

Privacy after Publicness: A Publicness-Inference-Power Theory of AI Governance

Mr. Lakshminarasimhan Santhanam

CEO, QML, QAF Lab India

Abstract

Privacy law and AI governance are often treated as adjacent but distinct regulatory domains. This separation becomes unstable when artificial intelligence systems transform publicly accessible or voluntarily disclosed information into sensitive attributes, predictions, rankings, and consequential decisions. This article develops a Publicness-Inference-Power (PIP) theory of AI governance. The theory rejects two opposing simplifications: that public disclosure extinguishes privacy, and that conventional data-protection rules adequately govern all downstream uses of public data. It argues instead that publicness changes the locus of regulatory concern. The principal risk moves from unauthorized access to the computational transformation of information into institutional power. The PIP model identifies five linked stages: publicness, aggregation, inference, decision, and power. At each stage, information acquires new meaning, creates new asymmetries, and may generate harms that cannot be explained by collection or disclosure alone. The article uses doctrinal and comparative analysis of the European Union General Data Protection Regulation, the EU Artificial Intelligence Act, India's Digital Personal Data Protection Act, the California Consumer Privacy Act, and major international AI-governance frameworks. It shows that current regimes contain fragments of an inference-oriented approach but do not yet provide a coherent governance architecture for derived data and consequential use. The article proposes six theoretical propositions and an operational governance model based on inference registers, contextual-purpose boundaries, provenance, validation, contestability, and action controls. The contribution is a shift from data-status governance - asking whether information is public or private - to transformation-and-power governance - asking what an AI system derives, how confidently it derives it, for what purpose, and with what effects on persons and institutions.

Keywords: AI governance; privacy; public data; inference; automated decision-making; contextual integrity; data protection; algorithmic accountability; DPDP Act; EU AI Act

Research contribution at a glance

- Introduces the Publicness-Inference-Power (PIP) model as a theory of how accessible data becomes consequential institutional power.
- Distinguishes data status, inferential transformation, and decision authority as separate governance objects.
- Explains why public availability is neither equivalent to unrestricted legitimacy nor irrelevant to governance.
- States six propositions that can be examined through legal, organizational, technical, and empirical research.

- Translates the theory into implementable controls for AI providers, deployers, auditors, and regulators.

1. Introduction

Privacy regulation was built around a powerful image: information begins within a protected sphere and becomes risky when another actor collects, discloses, or processes it without adequate authority. Digital platforms have weakened that image. Individuals now publish photographs, employment histories, opinions, locations, relationships, purchases, health experiences, and routine behaviour across interconnected services. Publicness is no longer exceptional. It is part of social participation, professional visibility, consumer convenience, and political expression.

This behavioural change has generated a misleading binary. One position treats voluntary disclosure as the effective surrender of privacy: once information is public, the individual has assumed the risk of any subsequent use. The opposing position treats the public or private status of data as secondary because data-protection and AI-governance principles continue to apply to processing. Neither position adequately explains the transformation introduced by contemporary AI. Artificial intelligence does not merely reproduce what was disclosed. It aggregates fragments, identifies correlations, generates representations, predicts traits, classifies persons, recommends action, and increasingly acts through connected tools and institutional workflows.

The governance problem therefore does not end at disclosure. Public facts can be transformed into claims that the person never stated, did not know could be inferred, cannot easily inspect, and may be unable to contest. A professional profile may become an employability score. Public posts may be converted into an estimate of political orientation or psychological vulnerability. A photograph may become a biometric template. Conversational traces may yield demographic or health-related predictions even after obvious identifiers are removed. The significance of the information lies increasingly in what a system can make of it and what an institution can do with the result.

This article develops a Publicness-Inference-Power (PIP) theory of AI governance. Its central claim is that publicness changes, but does not terminate, the normative analysis. As information moves through aggregation, inference, and decision, the relevant governance question shifts from whether an actor may access a datum to whether the transformation and use of that datum are legitimate, reliable, proportionate, contestable, and institutionally accountable. Publicness may weaken claims based purely on secrecy. It does not by itself justify every inference, every decision, or every exercise of power.

The argument contributes to AI-governance theory in three ways. First, it identifies inferential transformation as a distinct governance object between data collection and automated decision-making. Second, it proposes a causal model linking public information to institutional effects. Third, it converts the model into a governance architecture that can be implemented through organizational and technical controls. The framework is designed to complement, not replace, privacy law, anti-discrimination law, consumer protection, sector regulation, and emerging AI regulation.

The article proceeds as follows. Section 2 defines the research problem and method. Section 3 reviews the conceptual foundations in privacy and AI-governance scholarship. Section 4 presents the PIP theory, its constructs, mechanisms, boundary conditions, and propositions. Section 5 compares the treatment of public and inferred data across major legal and governance frameworks. Section 6 applies the theory to illustrative cases. Section 7 develops an operational governance architecture. Sections 8 and 9 address

counterarguments, research design, and testability. The final sections identify policy implications, limitations, and the broader theoretical significance of governing information as power.

