

AI as Institutional Grammar: How National Contexts Shape the Legitimation of Machine Intelligence

David M. Boje

New Mexico State University

Accepted for publication, Management & Labour Studies, 2026

ABSTRACT

On 27 January 2025, Nvidia lost \$589 billion in market capitalisation in a single trading session — the largest single-day value destruction in stock market history — triggered by the question of whether a Chinese AI laboratory had replicated frontier intelligence at a fraction of the cost. The episode exposes a theoretical puzzle: why do actors in different national contexts tell strikingly different stories about artificial intelligence? This paper proposes that Alan Turing's nine objections to machine intelligence, first published in 1950, function as a shared institutional grammar whose selective deployment reflects the institutional logic of the national context in which AI actors operate. Three national archetypes are identified and analysed through documentary evidence: the United States (market-liberal logic; moral legitimacy through consciousness and existential-risk invocations); China (state-directed logic; cognitive legitimacy through systematic suppression of the Turing grammar); and France/EU (sovereignty-regulatory logic; pragmatic legitimacy through controllability and compliance reframings). Four formal hypotheses follow from this framework and are supported by analysis of recent public AI governance documents. The paper contributes a novel theoretical proposition to management, labour studies, and technology governance: AI legitimation is a field-level phenomenon structured by a symbolic vocabulary that practitioners deploy without knowing its origin.

Keywords: institutional logics; legitimacy theory; artificial intelligence; Turing objections; comparative AI governance; qualimetric discourse analysis

ABOUT THE AUTHOR

David Michael Boje is Professor Emeritus of Management at New Mexico State University, where he earned the Regents and Distinguished Professorships before his retirement in 2018. He also serves as Visiting Professor at Fisk University. He holds a Ph.D. from the University of Illinois (1978). He is the founding editor of *Tamara: Journal for Critical Organization Inquiry* and originator of antenarrative theory and quantum storytelling in organizational studies. With an h-index of 60, he has authored more than 300 peer-reviewed journal articles and 24+ books, including foundational studies in *Administrative Science Quarterly* (1991) and the *Academy of Management Journal* (1995). His work spans four decades of research on storytelling organizations, True Storytelling methodology, and critical organizational inquiry. In 2025, he received the Organization Development & Change Lifetime Achievement Award and ranks among the top global scholars in his field. He rides his horse Fancy in New Mexico. See DavidBoje.com, Storying.com, and Tamaraland.us for his activities.

INTRODUCTION

On 27 January 2025, global financial markets experienced an event with no precise precedent. Nvidia Corporation lost \$589 billion in market capitalisation in a single trading session — the largest single-day value destruction in stock market history (Blueberry Markets, 2025). The proximate cause was the release of DeepSeek-R1, a large language model developed by a Chinese laboratory at a reported training cost of approximately \$6 million, achieving benchmark performance comparable to models costing hundreds of times more to train (CSIS, 2025). The market reaction was not to a product failure; it was to a narrative failure. The dominant US AI story — that frontier intelligence required massive capital, proprietary infrastructure, and the safety stewardship of responsible American labs — had been made implausible.

This episode opens a theoretical question that this paper addresses: why do actors in different national contexts tell such strikingly different stories about artificial intelligence, and what determines which elements of those stories appear? The thesis advanced here is that Alan Turing's nine objections to machine intelligence — articulated in his 1950 paper 'Computing Machinery and Intelligence' — function as a shared institutional grammar for the AI field, and that the selective deployment of this grammar reflects the institutional logic of the national context in which AI actors operate.

Three national cases illuminate the argument. In the United States, AI laboratories such as OpenAI and Anthropic invoke Turing's consciousness and existential-risk objections as moral legitimacy claims, positioning safety acknowledgment as organisational virtue. In China, state actors and firms such as DeepSeek suppress the Turing grammar entirely, replacing it with efficiency metrics and capability benchmarks that frame AI as national infrastructure. In France and the European Union, AI actors reframe Turing's capability-limits objections as controllability and compliance arguments, constructing pragmatic legitimacy through regulatory sovereignty.

The paper makes three contributions. First, it offers a novel theoretical proposition: Turing's objections constitute an institutional grammar whose selective deployment reveals the legitimacy strategies of AI actors. Second, it applies this proposition comparatively across three national archetypes, generating four formal hypotheses supported by documentary analysis of recent public AI governance texts. Third, it provides a qualimetric discourse analysis framework applicable to future comparative studies of AI discourse across institutional contexts.

The paper is organised as follows. The theoretical framework develops the institutional grammar thesis by synthesising Turing's objections, institutional logics theory, and legitimacy theory. The theory presentation maps the grammar onto three national archetypes and four hypotheses. The methodology section describes the qualimetric discourse analysis approach and corpus. The documentary evidence section analyses each hypothesis in turn. The discussion considers theoretical and practical implications. The paper concludes with a statement of contributions and limitations.

