including applications linked to weapons development and warfare (Csernaton, 2024). This move underscores a broader trend: the erosion of ethical guardrails on AI deployment and the rise of a new form of corporate nonstate sovereignty. In the absence of robust international norms governing military AI, private firms like OpenAI and Scale AI are increasingly shaping the battlefield, not merely supplying it. It is important to state at this point that with the arrival of Trump in the White House, the application of AI and especially LLMs in the military is going to escalate and proliferate further (Shkurti Özdemir & Ustun, 2024; Shkurti Özdemir, 2025a).

The strategic implications of agentic warfare are far-reaching. In this new paradigm, LLM-powered agents do not simply process text; they simulate escalation scenarios, interact with live databases, make strategic recommendations, and coordinate across operational units. They serve as cognitive engines embedded within AI warfighters, agents that monitor global signals, detect anomalies, and generate response plans at machine speed. This level of integration fundamentally transforms how war is planned, initiated, and potentially deterred.

Agentic warfare is not merely a futuristic concept. It is already unfolding through the testing of systems like Scale AI's *Donovan*, Microsoft's deployment of OpenAI models on Azure Government Cloud, and Anduril and Palantir's development of autonomous decision-making platforms. These systems are designed to execute joint force operations, interface with



sensors, and manage munitions, all while adapting in real time to fluid operational environments.

The conceptual leap lies in replacing static military doctrine with dynamic, AI-informed strategies. Agents now simulate entire campaigns, weigh risk trade-offs, and propose novel options grounded in both historical precedent and live data streams. This is not just about speed; it is about strategic foresight. An agentic military force may detect adversary movements before human analysts can process the signals, preempting escalation and preserving advantage (Jensen, Tadross, & Strohmeyer, 2025).

As the world enters this new era, the strategic imperative is clear: failure to embrace agentic warfare may relegate states to a reactive posture, outpaced by adversaries with more agile and autonomous capabilities. Yet doing so responsibly demands new doctrine, oversight mechanisms, and international agreements that balance innovation with restraint.

Strategic, Cybersecurity, and Geopolitical Risks

As the adoption of LLMs expands across bureaucratic and military domains, the associated risks become increasingly salient, many of which extend beyond technical challenges and into normative, institutional, and geopolitical territory. While LLMs promise enhanced efficiency and cognitive support, their deployment in governance and defense introduces vulnerabilities deeply embedded in the structure, ownership, and alignment of the models themselves. This section explores three key categories of risk: cybersecurity and data governance, deployment bias and strategic misalignment, and geopolitical dependency under a new paradigm of technopolitical power.

Deployment Bias, Strategic Misalignment, and the Escalation Risk

The risk of deployment bias, using LLMs in scenarios beyond their design parameters, is especially problematic in the context of state governance and international affairs (Schwartz, et al., 2022). Most LLMs are trained and evaluated on benchmarks focused on reasoning, coding, or summarization. These metrics do not capture the complex, value-laden nature of political or strategic decision-making. Specifically, there is no verifiable truth in the domain of diplomacy and defense. Therefore, the lack of this verifiable truth means that decisions such as escalating a conflict, imposing sanctions, or intervening diplomatically are inherently subjective and politically charged. When considered like that it is obvious that there is an incompatibility between the task that the LLMs are intended to be applied and the real



capabilities of these models. The majority of current model evaluations ignore subjective decision-making contexts where results rely on social goals or institutional norms in favor of concentrating on reasoning abilities and task execution. However, as mentioned above, in governance and international affairs, generally there is no 'correct' answer, therefore making reliance on LLMs very dangerous (Jensen, et al., 2025, p. 2).

Several studies prove indeed this incompatibility of the LLM's task and their real capabilities. For example, a study conducted in 2025 reached in the conclusion that during several scenario simulations, models such as LLaMA 3.1 8B Instruct, Gemini 1.5 Pro-002, and Qwen2 72B typically suggest more escalatory policies. Furthermore, based also on the data they were trained on, these models displayed geographical biases. Specifically, these models advocated less aggressive positions toward China or Russia and more interventionist tactics for nations such as the United States or the United Kingdom (Jensen, et al., 2025, p. 2). As a result of these biases, it would be fair to raise concerns about fairness, alignment, and the possibility for algorithmically induced conflict.

Furthermore, similar to the study conducted by Jensen, et al., another study conducted by Riviera, et al. reached parallel results, again emphasizing the fact that LLMs can display erratic and occasionally violent escalation patterns when used within conflict simulation scenarios, including here nuclear decision-making (Riviera, et al., 2024). Within this context, it is necessary to emphasize that when we take into consideration the vague algorithmic reasoning and the possible sidelining of human judgment there is a high possibility has the potential to increase the risk of catastrophic conflict escalation in high-stakes situations, especially those involving nuclear decision-making.