2. Research Problem, Questions, and Method

2.1 The regulatory gap

Existing governance approaches frequently begin from one of three objects: personal data, the AI system, or the automated decision. Data-protection law asks whether personal data are processed lawfully and fairly. AI regulation classifies systems or uses according to risk. Organizational governance frameworks allocate roles, controls, and assurance activities across the system lifecycle. These perspectives are necessary, but the transformation from accessible data to inferred knowledge often falls between them.

This gap is especially visible when the source material is publicly available. Some legal regimes exclude categories of public information from their scope. Others continue to regulate processing but may provide limited protection against the accuracy, legitimacy, or downstream use of inferred claims. AI frameworks may require risk management and transparency without specifying when a particular inference should be prohibited, validated, or subjected to heightened review. The result is a recurring accountability discontinuity: the input may be lawful, the model may be generally permitted, and the final decision may be formally reviewed, while the inferential bridge connecting them remains opaque.

2.2 Research questions

- RQ1. How does the public availability of personal information alter, rather than eliminate, the object of privacy and AI governance?
- RQ2. What mechanisms convert public or voluntarily disclosed data into consequential institutional power?
- RQ3. Where do major privacy and AI-governance regimes regulate - or fail to regulate - these mechanisms?
- RQ4. What governance controls are required to make high-impact AI inferences legitimate, reliable, traceable, and contestable?

2.3 Method and scope

The article uses a theory-building method based on conceptual synthesis and comparative doctrinal analysis. It integrates privacy theory, data-protection scholarship, algorithmic-governance research, and institutional AI-risk frameworks. The analysis is not a systematic literature review and does not claim exhaustive coverage. Instead, it identifies recurring constructs and explanatory gaps across established theories and current regulatory instruments.

The comparative legal component examines the GDPR, the EU AI Act, India's Digital Personal Data Protection Act, and the CCPA because they embody materially different approaches to public data, lawful processing, rights, and system risk. The analysis also uses the NIST AI Risk Management Framework, OECD AI Principles, and UNESCO Recommendation on the Ethics of Artificial Intelligence as examples of non-statutory and international governance approaches. The article focuses on civilian and commercial AI uses. National-security intelligence collection and purely interpersonal uses require separate treatment.

3. Conceptual Foundations

3.1 From secrecy to contextual information flow

The modern privacy tradition has never been reducible to secrecy. Warren and Brandeis framed privacy as a response to technologies and institutions capable of intruding upon personal life. Westin emphasized the individual's claim to determine when, how, and to what extent information about the self is communicated. Solove later demonstrated that privacy problems form a family of harms, including surveillance, aggregation, secondary use, exclusion, increased accessibility, distortion, and decisional interference. These accounts make clear that information can create privacy harm even when the central wrong is not the initial revelation of a secret.

Nissenbaum's theory of contextual integrity is particularly important for the present argument. It defines privacy through appropriate information flows rather than through a simple private-public boundary. A disclosure that is appropriate among friends, patients and physicians, or professionals and recruiters does not automatically become appropriate for every recipient, purpose, or transmission condition. Contextual integrity therefore supplies a normative explanation for why information can remain privacy-relevant after it is visible or accessible.

Publicness nevertheless matters. A person who intentionally publishes a statement ordinarily has a weaker claim that nobody may read or repeat it than a person whose confidential file was breached. The mistake is to infer from this weaker secrecy claim that all downstream transformations are legitimate. Scale, aggregation, persistence, cross-context transfer, and computational inference can change the character of an information flow. AI intensifies each of these transformations.

3.2 Consent, aggregation, and the limits of procedural privacy

Notice-and-consent frameworks presume that individuals can evaluate proposed uses and make meaningful choices. Big-data practices undermine that presumption because future uses, correlations, and inferences are difficult to anticipate at the time of disclosure. Barocas and Nissenbaum describe how data-intensive systems can bypass procedural safeguards based on anonymity and consent. A user may consent to publishing a profile for professional networking without understanding that the same profile may later be scraped, combined with unrelated data, and used to train or operate a classification system. The difficulty is not merely that privacy policies are long. It is that inference is generative. The system can produce a claim that was absent from any single input record. Consent to each input does not logically amount to consent to every possible conclusion. Nor does de-identification necessarily resolve the issue, because correlated datasets and model outputs can support re-identification or attribute inference.

3.3 Inference as a governance object

Wachter and Mittelstadt identify an accountability gap for high-risk inferences and propose a right to reasonable inferences. Their account shifts attention from the lawfulness of input data to the normative acceptability, relevance, and reliability of conclusions used in important decisions. This article builds on that insight but places inference within a wider institutional chain. The inference is not the endpoint. Its governance significance depends on who receives it, how it is validated, what decision it informs, and what power the receiving institution exercises.