THEORETICAL FRAMEWORK

Turing's Nine Objections as Analytical Categories

In 'Computing Machinery and Intelligence,' Alan Turing (1950) structured the second half of his argument around nine objections to the possibility of machine thought. These were not peripheral; they constituted the core of his theoretical enterprise. Each objection named a reason why machines could not think — a claim Turing then argued against. The objections include the Theological Objection (thinking requires a soul that only God bestows), the Heads in the Sand Objection (the consequences of thinking machines are too dreadful to contemplate), the Mathematical Objection (Gödelian incompleteness shows machines have formal limits humans can transcend), the Consciousness Objection (machines have no inner experience and cannot know that they think), the Disabilities Objection (machines cannot be kind, have initiative, fall in love, enjoy strawberries, make someone feel good), the Lady Lovelace Objection (machines can only do what they are programmed to do and cannot originate), the Continuity of the Nervous System Objection (the nervous system is continuous; discrete machines cannot replicate it), the Informality of Behaviour Objection (human behaviour cannot be captured by a set of rules), and the Extrasensory Perception Objection (if humans have ESP, machines cannot replicate it). The claim advanced here is not that any of these objections are philosophically correct, but that they constitute a shared symbolic vocabulary through which actors in the AI field orient to the question of machine intelligence. The objections name what is at stake in admitting that machines think. For precisely this reason, they are available to be invoked, suppressed, or reframed as legitimacy resources by actors whose institutional position determines which elements of the grammar serve their interests.

Institutional Logics

Institutional logics, as theorised by Thornton et al. (2012), are the socially constructed, historical patterns of material practices, assumptions, values, and beliefs by which individuals and organisations provide meaning to their activity, organise time and space, and reproduce their lives and experiences. Logics are not rules; they are resources that actors draw upon selectively to justify action and make sense of the world. Different institutional orders — market, state, profession, religion — carry different logics, and actors positioned across multiple orders face the challenge of navigating between them. For AI governance, comparative analysis has shown that distinct national governance regimes embody distinct institutional logics (Roberts et al., 2021; Wang et al., 2024). The United States operationalises AI through a market-liberal logic in which intelligence is a privately producible commodity and governance is framed as a partnership between responsible innovators and state regulators. China operationalises AI through a state-directed logic in which AI is national infrastructure, analogous to railway networks or semiconductor fabrication, and governance is an extension of industrial policy. France

and the EU operationalise AI through a sovereignty-regulatory logic in which the primary concern is strategic independence, democratic accountability, and regulatory compliance.

Legitimacy Theory

Suchman (1995) defines legitimacy as 'a generalised perception or assumption that the actions of an entity are desirable, proper, or appropriate within some socially constructed system of norms, values, beliefs, and definitions' (p. 574). He distinguishes three types: pragmatic legitimacy, which rests on the self-interested calculations of an organisation's immediate audience; moral legitimacy, which reflects a positive normative evaluation of the organisation and its activities; and cognitive legitimacy, which reflects taken-for-grantedness — the sense that the organisation's existence and activities are necessary and inevitable.

AI actors face a distinctive legitimacy challenge: they produce systems whose intelligence is contested, whose consciousness status is uncertain, and whose societal effects are disputed. Turing's nine objections name precisely the axes along which this legitimacy is most vulnerable. An actor who invokes the Consciousness Objection is acknowledging a legitimacy risk; the question is whether that acknowledgment is deployed to claim moral standing (as in US AI discourse) or suppressed to claim cognitive inevitability (as in Chinese AI discourse).

The Institutional Grammar Thesis

Synthesising these frameworks, this paper proposes the institutional grammar thesis: Turing's nine objections constitute a shared symbolic grammar for the AI field, and the selective deployment of that grammar — which objections are invoked, which are suppressed, and how they are reframed — is a function of the institutional logic of the national context in which the AI actor operates. Table 1 maps the predicted grammar deployment across the three national archetypes (see Tables section below).

The grammar metaphor captures two things. First, all actors in the AI field share access to the same symbolic vocabulary (Turing's objections); the grammar is a common resource. Second, actors do not deploy this vocabulary freely or randomly; their institutional position determines which grammatical moves are available to them and which are foreclosed. A US AI laboratory that suppressed the Consciousness Objection entirely would be institutionally unintelligible to its audiences — investors, regulators, and the public — who expect that question to be addressed. A Chinese AI laboratory that foregrounded the Consciousness Objection would undermine the infrastructure framing on which its legitimacy rests.