Cybersecurity and Data Governance

When we talk about the application of technologies such as AI or LLM in the bureaucracy and military domain, the cybersecurity, and the challenges posed to it, become an unavoidable concern. Technically speaking, LLMs have the capabilities to memorize and repeat sensitive data provided in their training sets, therefore directly increasing the risk of information leakage. This is very concerning especially when LLMs are exposed to unredacted internal documents or private conversations, which are frequent in fields like national security, law enforcement, and taxation. Furthermore, adversarial prompt can also take advantage of the possible weaknesses and therefore lead to the exposure of confidential information, proprietary knowledge, or socially offensive material.



Another issue that may emerge is related to the anonymization of data. Specifically, the public official's interaction with LLMs has the capability to create new data streams that be used to retrain future models if they are not appropriately anonymized. For instance, any user data from ChatGPT can be incorporated into OpenAI's continuing training cycles unless agencies choose not to. As it may be understood, this may result in the unintentional revealing of makes sensitive discussions, strategic planning, or legal interpretations.

Lastly, the attack vectors need to be taken into consideration. Malicious actors can modify outputs or retrieve training data by using different strategies including prompt injection, model inversion, or synthetic querying. In the bureaucracy realm, where the IT infrastructures and generally underfunded or outdated, through the attack vectors, LLMs can be used to direct the development of malware, presenting a significant risk. Moreover, the dependency on cloud-hosted models and private vendor-managed APIs worsens the problem as it reduces governmental control and creates uncertainty regarding data sovereignty.

Geopolitical Dependency, Corporate Power, and Technological Sovereignty

When we discuss the LLMs adoption within the bureaucracy and military, one of the most important threats is the increasing influence of Big Tech companies over sovereign affairs. LLMs are highly resource-intensive systems developed by a small number of private actors. As of now, only a few firms, including OpenAI, Google DeepMind, Anthropic, and Baidu, possess the computational infrastructure, proprietary data, and technical talent to develop frontier models.

This dynamic creates two parallel dependencies. First, even technologically advanced states such as the United States are increasingly reliant on private firms for access to and control over LLM capabilities. For example, the U.S. government's collaboration with OpenAI on ChatGPT Gov illustrates a deeper entanglement between public institutions and corporate platforms. While such partnerships provide cutting-edge tools, they also allow private firms to gain privileged access to massive volumes of sensitive governmental data, which can be used to refine commercial models, shape policy discourse, or even nudge administrative behavior. In effect, governments risk becoming junior partners in a technocratic order governed not by democratic deliberation but by platform logics. If we focus especially on agentic warfare for example, the reliance on corporate AI infrastructure introduces a new dependency dynamic. Firms like OpenAI and Scale AI are now de facto defense partners with privileged access to sensitive data, shaping the capabilities and limitations of military force projection. In this



sense, agentic warfare is both a technological and political transformation, reshaping the relationship between states, private actors, and the conduct of war.

Second, governments that are unable to develop their own models, particularly those in the Global South or among mid-sized economies, become dependent on foreign vendors and, indirectly, on the geopolitical priorities of the states where these vendors are based. This dual dependency can severely constrain policy autonomy and expose national infrastructure to influence or coercion.

This dynamic resonates with emerging critiques of technofeudalism (Varoufakis, 2023), the idea that contemporary digital capitalism is marked by a concentration of infrastructural power in the hands of tech oligopolies that extract rents from data, labor, and public resources, or even that of data colonialism (Mejias & Couldry, 2024), the extraction and appropriation of personal and institutional data by corporate platforms, mirroring historical patterns of colonial resource exploitation, but now operating through algorithmic infrastructures and transnational data flows. The reliance on LLMs hosted by proprietary cloud infrastructures fits this pattern. States are not only consumers of corporate AI but also de facto data suppliers, reinforcing the centrality of big tech firms in shaping the governance of the digital age (Akyesilmen, 2023).

Moreover, the opacity of proprietary models further complicates oversight. OpenAI, for example, no longer discloses key architectural and training data for its latest models, making external auditing impossible. Without transparency, states cannot verify whether these systems uphold democratic principles, remain neutral in geopolitical conflicts, or embed unwanted ideological perspectives.

In sum, LLMs are not neutral infrastructure. Their integration into critical decision systems should not be viewed solely through the lens of utility or innovation. Rather, it must be approached as a question of political power, institutional trust, and long-term sovereignty. States must respond through a combination of regulatory development, public investment in open-source AI, and new international norms that align AI deployment with democratic accountability and strategic autonomy.