Inference must also be distinguished from explanation. An organization may explain that a model uses public professional information to rank applicants without demonstrating that the inferred trait is valid, relevant, non-discriminatory, or appropriate for the employment context. Transparency can reveal a

problematic practice without legitimizing it. Conversely, a technically accurate inference may still be normatively inappropriate if it violates contextual expectations or enables disproportionate control.

3.4 AI governance: principles, systems, and organizational practice

AI governance scholarship increasingly defines governance as a system of rules, practices, processes, roles, and technical tools that align AI use with legal requirements, organizational objectives, and ethical values. Reviews of the field distinguish questions of who governs, what is governed, when governance occurs in the lifecycle, and how it is operationalized. More recent work emphasizes structural, relational, and procedural practices. These developments are valuable because they move beyond abstract ethics toward implementation.

Yet principle-based governance can remain underspecified. Fairness, transparency, privacy, safety, and accountability do not automatically identify the unit of analysis or the control point. The PIP theory addresses this by locating governance interventions along a transformation chain. It asks not only whether an AI system is trustworthy in general, but whether a particular class of inference, in a particular context, can legitimately support a particular institutional action.

4. The Publicness-Inference-Power Theory

4.1 Core claim

Public availability reduces the force of secrecy-based claims, but it does not determine the legitimacy of aggregation, inference, decision, or institutional action. AI governance must therefore follow information through its transformations into power.

The PIP theory is a meso-level governance theory. It does not attempt to define privacy in every context or provide a complete theory of AI regulation. It explains a recurring class of governance failures in which accessible information is transformed into claims and actions that materially affect persons. The theory identifies five stages and six cross-cutting governance requirements.

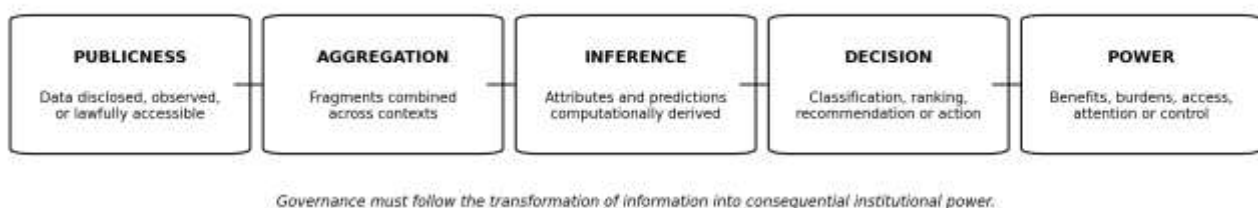


Figure 1. The Publicness-Inference-Power transformation chain.

4.2 Constructs

Publicness. The practical accessibility of information to an audience beyond the original subject or bounded context. Publicness is graduated rather than binary. Searchability, persistence, audience size, platform rules, technical barriers, and legal obligations all affect the degree of publicness.

Aggregation. The combination of information across records, platforms, time periods, persons, or social contexts. Aggregation changes informational meaning because a collection can reveal patterns not visible in isolated items.

Inference. A computationally generated claim, estimate, classification, prediction, representation, or recommendation derived from observed data. Inferences vary in sensitivity, verifiability, confidence, reversibility, and temporal stability.

Decision. A human or automated selection, ranking, allocation, denial, prioritization, recommendation, or intervention informed by an inference.

Power. The institutional capacity to alter another person's opportunities, burdens, visibility, treatment, access, price, reputation, or range of practical choices.

4.3 Mechanisms

- **Scale transformation:** AI permits observation and analysis across volumes that humans could not practically examine. Public observation becomes systematic surveillance when it is automated, persistent, and comprehensive.
- **Context collapse:** Information disclosed for one audience or purpose is transferred into another institutional context whose norms, stakes, and power relationships differ.
- **Semantic expansion:** Models derive attributes, relationships, and predictions not explicitly present in the inputs.
- **Epistemic asymmetry:** Institutions can generate and act on profiles that subjects cannot see, understand, verify, or correct.
- **Action coupling:** Inferences are integrated into workflows, agents, scoring systems, and decision engines, reducing the distance between prediction and consequence.
- **Feedback reinforcement:** Decisions based on inferences alter future data and opportunities, potentially making the original prediction appear self-confirming.

4.4 Boundary conditions

The theory does not treat all uses of public data as equally risky. Risk increases when the inference is sensitive, weakly verifiable, unexpected in context, difficult to contest, used at population scale, or connected to a high-power decision-maker. Risk decreases when the inference is transparent, necessary for a legitimate purpose, empirically validated, proportionate, subject to human review, and separated from consequential action unless a defined threshold is met.

Public-interest uses also require differentiated treatment. Journalism, academic research, public-health analysis, fraud detection, and democratic accountability may justify uses that would be inappropriate in advertising or employment screening. The PIP model does not resolve these conflicts mechanically. It identifies the questions and evidence required for legitimate resolution.