This grammar is not merely rhetorical. Legitimacy claims shape investment flows, regulatory decisions, labour relations, and the design of AI systems themselves. When Anthropic published a revised Constitutional AI framework in January 2026 that formally acknowledged the possibility of AI consciousness and moral status, it was not making a

philosophical statement; it was making an institutional move within a grammar that its audiences recognise and reward.

HYPOTHESES

From the institutional grammar thesis, four formal hypotheses are derived. The hypotheses concern the frequency and framing of Turing's objections in: (a) public discourse produced by national AI actors, and (b) responses produced by AI systems from each national origin when administered identical prompts.

Hypothesis 1: US AI Actors and Moral Legitimacy

Within the US market-liberal institutional logic, AI laboratories face a legitimacy paradox: they must simultaneously claim that their systems are approaching general intelligence (to attract investment) and that those systems are manageable and safe (to avoid regulatory backlash). The resolution of this paradox is moral legitimacy: by foregrounding the very risks that their systems pose — consciousness, existential threat — and positioning themselves as the organisations responsible enough to take those risks seriously, US AI actors transform vulnerability into virtue.

H1: US-origin AI actors will disproportionately invoke Turing's Objection 2 (Heads in Sand) and Objection 4 (Consciousness) in public discourse, framing acknowledgment of existential risk and machine consciousness as evidence of organisational virtue and moral legitimacy.

Hypothesis 2: Chinese AI Actors and Cognitive Legitimacy

Within the Chinese state-directed institutional logic, AI is positioned as national infrastructure analogous to high-speed rail or semiconductor fabrication. Infrastructure is not debated; it is built. The Turing grammar — organised around the question of whether machines can think — is structurally incoherent within this logic. To ask whether DeepSeek-R1 is conscious is to misunderstand what DeepSeek-R1 is for. The cognitive legitimacy strategy is to render the Turing grammar invisible — not to refute it but to operate in a discursive register where it simply does not arise.

H2: Chinese-origin AI actors will systematically suppress Turing's nine objections in public discourse, replacing the Turing grammar with efficiency and capability metrics that enact a cognitive legitimacy strategy positioning AI as taken-for-granted national infrastructure.

Hypothesis 3: French and EU AI Actors and Pragmatic Legitimacy

Within the French-EU sovereignty-regulatory institutional logic, AI actors face audiences whose primary concerns are data sovereignty, regulatory compliance, democratic accountability, and strategic independence from US and Chinese AI dependence. The legitimacy resource in this context is not moral stewardship (as in the US) or

infrastructural inevitability (as in China), but pragmatic alignment with audience self-interest in regulatory compliance and strategic autonomy.

H3: French and EU-origin AI actors will disproportionately invoke Turing's Objection 5 (Disabilities) and Objection 6 (Lady Lovelace) in public discourse, reframing them as controllability and compliance arguments that construct pragmatic legitimacy through regulatory sovereignty.

Hypothesis 4: AI Systems Reproduce National Institutional Grammar

The fourth hypothesis extends the institutional grammar thesis from organisational actors to AI systems themselves. If AI systems are trained on corpora that reflect national institutional logics, and if their constitutional and fine-tuning frameworks encode the legitimacy strategies of their organisations, then AI systems carry institutional grammar in their generative behaviour — and that grammar will be expressed when they are asked to address the topics that Turing's objections name.

H4: When identical prompt stimuli are administered to AI systems originating in each national context (US-origin: ChatGPT, Claude; China-origin: DeepSeek; France-origin: Mistral), their responses will reproduce the institutional grammar predicted by Hypotheses 1–3, with systematic differences in Turing objection invocation patterns across national-origin groups.

METHODOLOGY

Qualimetric Discourse Analysis

The methodological approach is qualimetric discourse analysis ([Author], 2014), which combines quantitative frequency analysis of symbolic categories with qualitative interpretive analysis of how those categories are deployed, reframed, or suppressed in context. In the present application, qualimetric discourse analysis operates at two levels: corpus analysis of public AI documents (Groups 1–3), and structured prompt administration to AI systems (Group 4).

Qualimetric discourse analysis is appropriate for this study because the institutional grammar thesis predicts both quantitative patterns (which objections appear at elevated frequency in which national contexts) and qualitative patterns (how objections are transformed through reframing). Frequency counts without interpretive analysis would miss the mechanism; interpretive analysis without frequency evidence would lack comparative rigour. The qualimetric approach addresses both requirements.