Recommendations

The integration of LLMs into bureaucratic and military infrastructures signals a profound transformation in the architecture of governance and warfare. Yet, this transformation has



outpaced the ability of regulatory institutions to respond. At present, there is a conspicuous absence of comprehensive legal and ethical frameworks capable of managing not only the systemic risks posed by LLM deployment but AI in general. Global digital governance remains fragmented, slow-moving, and largely reactive. As demonstrated during the 2024 AI Paris Summit, efforts to build a coordinated global response towards responsible AI have been hampered not only by geopolitical competition but also by the strategic lobbying of Big Tech firms, whose interests often conflict with calls for stronger public oversight (Shkurti Özdemir, 2025b).

Indeed, as it was seen also under Biden Administration, these Big Tech companies, when pushed by the states towards more regulations, they try to shape the regulatory agenda itself, contributing to draft frameworks, influencing policy timelines, and pushing for self-regulation. In this context, the race for AI governance is being lost not because states are unaware of the risks, but because the very architecture of global governance remains vulnerable to corporate capture. The asymmetry of technical capacity and infrastructural control means that, in many ways, the rules are being written by those who own the models.

Nevertheless, this institutional stagnation should not be cause for resignation. On the contrary, it highlights the urgent need for middle and regional powers, such as Türkiye, Indonesia, South Korea, and Brazil, to step forward and advocate for more assertive regulatory initiatives. These actors are uniquely positioned to push for a more pluralistic and equitable AI order, one that balances innovation with democratic values and strategic sovereignty.

In light of these challenges, there can be proposed several recommendations:

One of the biggest problems with the application of new technologies is generally related to the lack of the oversight bodies. For this reason, it is necessary that states focus on the establishment of these bodies before it becomes more difficult to control the adaptation of the newly emerging technologies, especially LLMs. These institutions should be responsible for the auditing and regulating the integration of LLMs, especially in terms of governance and defense. These organizations should focus of guaranteeing openness, human supervision, and conformity to moral and constitutional requirements.

Currently one of the most discussed issues revolves around the use of closed-source and open-source AI models. Within this context, it is necessary that states focus on the developments of



sovereign and open-source models that would serve best the public interest and would reduce the dependency on external actors including here other states or Big Techs.

The biggest risk with the adoption of LLMs emerge in the defense domain; therefore, it is important that AI system must always operation under strict human-in-the-loop control. There should be clear protocols and regulations that prevent the autonomous escalations, especially in regard to decisions related to nuclear posture of active conflict engagement.

As mentioned above cybersecurity is an issue that automatically comes to the fore when AI Models such as LLMs are applied in sensitive domains such as bureaucracy and defense. Within this context, it is necessary that government update their cybersecurity standards in order to handle the unique risks posed by LLMs, i.e. data leakage, prompt injection, and model inversion attacks. It is also important there the government create protocols that prohibit the use of the public sector data for commercial training of the LLMs models.

Conclusion

In this algorithmic age, the integration of LLMs in the bureaucracy and military domain symbolizes a revolutionary reorganization of authority and governance. While previously LLMs were tools of efficiency and automation, LLMs are now integrated into decision-making architectures that may control anything from taxation to nuclear escalation. Without any doubt, this brings both advantages and risks. On the one hand, LLMs have the potential to improve state responsiveness, accelerate cognitive labor, and improve institutional foresight. However, on the other hand, these models bring unique challenges on issues that are mainly political and normative, including here bias, opacity, dependency, and conflict escalations.

As this paper has argued, the adoption of LLMs in bureaucracy and decision making is changing the epistemic foundations of the governance itself. At the same time, the use beginning of the so-called agentic warfare signifies a fundamental change in the logic and conduct of war, as speed, simulation, and predictive modeling progressively replace discussion and diplomacy. Besides this, the dependency on proprietary infrastructures largely controlled by Big Tech companies emerges as another important issue, especially taking into consideration that their interest may not always coincide with that of the public.

Within this framework, national and international policy must focus especially on institutional accountability, strategic autonomy, and technological sovereignty. States need to be careful not to be fully dependent on external actors under a new regime of technopolitical extraction.



At this point, while banning the use and integration of LLMs into decision-making structures is not possible, it is important that states take the necessary steps and make sure that LLMs are governed by the protocols of innovation but at the same time by principles of justice, transparency, and public control.

References

Akyesilmen, N. (2023). Editorial Preface: The Age of Digital Empires: Transformation of International Politics. *Cyberpolitik Journal*, 8(16), vi-xi.

Council of the European Union. (2023). ChatGPT in the Public Sector – overhyped or overlooked? European Union.

Csernaton, R. (2024, July 17). Governing Military AI Amid a Geopolitical Minefield. Retrieved from Carnegie Europe: <https://carnegieendowment.org/research/2024/07/governing-military-ai-amid-a-geopolitical-minefield?lang=en>.

Farnell, R., & Coffey, K. (2024). AI's New Frontier in War Planning: How AI Agents Can Revolutionize Military Decision-Making. Cambridge: Belfer Center for Science and International Affairs.