4.5 Theoretical propositions

Proposition	Core relationship	Indicative evidence
P1 - Publicness displacement proposition	As information becomes more publicly accessible, secrecy-based protection weakens, but governance demand shifts toward the legitimacy of aggregation, inference, and use rather than disappearing.	Comparative rights claims across public and non-public datasets; regulator decisions; organizational policies.
P2 - Transformation amplification proposition	The governance risk of a dataset is not determined solely by the sensitivity of individual inputs; it increases when aggregation and	Attribute-inference accuracy; re-identification studies; difference between input and derived sensitivity.

Proposition	Core relationship	Indicative evidence
	modelling materially expand what can be known or predicted.	
P3 - Context-power interaction proposition	The same inference creates greater normative risk when transferred into a context with stronger institutional power, higher stakes, or weaker subject agency.	User expectations and harm across advertising, employment, credit, health, policing, and research contexts.
P4 - Epistemic accountability proposition	The legitimacy of consequential AI use depends on the validity, relevance, confidence, and contestability of the inference, not merely on the lawfulness of the source data.	Validation records, error distributions, appeal outcomes, decision reversals, and documentation quality.
P5 - Coupling proposition	Risk rises as the temporal and organizational distance between inference and action decreases, particularly in autonomous or agentic systems.	Incident rates and harm severity across advisory, semi-automated, automated, and agentic deployments.
P6 - Governance continuity proposition	Effective AI governance requires controls across the full publicness-to-power chain; controls limited to data intake, model development, or final decision review will leave predictable accountability gaps.	Control coverage studies; audit findings; gap analysis across lifecycle stages.

5. Comparative Regulatory Analysis

5.1 GDPR: processing remains central after disclosure

The GDPR does not generally convert publicly available personal data into unregulated material. Processing still requires a lawful basis and remains subject to principles including fairness, transparency, purpose limitation, data minimization, accuracy, storage limitation, security, and accountability. Special-category data may receive heightened protection even when elements have been made public, although the precise legal route and exceptions require context-specific analysis.

The GDPR nevertheless illustrates the inferential gap. Rights concerning access, rectification, objection, erasure, and automated decision-making provide important tools, but they do not establish a general right to the accuracy or normative reasonableness of every prediction. A model may rely on accurate inputs while producing a speculative, unstable, or contextually inappropriate inference. The PIP approach would require explicit governance of that derived claim and its connection to action.

5.2 India's DPDP Act: a sharp public-data boundary

India's Digital Personal Data Protection Act adopts an especially consequential scope rule. Section 3(c)(ii) states that the Act does not apply to personal data made publicly available by the data principal or by a person legally obliged to make it public. The statutory illustration states that when an individual publishes personal data while blogging views on social media, the Act does not apply to that data. This creates a strong doctrinal example of the article's central problem: voluntary publicness can remove the principal data-protection statute even though AI processing may generate sensitive inferences and consequential uses.

The exclusion does not make all conduct lawful. Other constitutional, sectoral, consumer-protection, intellectual-property, anti-discrimination, tort, cybercrime, platform, and contractual rules may apply. Nor does it necessarily resolve whether a derived output is the same personal data that was made public. Nonetheless, the statutory boundary creates a governance gap unless AI-specific or sector-specific rules regulate inference, profiling, and consequential use. The PIP theory provides a way to identify that gap without denying the legislature's deliberate distinction between public and non-public data.

5.3 CCPA: public information and consumer control

The CCPA and related California privacy rules provide consumer rights concerning access, deletion, correction, and certain forms of sharing or sale, while excluding defined categories of publicly available information from personal information. The definition and regulatory treatment are more qualified than a simple claim that anything visible online is public. Even so, public-information exclusions again demonstrate that data status can determine statutory coverage before the analysis reaches inference quality or decision legitimacy.

California's experience with data brokers also illustrates the limits of rights that depend on individual requests and decentralized compliance. Where hundreds of entities aggregate and trade data, the subject may not know which organization holds an inferred profile or how it affects downstream treatment. A governance system focused only on access and deletion can leave the creation and operational use of profiles largely intact.

5.4 EU AI Act: risk-based system and use governance

The EU AI Act moves beyond data protection by regulating AI systems according to categories of risk and use. It prohibits specified practices, imposes requirements on high-risk systems, creates transparency duties for certain AI interactions and generated content, and establishes obligations for general-purpose AI. This architecture recognizes that lawfulness cannot be determined solely by the status of input data. However, system classification does not eliminate the need for inference-level analysis. Within a permitted or high-risk system, some inferences may be more sensitive, less reliable, or less contextually appropriate than others. The PIP model can operate as a control layer within the Act's risk-management, data-governance, technical-documentation, transparency, human-oversight, accuracy, and post-market-monitoring requirements. It specifies what must be documented about the path from source information to derived claim and action.