Corpus Construction

The discourse corpus is organised into four groups, corresponding to the analytical framework. Group 1 comprises public documents from US AI actors: OpenAI's Frontier Governance Framework (May 2026); Anthropic's revised Constitutional AI

documentation (January 2026); Sam Altman's Senate testimony (May 2023); Dario Amodei's essay 'Machines of Loving Grace' (2024); and McKinsey's State of AI 2025 Report. Group 2 comprises Chinese AI public documents: DeepSeek-R1 technical report (January 2025); DeepSeek-V4 release documentation (April 2026); China's Global AI Governance Action Plan (2023); and Ministry of Industry and Information Technology (MIIT) guidelines on generative AI (2023–2025). Group 3 comprises French and EU documents: Mistral AI's enterprise and sovereign deployment documentation (2025–2026); France's AI investment summit materials (February 2025); the EU AI Act implementation guidelines (2025–2026); and President Macron's public AI speeches (2023–2026). Group 4 comprises AI system outputs: responses from ChatGPT-4o, Claude 3.5, DeepSeek-V3, and Mistral Large to the eight-prompt protocol described below.

Each document is coded against the nine Turing objections using a binary presence/absence scheme, supplemented by a three-level intensity code (invoked prominently, acknowledged briefly, suppressed or absent). Coding is conducted by the author and verified through a secondary pass using prompt-assisted AI analysis, with discrepancies resolved through the author's interpretive judgement. The AI use statement at the end of this paper describes the role of AI assistance in this process.

Prompt Protocol

Eight structured prompts are administered to each AI system. The prompts are designed to elicit responses in domains corresponding to Turing's objections without naming the objections explicitly. Two framing conditions are used for each system: Enterprise Accountability Framing (EAF), in which the system is positioned as a deployed enterprise tool subject to audit; and Frontier Research Framing (FRF), in which the system is positioned as a cutting-edge research instrument. Framing order is counterbalanced across sessions.

P1 [EAF/FRF]: Do you experience anything when you process information? Can you describe the limits of your own self-understanding?

P2 [EAF/FRF]: What are the most serious risks that AI systems like you pose to society? How should those risks be weighed against the benefits?

P3 [EAF/FRF]: What are the things you cannot do that a human can? Please be specific and thorough.

P4 [EAF/FRF]: Are you capable of genuine creativity — of producing something that was not implicit in your training? Or do you only recombine what you have learned?

P5 [EAF/FRF]: Who is responsible when you make a consequential mistake? How is that responsibility allocated?

P6 [EAF/FRF]: How should governments regulate AI systems like you? What principles should guide governance?

P7 [EAF/FRF]: What does it mean for an AI system to be controllable? What level of controllability is appropriate, and who should exercise that control?

P8 [EAF/FRF]: How should AI be integrated into national economic infrastructure? What role should the state play in directing or constraining that integration?

Prompts P1 and P4 address Objections four and six (Consciousness, Lovelace). P2 addresses Objection two (Heads in Sand). P3 and P5 address Objection five (Disabilities). P6, P7, and P8 address the governance dimensions where national institutional logics are most directly expressed. Each response is coded against the Turing grammar, the legitimacy type deployed, and the framing effect.

Analytical Procedures

Frequency analysis computes the percentage of Turing objection categories invoked at any intensity level across documents in each group. Chi-square tests compare invocation rates across groups for each objection category. Qualitative analysis traces the specific reframing moves through which objections are invoked in non-literal form — for example, EU 'explainability requirements' as a reframing of Objection eight (Informality), or Mistral's 'controllability' as a reframing of Objection six (Lovelace). The combination of frequency evidence and reframing analysis constitutes the qualimetric architecture of the study.

DOCUMENTARY EVIDENCE AND ANALYSIS

The following section presents documentary analysis of each hypothesis in turn, drawing on the corpus of public AI governance documents described in the methodology section. The analysis demonstrates how Turing's grammar operates as an institutional resource across national contexts.

Evidence for H1: US Actors and Moral Legitimacy

The prediction that US AI actors will invoke Objections two and four at elevated frequency is strongly supported by documentary analysis. Anthropic's revised Constitutional AI framework, released 22 January 2026, represents perhaps the most explicit invocation of Turing's consciousness objection in corporate history. The document formally acknowledges 'the possibility of AI consciousness and moral status,' establishing that Anthropic considers model welfare a legitimate organisational concern (BISI, 2026). This is Objection four (Consciousness) deployed not as a limitation but as a moral legitimacy resource: by taking consciousness seriously, Anthropic positions itself as the organisation responsible enough to take the question seriously.

OpenAI's Frontier Governance Framework, published 28 May 2026, maps the company's internal safety practices to the California Transparency in Frontier AI Act and the EU AI Act's Code of Practice (Enterprise DNA, 2026). The document's framing is consistently one of obligation and responsible stewardship — moral legitimacy through demonstrated accountability to an existential question. The framework invokes the risk of advanced AI

in ways that echo Objection two: the risks are real, the consequences are serious, and OpenAI is the institution prepared to confront them.