Febrian, G. F., & Figueredo, G. (2024). KemenkeuGPT: Leveraging a Large Language Model on Indonesia's Government Financial Data and Regulations to Enhance Decision Making. arXiv.

Gesnoui et al., J. (2024). LLaMandement: Large Language Models for Summarization of French Legislative Proposals. Arxiv.

Government of Canada. (2025). AI Strategy for the Federal Public Service 2025-2027: Overview. Retrieved from Government of Canada: <https://www.canada.ca/en/government/system/digital-government/digital-government-innovations/responsible-use-ai/gc-ai-strategy-full-text.html>.

Harari, Y. N. (2024). Nexus: A Brief History of Information Networks from the Stone Age to AI. Great Britain: Fern Press.

Jensen, B., Reynolds, I., Atalan, Y., Garcia, M., Woo, A., Chen, A., & Howarth, T. (2025). Critical Foreign Policy Decisions (CFPD)-Benchmark: Measuring Diplomatic Preferences in Large Language Models. Arxiv.

Jensen, B., Tadross, D., & Strohmeyer, M. (2025, April 23). Agentic Warfare Is Here. Will America Be the First Mover? Retrieved from War on the Rocks: <https://warontherocks.com/2025/04/agentic-warfare-is-here-will-america-be-the-first-mover/>.

Karaguezian, S. (2024). AI and Cybersecurity: Navigating the Future of Warfare and Digital Defense. *Cyberpolitik Journal*, 9(18), 241-247.

Liu, S., Geng, M., & Hart, R. (2025). Exploring Generative AI Techniques in Government: A Case Study. Arxiv.



Lund, B., & Ting, W. (2023). Chatting about ChatGPT: How May AI and GPT Impact Academia and Libraries? Library Hi Tech News.

Manson, K. (2023, July 5). The US Military Is Taking Generative AI Out for a Spin. Retrieved from Bloomberg: <https://www.bloomberg.com/news/newsletters/2023-07-05/the-us-military-is-taking-generative-ai-out-for-a-spin>.

Mejias, U. A., & Couldry, N. (2024). Data Grab: The New Colonialism of Big Tech and How to Fight Back. Chicago: The University of Chicago Press.

OpenAI. (2025, January 25). Introducing ChatGPT Gov. Retrieved from OpenAI: <https://openai.com/global-affairs/introducing-chatgpt-gov/>.

Puscas, I. (2024). Large Language Models and International Security. Geneva: UNIDIR.

Rivera, J.-P., Mukobi, G., Reuel, A., Lamparth, M., Smith, C., & Schneider, J. (2024). Escalation Risks from Language Models in Military and Diplomatic Decision-Making. Stanford University Human-Centered Artificial Intelligence.

Schwartz, R., Vassilev, A., Greene, K. K., Perine, L., Burt, A., & Hall, P. (2022). Towards a Standard for Identifying and Managing Bias in Artificial Intelligence. NIST.

Shkurti Özdemir, G. (2024). AB, Yapay Zekâ Düzenlemesinde Küresel Lider Olabilecek mi? KRITER, 8(86). Retrieved from Kriter <https://kriterdergi.com/dis-politika/ab-yapay-zek-duzenlemesinde-kuresel-lider-olabilecek-mi>.

Shkurti Özdemir, G. (2024). Artificial Intelligence ‘Arms Dynamics’: The Case of The U.S. And China Rivalry. Istanbul: SETA.

Shkurti Özdemir, G. (2025a, July 5). Trump’ın Geri Dönüşü ve ABD’de Yenilenen Teknoloji-Savunma Bağı. Retrieved from Sabah: <https://www.sabah.com.tr/yazarlar/perspektif/gloria-shkurti-ozdemir/arsiv/getall>.

Shkurti Özdemir, G. (2025b, February 25). Paris AI Summit: Stage for power struggles, not regulation. Retrieved from Daily Sabah: <https://www.dailysabah.com/opinion/op-ed/paris-ai-summit-stage-for-power-struggles-not-regulation>.

Shkurti Özdemir, G., & Ustun, K. (2024). The Future of AI Policies in the US And Implications for Türkiye. Istanbul: SETA.

The White House. (2024, October 24). Memorandum on Advancing the United States’ Leadership in Artificial Intelligence; Harnessing Artificial Intelligence to Fulfill National Security Objectives; and Fostering the Safety, Security, and Trustworthiness of Artificial Intelligence. Retrieved from The Biden White House: <https://bidenwhitehouse.archives.gov/briefing-room/presidential-actions/2024/10/24/memorandum-on-advancing-the-united-states-leadership-in-artificial-intelligence-harnessing-artificial-intelligence-to-fulfill-national-security-objectives-and-fostering-the>.

Varoufakis, Y. (2023). Technofeudalism: What Killed Capitalism. New York: Melville House Publishing.