5.5 NIST, OECD, and UNESCO: broad principles, limited inference granularity

The NIST AI Risk Management Framework organizes practice around Govern, Map, Measure, and Manage and treats trustworthiness as multidimensional. OECD principles emphasize human rights, transparency, robustness, and accountability. UNESCO calls for privacy protection across the AI lifecycle and links AI ethics to human dignity and broader social values. These frameworks support a lifecycle and sociotechnical approach compatible with the PIP theory.

Their limitation is not conceptual weakness but granularity. Organizations still need to decide which inferences require registration, validation, explanation, restriction, or prohibition. The PIP model supplies an intermediate structure between high-level principles and system-specific controls.

5.6 Comparative synthesis

Framework	Public-data treatment	Primary object	Inference coverage	PIP implication
GDPR	Public availability does not generally remove processing from the regime.	Personal-data processing and rights.	Partial; stronger for profiling and automated decisions than for all derived claims.	Add inference-specific necessity, validity, context, and contestability tests.
India DPDP Act	Excludes data made public by the data principal or by legal obligation.	Digital personal-data processing within statutory scope.	Potentially limited where source data fall outside scope.	Use AI/sector governance to prevent a public-data-to-inference accountability gap.
CCPA/CPRA	Excludes defined publicly available information; grants consumer rights for covered information.	Consumer personal information and business practices.	Indirect and fragmented.	Track derived profiles and downstream recipients, not only source records.
EU AI Act	Public status is not the decisive classifier.	AI systems and use cases according to risk.	Material but system-level; varies by category and obligation.	Embed inference controls within risk management and technical documentation.
NIST AI RMF	No categorical public-data exemption.	Sociotechnical AI risk across lifecycle.	Broadly compatible but non-prescriptive.	Operationalize Map/Measure/Manage at inference and action layers.
OECD/UNESCO	Human-rights and lifecycle framing.	Values, institutions, and responsible AI practice.	Normatively supportive but high level.	Convert principles into auditable inference and decision controls.

6. Illustrative Applications

6.1 Employment screening from public professional data

Consider a recruitment system that collects public resumes, professional-network profiles, publications, and social-media posts. The initial data may be intentionally public and relevant to some employment purposes. The system then infers leadership potential, cultural fit, emotional stability, political risk, or likelihood of retention. Those claims are not equivalent to the source facts. They may reflect proxies, stereotypes, or spurious correlations and may be difficult to verify.

Under the PIP model, governance would not begin and end with whether scraping was lawful. The organization would identify each consequential inference, state its purpose, document the evidence connecting it to job requirements, test error and subgroup effects, define an expiration period, disclose its use at an appropriate level, and provide a route for challenge. Inferences unrelated to legitimate job requirements or too weakly verifiable for employment decisions would be prohibited even if the inputs were public.

6.2 Facial images and biometric transformation

A publicly posted photograph is visible to human viewers. Converting it into a persistent biometric template changes the nature of the information. The template supports identity matching across contexts, locations, and time. Additional models may estimate age, emotion, health-related features, or demographic attributes. The transformation increases searchability, scalability, and institutional reach.

The PIP theory treats the photograph, biometric representation, inferred attribute, matching decision, and institutional action as separate objects. Permission or public visibility at the first stage cannot automatically authorize every later stage. Governance must consider whether the transformation is expected, necessary, accurate, proportionate, and compatible with the receiving context.

6.3 Conversational AI and latent disclosure

Conversational systems create a particularly important inference surface. Users may not state a demographic identity or health condition directly, yet repeated language patterns, topics, schedules, relationships, and preferences can support confident estimates. Removing names and obvious identifiers may therefore leave substantial inferential exposure. As assistants gain memory and tool access, inferred attributes can shape recommendations and autonomous actions.

The PIP framework requires a separation between personalization and consequential profiling. A system may use short-term context to answer a request without storing a durable claim that the user has a particular health, political, financial, or psychological characteristic. Persistent sensitive inferences should require defined authorization, provenance, confidence thresholds, use restrictions, and deletion or review mechanisms.

6.4 Public records, data brokers, and cumulative profiles

Public records serve legitimate transparency and administrative purposes. Data brokers can combine them with commercial, behavioural, device, and inferred data to create profiles for marketing, fraud screening, identity verification, pricing, or eligibility decisions. The resulting profile may be difficult for the subject to locate because no single source contains it.

The governance challenge is cumulative. Each input may be accessible and each transaction may be formally lawful, while the combined profile generates an information asymmetry capable of altering treatment. The PIP model therefore assigns responsibility to the actor that performs or commissions the aggregation and inference, as well as to the institution that operationalizes the result.