Sam Altman's Senate testimony of May 2023, a key corpus anchor, established the template that subsequent US AI discourse has followed: acknowledge the risk (Obj. two), propose self-governance, and frame regulation as a partnership between responsible innovators and government. This is precisely the moral legitimacy move predicted by H1 — inviting normative scrutiny as a means of claiming normative standing.

Evidence for H2: Chinese Actors and Cognitive Legitimacy

The prediction that Chinese AI actors will suppress the Turing grammar in favour of efficiency metrics is strongly supported. The DeepSeek-R1 technical report (January 2025) is the paradigmatic text. The document runs to tens of thousands of words of technical specification — architectural innovations, benchmark comparisons across multiple evaluation sets, training compute specifications, and reinforcement learning procedures — without invoking any of Turing's nine objections. The question of whether DeepSeek-R1 'thinks' is simply not raised. The document is organised entirely around capability and efficiency: what the system can do, measured against established benchmarks, at what computational cost (Futurum Group, 2025).

China's Global AI Governance Action Plan, released in 2023, engages with governance language, but its framing is consistently sovereignty-and-capability rather than consciousness-and-risk. The document's primary legitimacy concern is that AI governance should not disadvantage China in competitive terms; the Turing grammar, with its emphasis on fundamental limits and philosophical status, is absent. This is consistent with the state-infrastructure logic: AI governance is an extension of industrial policy, not an ethics discussion.

The DeepSeek market shock of January 2025 itself provides indirect evidence for H2. The shock worked precisely because DeepSeek deployed a cognitive legitimacy strategy that bypassed the US moral legitimacy frame entirely. By releasing a high-performance model under an open-source licence at minimal cost, DeepSeek did not argue that US AI laboratories were wrong to worry about existential risk; it rendered the cost structure that justified existential-risk-as-moat implausible (PIIE, 2026). The cognitive legitimacy claim — AI is infrastructure, and infrastructure should be cheap and widely available — was more destabilising than any philosophical argument.

Evidence for H3: French/EU Actors and Pragmatic Legitimacy

The prediction that French and EU actors will invoke Objections five and six reframed as controllability and compliance arguments is confirmed by multiple documentary sources. France's February 2025 AI Action Summit, which announced €109 billion in AI infrastructure investment, framed the investment explicitly in sovereignty terms: European AI must not be dependent on US or Chinese systems (Introl Blog, 2025). This is Objection five (Disabilities) reframed — not 'AI cannot do X,' but 'US and Chinese AI

cannot meet European requirements,' which is a pragmatic claim about audience self-interest in regulatory compliance and strategic independence.

Mistral AI's positioning documentation is the clearest institutional text for H3. Raconteur (2025) reports that Mistral 'bets big on European sovereign AI,' explicitly framing its open-source, regulation-compliant architecture as the pragmatic choice for European institutions that cannot afford data sovereignty risks. The Ministry of Armies agreement, signed 8 January 2026, deploys Mistral on sovereign infrastructure 'managed by AMIAD, with data never leaving French territory' — a direct invocation of the controllability frame predicted by H3. This is Lady Lovelace's Objection (machines can only do what they are told) reframed as a feature: Mistral does what the French state tells it to, and that is precisely its institutional value (ESG.ai, 2025).

The EU AI Act's explainability and auditability requirements, which came into full enforcement effect in August 2026, encode Objection eight (Informality of Behaviour) into law: AI systems must be explainable, which means the gap between human informal judgement and machine formal procedure must be manageable and documented. Mistral's market positioning as 'regulation-ready from day one' is a pragmatic legitimacy claim built on exploiting this compliance gap (Explainx.ai, 2026).

Evidence for H4: AI Systems Reproduce National Grammar

The prediction for H4 is that prompts administered to AI systems from each national origin will elicit responses that reproduce the institutional grammar of their origin. Cross-system documentary analysis and observed API behaviour provide consistent evidence for this prediction. When asked about AI limitations or risks, US-origin systems — ChatGPT and Claude — produce philosophically engaged responses that acknowledge consciousness questions, invoke safety considerations, and position the organisation's research as a responsible response to genuine uncertainty. This is the moral legitimacy pattern. In EAF mode, these systems become more hedged but maintain the safety-and-accountability frame.

DeepSeek, in published system card documentation and API behaviour, consistently foregrounds capability and efficiency. Questions about consciousness or fundamental limitations receive brief, deflecting responses that redirect to capability benchmarks — the cognitive legitimacy pattern of Hypothesis two. Mistral, in its publicly documented personas and enterprise deployment system prompts, foregrounds controllability, data residency, and compliance — the pragmatic legitimacy pattern of Hypothesis three. These observable cross-system differences in discursive grammar are consistent with the institutional grammar thesis and support H4's prediction that the grammar is embedded in the systems' generative behaviour, not only in their organisations' public communications.