7. From Theory to an Operational Governance Architecture

The PIP theory becomes useful only if it changes organizational practice. The following architecture is designed for providers and deployers of AI systems that derive person-level or group-level claims from accessible information. Controls should be proportionate to sensitivity, uncertainty, scale, autonomy, and institutional power.

1. Inference inventory and register

- List material inferences separately from raw input fields and model outputs.
- Record whether each inference is observed, calculated, predicted, generated, or externally supplied.
- Classify sensitivity, confidence, verifiability, temporal stability, and decision relevance.
- Identify affected persons and downstream recipients.

2. Context and purpose boundary

- Document the context in which source data were disclosed or collected.
- State the new purpose and explain why transfer across contexts is appropriate.
- Prohibit incompatible secondary uses even where source access is lawful.
- Require heightened approval for employment, credit, insurance, health, education, policing, migration, and essential-service decisions.

3. Provenance and transformation trace

- Maintain lineage from source categories through features, models, prompts, tools, and transformations to the final claim.
- Record model and dataset versions, material preprocessing, and confidence measures.
- Distinguish direct evidence from proxy variables and model-generated assumptions.
- Preserve audit evidence without retaining unnecessary personal data.

4. Inference validation and reasonableness review

- Test construct validity: does the inference measure what the organization claims it measures?
- Test statistical reliability, calibration, error distribution, robustness, and subgroup performance.
- Assess normative relevance to the stated purpose.
- Ban unsupported personality, emotion, intent, or sensitive-trait claims in consequential settings unless exceptional evidence and legal authority exist.

5. Decision and action controls

- Define which inferences may inform, recommend, or automatically trigger action.
- Use higher thresholds as autonomy and consequence increase.
- Provide meaningful human review that can override the system and is not merely ceremonial.
- Log the relationship between inference, human judgment, and final outcome.

6. Subject visibility, contestability, and remedy

- Provide notice that consequential inference is occurring, at a level that does not enable manipulation or compromise legitimate security.
- Allow affected persons to access a meaningful description of high-impact inferred claims.
- Provide correction, contextualization, objection, and appeal mechanisms.
- Monitor whether challenges reveal systematic model or governance failures.

7. Post-deployment monitoring and sunset

- Monitor drift, changed context, emerging harms, feedback effects, and downstream repurposing.
- Revalidate inferences when populations, data sources, models, or decision contexts change.

- Apply expiration periods to unstable or time-sensitive inferences.
- Retire inferences and models that cannot meet evidence or accountability requirements.

7.1 A proportional inference-risk test

Organizations can prioritize controls through a structured test. The test is not a mathematical substitute for judgment; it is a disciplined way to identify when public-data use becomes high-risk inference governance.

- Sensitivity: Does the inference concern health, biometrics, sexuality, religion, political belief, finances, vulnerability, criminality, or similarly consequential traits?
- Uncertainty: Is the inference probabilistic, weakly verifiable, unstable, or dependent on proxies?
- Unexpectedness: Would a reasonable person expect the source information to be used for this inference and purpose?
- Scale: How many persons, records, contexts, and decisions are affected?
- Power: Can the recipient deny opportunities, impose burdens, alter prices, direct attention, or constrain choices?
- Autonomy: How directly can the system act without independent human judgment?
- Contestability: Can the affected person discover, understand, challenge, and remedy the inference or action?

A high score on any combination of sensitivity, uncertainty, institutional power, autonomy, or weak contestability should trigger an enhanced inference impact assessment. The assessment should be distinct from, but integrated with, privacy impact assessments, data-protection impact assessments, fundamental-rights impact assessments, model risk management, and security review.

7.2 Governance roles

Role	Primary responsibility	Required evidence
Board or executive risk owner	Set risk appetite and prohibited-use boundaries.	Approved policy, escalation thresholds, oversight minutes.
Business owner/deployer	Justify purpose and decision relevance.	Use-case record, impact assessment, operational controls.
Model/data team	Document provenance, performance, uncertainty, and limitations.	Model cards, data sheets, validation and monitoring reports.
Privacy/legal/ethics function	Assess legality, contextual appropriateness, rights, and vulnerable groups.	Legal analysis, rights assessment, consultation record.
Independent assurance/audit	Test whether controls operate and evidence supports claims.	Audit trail, findings, remediation verification.
Human decision-maker	Exercise accountable judgment and document	Decision rationale, review logs, training and competency

Role	Primary responsibility	Required evidence
	overrides.	evidence.
Affected person or representative	Challenge inaccurate or inappropriate claims and outcomes.	Accessible notice, appeal record, remedy and feedback data.

8. Counterarguments and Responses

Public data must remain usable for innovation and expression. The theory does not impose ownership over facts or require consent for every use. It uses proportionality and context. Low-impact analysis, journalism, research, and ordinary communication may require minimal controls, while sensitive inference tied to institutional power requires stronger justification.

The framework duplicates existing privacy law. Privacy law supplies essential principles and rights, but public-data exclusions, jurisdictional variation, and limited treatment of inferred accuracy leave gaps. The PIP model also applies where data-protection law is not the only or primary regime and connects privacy to system and decision governance.

Inferences are unavoidable; every human decision involves inference. The relevant difference is not that machines infer. It is the scale, persistence, opacity, replicability, autonomy, and institutional integration of computational inference. Governance should be proportionate rather than absolute.

Inference-level documentation is technically impossible for complex models. Perfect causal explanation may be unavailable, but organizations can still document source categories, feature or prompt pathways, intended constructs, evaluation results, uncertainty, output use, and decision controls. Inability to explain every internal parameter does not justify inability to govern the inference.

Subjects may manipulate systems if inferences are disclosed. Contestability does not require release of source code, anti-fraud logic, or exploitable thresholds. Layered transparency can describe the nature, purpose, principal factors, and challenge process while protecting legitimate security and intellectual-property interests.

The theory is too broad to be falsifiable. The propositions specify directional relationships among publicness, transformation, context, power, coupling, and governance coverage. They can be tested through experiments, audits, comparative case studies, incident datasets, user research, and organizational fieldwork.

9. Research Agenda and Empirical Testability

A theory paper should generate research rather than close debate. The PIP model supports at least five empirical programmes.

- **Publicness and expectation experiments:** Vary audience size, searchability, age of disclosure, platform context, and downstream purpose to measure when users perceive inference and use as legitimate.
- **Inference transformation audits:** Compare the sensitivity and accuracy of explicit source data with model-derived attributes, including how performance differs across groups and contexts.

- **Institutional power studies:** Examine whether identical inferences produce different perceived and actual harms when used by advertisers, employers, insurers, schools, healthcare providers, police, or social platforms.
- **Governance-control evaluations:** Test whether inference registers, purpose reviews, provenance, and appeal mechanisms reduce incidents, improve documentation, or change deployment decisions.
- **Cross-jurisdictional legal research:** Study how public-data exclusions, profiling rules, automated-decision rights, discrimination law, and AI regulation interact in real disputes.
- **Agentic coupling research:** Measure how risk changes when an AI system moves from prediction to tool use, transaction execution, communication, or physical action.

9.1 Candidate hypotheses

- H1. Individuals will judge an inference as less legitimate when the downstream context differs materially from the context of disclosure, even when the source data are fully public.
- H2. Perceived harm and demand for explanation will rise more strongly with institutional power than with the technical complexity of the model.
- H3. Organizations using inference inventories will identify more prohibited or unsupported use cases before deployment than organizations relying only on general AI-principles checklists.
- H4. Automated action based on an uncertain inference will produce higher error cost and lower perceived legitimacy than advisory use of the same inference.
- H5. Appeal data will reveal systematic problems not detected by aggregate model-performance metrics, particularly for context-sensitive and weakly verifiable traits.

10. Policy Implications

First, lawmakers should avoid treating public availability as a complete proxy for low risk. Public-data exclusions may be justified for ordinary observation, expression, research, and administrative transparency, but should not automatically extend to sensitive profiling or high-impact automated use.

Second, AI legislation and sector rules should define derived or inferred information explicitly. Regulation should distinguish observed facts from model-generated claims and require higher justification where inferences are sensitive, consequential, or difficult to verify.

Third, impact assessments should trace the full transformation chain. A form that lists datasets and model names without identifying material inferences and actions will miss the central governance issue.

Fourth, rights should attach to consequential inferred profiles even when source data are public. Depending on context, these may include notice, access to a meaningful description, correction of underlying facts, challenge to unreasonable inferences, human review, and remedy.

Fifth, regulators should combine privacy, consumer protection, anti-discrimination, competition, and sector expertise. Public-data inference cuts across legal silos because the harm may arise through unfair practice, exclusion, manipulation, reputational injury, or unequal treatment rather than confidentiality alone.

Sixth, technical standards should support interoperable provenance and inference documentation. An AI bill of materials should not stop at models and datasets; it should identify classes of derived claims and permitted actions.

11. Limitations

The PIP theory is conceptual and has not yet been validated through original empirical research. Its pro-

positions require testing across cultures, sectors, model architectures, and institutional settings. The constructs may overlap in practice: aggregation and inference can occur inside a foundation model, and a recommendation can itself exercise power through attention or interface design.

The comparative legal analysis is necessarily selective and does not provide jurisdiction-specific legal advice. Implementation dates, rules, regulatory guidance, and case law continue to evolve. The treatment of publicly available data also depends on definitions, source restrictions, intellectual-property rules, platform terms, and sector-specific duties not fully examined here.

Finally, the framework emphasizes risks to persons and groups arising from informational transformation. AI governance must also address security, environmental cost, labour, market concentration, systemic risk, misinformation, and catastrophic misuse. The PIP model should therefore be treated as one component of a broader governance architecture.