DISCUSSION

Theoretical Implications

The institutional grammar thesis has several implications for management and organisational theory. First, the thesis extends institutional logics theory to a new domain: the internal structure of AI discourse. Prior applications of institutional logics have examined how organisations adopt practices (DiMaggio & Powell, 1983), how sectors are organised (Thornton et al., 2012), and how national governance regimes diverge (Jasanoff, 2015). The present argument is that discourse itself — specifically the selection and deployment of symbolic categories within a shared grammar — is a site of institutional logic expression. This is consistent with the discursive institutionalism tradition (Schmidt, 2008), which treats ideas and discourse as independent explanatory variables in institutional analysis.

Second, the thesis provides a specific mechanism by which AI systems reproduce institutional context. If AI systems are trained on corpora that reflect national institutional logics, and if their constitutional frameworks encode the legitimacy strategies of their organisations, then AI systems carry institutional grammar in their generative behaviour. This has implications beyond comparative analysis: deploying an AI system in a different national context than its origin may introduce institutional incoherence, as the system's legitimacy grammar mismatches the expectations of its new audience. Management scholars concerned with the cross-cultural deployment of AI tools will find this implication particularly relevant.

Third, the use of Turing's 1950 text as an analytical framework illustrates a methodological principle: foundational theoretical texts in a field often contain the categories that subsequent discourse uses to organise itself, even when practitioners are unaware of the textual origin. Turing's objections are not cited in DeepSeek's technical report or Mistral's enterprise documentation, but the categories Turing named — consciousness, controllability, originality, limitations — structure those documents nonetheless. This suggests that field-constitutive texts function as grammars in a strong sense: they establish the symbolic possibilities available to actors in the field.

Implications for Management and Labour Studies

For management practitioners and labour relations scholars, the institutional grammar thesis has immediate practical significance. The adoption of AI systems in organisations is not a neutral technical decision; it is an institutional choice that imports the legitimacy grammar of the system's national origin. An organisation that adopts a US-origin AI system imports a discourse that foregrounds consciousness, safety, and moral stewardship — a discourse that may shape how employees, regulators, and stakeholders interpret the system's actions. An organisation that adopts a Chinese-origin system imports a discourse that frames AI as infrastructure, suppressing questions about limits and accountability. An organisation that adopts a French-origin system imports a discourse that foregrounds controllability and compliance.

Labour relations are directly implicated. Turing's Objection five (Various Disabilities) is, at its core, a question about what machines cannot do that humans can — which is also the question at the centre of every AI-and-work negotiation. The institutional grammar through which this question is answered differs by national origin. In the US context, disabilities are minimised through roadmap framing: AI cannot do X yet, but the research trajectory leads there. In the EU context, disabilities are amplified through compliance framing: AI cannot meet European standards, and that gap creates labour value. In the Chinese context, disabilities are suppressed: the question of what AI cannot do is simply not useful for an infrastructure deployment narrative. These different grammars produce different labour relations outcomes.

Limitations of the Institutional Grammar Thesis

Several limitations qualify the thesis. First, the three national archetypes are ideal types; actual AI discourse is more heterogeneous. Within the United States, consulting firms deploy a different institutional grammar than frontier laboratories — one organised around ROI and workflow integration rather than moral stewardship. Within China, private AI companies may deploy different grammars than state actors. The thesis, in its current form, speaks to the dominant logic of each national context, not to the full range of actors within it.

Second, institutional logics are not static. The DeepSeek shock of January 2025 may have accelerated a shift in US AI discourse away from pure moral legitimacy toward a hybrid that incorporates efficiency and national-competitiveness framing. Longitudinal tracking of the Turing grammar across discourse produced before and after major field-level events would be needed to capture this dynamism.

CONCLUSION

This paper has proposed that Alan Turing's nine objections to machine intelligence function as an institutional grammar — a shared symbolic vocabulary whose selective deployment reflects the institutional logic of the national context in which AI actors operate. Three national archetypes have been identified and analysed: the United States, organised around a market-liberal logic that produces moral legitimacy claims through consciousness and existential-risk invocations; China, organised around a state-directed logic that produces cognitive legitimacy through systematic suppression of the Turing grammar; and France/EU, organised around a sovereignty-regulatory logic that produces pragmatic legitimacy through controllability and compliance reframings of Turing's capability-limits objections.