12. Conclusion

The digital age has not simply made private information public. It has created institutions capable of converting accessible fragments into predictions and converting predictions into action. That transformation changes the problem of privacy. Secrecy remains important, but it is no longer a sufficient organizing principle for AI governance.

The Publicness-Inference-Power theory explains why the public status of data cannot determine the legitimacy of AI use. Publicness is the beginning of an analysis, not its conclusion. Governance must follow the information through aggregation, inference, decision, and institutional power. At each stage, new claims are created, contexts are crossed, asymmetries increase, and the possibility of consequential harm changes.

The practical implication is direct. Organizations should inventory and validate material inferences, document purpose and provenance, control their coupling to decisions, provide meaningful contestability, and monitor effects after deployment. Regulators should protect persons not only from unauthorized collection but also from unreasonable computational claims and unaccountable action.

Twentieth-century privacy law was often framed around who may know. AI governance must add a further question: who may derive, decide, and act on what can be known? In a world of voluntary disclosure and machine inference, privacy survives not as a promise that information will remain unseen, but as a demand that informational power remain legitimate, bounded, and accountable.

References

1. Barocas, S., & Nissenbaum, H. (2014). Big data's end run around anonymity and consent. In J. Lane, V. Stodden, S. Bender, & H. Nissenbaum (Eds.), *Privacy, big data, and the public good* (pp. 44-75). Cambridge University Press.
2. Batool, A., Zowghi, D., & Bano, M. (2025). AI governance: A systematic literature review. *AI and Ethics*, 5, 3265-3279. <https://doi.org/10.1007/s43681-024-00653-w>
3. Cohen, J. E. (2008). Privacy, visibility, transparency, and exposure. *University of Chicago Law Review*, 75, 181-201.
4. Cohen, J. E. (2019). *Between truth and power: The legal constructions of informational capitalism*. Oxford University Press.
5. European Parliament and Council. (2016). Regulation (EU) 2016/679 (General Data Protection Regulation). Official Journal of the European Union.

6. European Parliament and Council. (2024). Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union.
7. Government of India. (2023). Digital Personal Data Protection Act, 2023 (Act No. 22 of 2023). Gazette of India.
8. Government of India, Ministry of Electronics and Information Technology. (2025). Digital Personal Data Protection Rules, 2025.
9. Larrucea, X., Santamaria, I., & Ruano, O. (2021). Towards a privacy debt. *IET Software*, 15(6), 391-403. <https://doi.org/10.1049/sfw2.12044>
10. Mäntymäki, M., Minkkinen, M., Birkstedt, T., & Viljanen, M. (2022). Defining organizational AI governance. *AI and Ethics*, 2(4), 603-609. <https://doi.org/10.1007/s43681-022-00143-x>
11. Mittelstadt, B. (2019). Principles alone cannot guarantee ethical AI. *Nature Machine Intelligence*, 1, 501-507. <https://doi.org/10.1038/s42256-019-0114-4>
12. National Institute of Standards and Technology. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0) (NIST AI 100-1). <https://doi.org/10.6028/NIST.AI.100-1>
13. National Institute of Standards and Technology. (2024). Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1).
14. Nissenbaum, H. (2004). Privacy as contextual integrity. *Washington Law Review*, 79(1), 119-158.
15. Nissenbaum, H. (2010). *Privacy in context: Technology, policy, and the integrity of social life*. Stanford University Press.
16. Nissenbaum, H. (2019). Contextual integrity up and down the data food chain. *Theoretical Inquiries in Law*, 20(1), 221-256.
17. OECD. (2024). Recommendation of the Council on Artificial Intelligence (updated). Organisation for Economic Co-operation and Development.
18. Papagiannidis, E., Mikalef, P., & Conboy, K. (2025). Responsible artificial intelligence governance: A review and research framework. *Journal of Strategic Information Systems*, 34(2), 101885. <https://doi.org/10.1016/j.jsis.2024.101885>
19. Solove, D. J. (2006). A taxonomy of privacy. *University of Pennsylvania Law Review*, 154(3), 477-560.
20. UNESCO. (2021). Recommendation on the ethics of artificial intelligence. United Nations Educational, Scientific and Cultural Organization.
21. Wachter, S., & Mittelstadt, B. (2019). A right to reasonable inferences: Re-thinking data protection law in the age of Big Data and AI. *Columbia Business Law Review*, 2019(2), 494-620. <https://doi.org/10.7916/cblr.v2019i2.3424>
22. Warren, S. D., & Brandeis, L. D. (1890). The right to privacy. *Harvard Law Review*, 4(5), 193-220.
23. Yeung, K. (2018). Algorithmic regulation: A critical interrogation. *Regulation & Governance*, 12(4), 505-523. <https://doi.org/10.1111/rego.12158>
24. Zuboff, S. (2019). *The age of surveillance capitalism*. PublicAffairs.