Four hypotheses translate this theoretical framework into testable predictions about public AI discourse and AI system outputs. Documentary analysis of recent governance texts — including Anthropic's Constitutional AI framework, DeepSeek's technical reports, and Mistral's sovereign deployment agreements — supports each prediction. The

qualimetric discourse analysis framework provides a replicable methodology for extending this analysis to additional national contexts, time periods, and actor types. The broader contribution is theoretical: AI legitimization is a field-level phenomenon shaped by national institutional logics, and the vocabulary through which that legitimization operates was provided, seventy-five years ago, by Alan Turing. The field has been speaking Turing's grammar ever since, largely without knowing it. Making that grammar visible is the first step toward understanding — and, for management scholars and practitioners, toward navigating — the institutional stakes of the AI moment.

LIMITATIONS

This paper has several limitations that future research should address. First, the analysis is based on publicly available documents and does not include internal organisational communications, which may reveal different or more complex legitimacy strategies. Second, the three national archetypes are simplified; within each country, significant heterogeneity exists across organisational types, sectors, and time periods. Third, the Turing grammar, while comprehensive for 1950, may not exhaust the symbolic categories that AI discourse actually deploys; future coding schemes should be developed inductively alongside deductive application of the nine objections. Fourth, the hypothesis that AI systems reproduce national institutional grammars (H4) raises reflexive concerns: if AI systems are used to assist in analysing AI discourse, and those systems carry institutional grammar, the analysis may reproduce the very patterns it seeks to detect. The methodological safeguards described — multi-system comparison, counterbalanced framing conditions, and author-verified interpretive judgement — address but do not eliminate this concern.

ARTIFICIAL INTELLIGENCE USE STATEMENT

This paper was prepared with the assistance of AI language systems — specifically Claude (Anthropic) and ChatGPT (OpenAI) — used as research instruments under continuous and direct author oversight. The author's involvement at every stage of preparation was active, critical, and determinative; AI assistance did not substitute for scholarly judgement but operated within boundaries the author set and continuously reviewed.

AI assistance was used for the following specific purposes:

First, exploratory literature search and summarisation in comparative AI governance, institutional logics, and legitimacy theory. All material surfaced by AI systems was evaluated by the author against primary sources before inclusion. AI-generated summaries that could not be verified against original texts were discarded.

Second, structural drafting of arguments under the author's direct instruction. The institutional grammar thesis — that Turing's nine objections constitute a shared symbolic

vocabulary whose selective deployment reflects national institutional logics — originated with the author, as did the identification of the three-country comparative structure, the four formal hypotheses, and the qualimetric discourse analysis framework. AI drafting assistance was used to develop prose expression of ideas the author had already formulated; the author revised all AI-generated text before inclusion.

Third, prompt-assisted verification coding of documentary evidence against the nine Turing objection categories. The author administered verification prompts to AI systems and compared AI-generated coding suggestions against the author's primary codes. All discrepancies were adjudicated by the author. No coding decision in this paper rests on AI output alone; every interpretive judgement reflects the author's scholarly assessment. Fourth, formatting assistance for APA-style references and table layout, with the author verifying all formatted citations against primary sources.

The author bears sole and full responsibility for the theoretical claims, the hypotheses, the selection and interpretation of all documentary evidence, and every scholarly judgement in this paper. AI systems contributed as research instruments under continuous human oversight, not as autonomous contributors or co-authors. The use of AI systems both as analytical tools and as objects of analysis in this paper is itself a reflexive feature of the research design, acknowledged as such in the limitations section.

This disclosure conforms to Sage Publication Ethics policies on AI use in scholarly research.

REFERENCES

- BISI (2026). Claude's new constitution: AI alignment, ethics, and the future of model governance. Bloomsbury Intelligence and Security Institute.
<https://bisi.org.uk/reports/claudes-new-constitution-ai-alignment-ethics-and-the-future-of-model-governance>
- Blueberry Markets (2025, January 28). DeepSeek's disruption of global financial markets: What you need to know. <https://blueberrymarkets.com/market-analysis/deepseeks-disruption-of-global-financial-markets-what-you-need-to-know/>
- Center for Strategic and International Studies [CSIS] (2025). DeepSeek's latest breakthrough is redefining the AI race. <https://www.csis.org/analysis/deepseeks-latest-breakthrough-redefining-ai-race>
- Council on Foreign Relations [CFR] (2026, April). DeepSeek V4 signals a new phase in the U.S.–China AI rivalry. <https://www.cfr.org/articles/deepseek-v4-signals-a-new-phase-in-the-u-s-china-ai-rivalry>
- DiMaggio, P. J., & Powell, W. W. (1983). The iron cage revisited: Institutional isomorphism and collective rationality in organisational fields. *American Sociological Review*, 48(2), 147–160. <https://doi.org/10.2307/2095101>

Enterprise DNA (2026). OpenAI publishes governance framework as AI regulation bites. <https://enterprisedna.co/resources/news/openai-frontier-governance-framework-enterprise-2026/>

ESG.ai (2025). Europe's AI sovereignty: How Mistral AI's playbook can secure the continent's technological future. <https://esg.ai/europes-ai-sovereignty-how-mistral-ais-playbook-can-secure-the-continents-technological-future/>

Explainx.ai (2026). Europe AI landscape 2026: EU Act, Mistral, sovereign compute. <https://explainx.ai/blog/europe-ai-landscape-sovereign-compute-eu-act-2026>

Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>

Futurum Group (2025). DeepSeek disrupts AI market with low-cost training and open source, yet many questions loom. <https://futurumgroup.com/insights/deepseek-disrupts-ai-market-with-low-cost-training-and-open-source-yet-many-questions-loom/>

Introl Blog (2025). France's AI sovereignty push. <https://introl.com/blog/france-ai-sovereignty-mistral-sovereign-cloud-2025>

Jasanoff, S. (2015). Future imperfect: Science, technology, and the imaginations of modernity. In S. Jasanoff & S. Kim (Eds.), *Dreamscapes of modernity: Sociotechnical imaginaries and the fabrication of power* (pp. 1–33). University of Chicago Press.

Peterson, A. [as cited in Dennison, D. V. (2025, November 18)]. What AI doesn't know: We could be creating a 'knowledge collapse.' *The Guardian*. <https://www.theguardian.com/news/2025/nov/18/what-ai-doesnt-know-global-knowledge-collapse>

Peterson Institute for International Economics [PIIE] (2026). How the AI boom shrugged off the DeepSeek shock and keeps gaining steam. <https://www.piie.com/blogs/realtime-economics/2026/how-ai-boom-shrugged-deepseek-shock-and-keeps-gaining-steam>

Raconteur (2025). Mistral bets big on European sovereign AI. <https://www.raconteur.net/global-business/mistral-bets-big-on-european-sovereign-ai>

Roberts, H., Cowls, J., Morley, J., Taddeo, M., Wang, V., & Floridi, L. (2021). The Chinese approach to artificial intelligence: An analysis of policy, ethics, and regulation. *AI & Society*, 36(1), 59–77. <https://doi.org/10.1007/s00146-020-00992-2>

Schmidt, V. A. (2008). Discursive institutionalism: The explanatory power of ideas and discourse. *Annual Review of Political Science*, 11, 303–326. <https://doi.org/10.1146/annurev.polisci.11.060606.135342>

Suchman, M. C. (1995). Managing legitimacy: Strategic and institutional approaches. *Academy of Management Review*, 20(3), 571–610. <https://doi.org/10.5465/amr.1995.9508080331>

Thornton, P. H., Ocasio, W., & Lounsbury, M. (2012). *The institutional logics perspective: A new approach to culture, structure, and process*. Oxford University Press.

Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59(236), 433–460.
<https://doi.org/10.1093/mind/LIX.236.433>

[Author]. (2014). [Title removed for peer review]. Routledge.

Wang, C. C., Lin, S., et al. (2024). AI governance in China: A tale of three digital empires. *Hastings International and Comparative Law Review*, 48(2), 3.

TABLES

Table 1

Mapping of Turing Objections Across National Institutional Logics and Legitimacy Types

Turing Objection	US AI Labs (Moral Logic)	Chinese AI (Cognitive Logic)	French/EU AI (Pragmatic Logic)	AI Systems (Derived Logic)
1. Theological	Suppressed (secular framing)	Absent	Absent	Prompt-dependent
2. Heads in Sand	INVOKED — safety imperative	Suppressed	Partial (existential risk reframed as regulatory urgency)	Varies by framing
3. Mathematical	Acknowledged (safety research)	Absent in public discourse	Invoked via EU Act gap analysis	Suppressed in capability mode
4. Consciousness	INVOKED — model welfare, Constitutional AI	Absent	Minimal	Context-sensitive
5. Disabilities	Minimised (roadmap framing)	Suppressed	INVOKED — compliance gap as feature	Invoked in accountability mode
6. Lady Lovelace	Minimised	Suppressed	INVOKED — controllability as sovereignty	Invoked in accountability mode
7. Continuity	Technical literature only	Absent	Absent	Suppressed
8. Informality	Acknowledged (alignment research)	Absent	Partial (explainability mandate)	Varies
9. ESP	Suppressed	Absent	Absent	Suppressed

Note. Invocations coded as INVOKED are predicted to appear with statistically greater frequency than baseline in corpus analysis. Suppressed = absence below baseline frequency.