



GOVERNING ARTIFICIAL INTELLIGENCE IN CATASTROPHE RISK INTELLIGENCE: THE TRUST FRAMEWORK FOR TRUSTWORTHY AND EXAMINATION-READY AI GOVERNANCE IN U.S. PROPERTY AND CASUALTY INSURANCE

Jahtel Philips¹

Affiliations:

¹ Independent Researcher, AI Governance and Enterprise Risk Management, Property and Casualty Insurance, United States

Email: ezewu.onajite@yahoo.com

Corresponding Author's Email

¹ ezewu.onajite@yahoo.com

License:



Abstract

Artificial intelligence (AI) is becoming an operating layer beneath underwriting, pricing, catastrophe risk assessment, and claims administration in the U.S. Property and Casualty (P&C) insurance sector, yet the governance frameworks available to insurers remain either sector-agnostic or limited to compliance-level guidance. This paper introduces the Transparency, Risk ownership, Underlying data integrity, Surveillance, and Traceability (TRUST) Framework, a five-pillar, system-level governance model developed specifically for AI deployed in catastrophe-exposed insurance operations, together with three complementary original contributions: an AI Governance Maturity Model spanning five organizational stages, a Catastrophe AI Governance Lifecycle describing the continuous operational cycle through which TRUST is applied, and an Executive Decision Model linking governance quality to enterprise and national resilience outcomes. Drawing on literature spanning AI governance, explainable AI, algorithmic accountability, model risk management, catastrophe risk modelling, insurance regulation, operational resilience, and model lifecycle engineering, the paper identifies a research gap: existing frameworks address governance principles or regulatory compliance in the abstract but do not integrate system-level AI governance with the specific non-stationary risk conditions and regulatory examination trajectory characteristic of catastrophe-exposed insurance. Using qualitative conceptual framework development, comparative regulatory analysis, and case-study illustration grounded in the market disruption following Hurricane Ian, the paper positions TRUST as a framework that operationalizes, rather than replaces, existing standards including the NIST AI Risk Management Framework, ISO/IEC 42001, the OECD AI Principles, the NAIC Model Bulletin, COSO Enterprise Risk Management, and Federal Reserve model risk guidance. Limitations, practical implications, and directions for future empirical research are discussed.

Keywords: Artificial Intelligence Governance, Catastrophe Risk Intelligence, Insurance Regulation, Model Risk Management, Explainable AI, Algorithmic Accountability, Operational Resilience, Enterprise Resilience, Property and Casualty Insurance.

I. INTRODUCTION

Artificial intelligence (AI) has moved from a peripheral analytical enhancement to a central operating layer within catastrophe risk intelligence, shaping underwriting segmentation, portfolio aggregation analysis, nowcasting during active events, and claims triage in the U.S. Property and Casualty (P&C) insurance sector [1, 9]. This diffusion of AI into consequential, capital-relevant decisions has occurred more quickly than the development of governance structures specific to the operating conditions under which catastrophe-exposed insurance AI systems function. General-purpose AI governance instruments, including the National Institute of Standards and Technology (NIST) Artificial Intelligence Risk Management Framework (AI RMF) [18], the International Organization for Standardization (ISO)/IEC 42001 AI management system standard [19], and the Organisation for Economic Co-operation and Development (OECD) AI Principles [20], establish



valuable, broadly applicable risk-management structures, but were not designed to address the fundamental non-stationarity of catastrophe risk, the dense third-party data dependency characteristic of catastrophe modeling, or the distinctive, increasingly active state-based regulatory examination regime exemplified by the National Association of Insurance Commissioners (NAIC) Model Bulletin on the Use of Artificial Intelligence Systems by Insurers [21].

This paper introduces the Transparency, Risk ownership, Underlying data integrity, Surveillance, and Traceability (TRUST) Framework as an original, system-level governance model designed specifically for AI deployed in catastrophe-exposed insurance operations. The paper additionally introduces three complementary original contributions that extend TRUST from a governance standard into a practically applicable research program: an AI Governance Maturity Model, allowing organizations to assess governance sophistication across five discrete developmental stages; a Catastrophe AI Governance Lifecycle, describing the continuous operational cycle—from data collection through continuous improvement—through which TRUST is applied to an individual AI system over its operational life; and an Executive Decision Model, articulating the causal chain connecting AI governance quality to enterprise decision-making, regulatory confidence, public trust, and, ultimately, national resilience to catastrophic loss. Together, these four contributions constitute this paper's central scholarly argument: that trustworthy AI governance in catastrophe risk intelligence requires a framework purpose-built for the sector's distinctive risk and regulatory conditions, rather than a direct application of sector-agnostic governance principles.

The remainder of this paper is organized as follows. Section II reviews literature spanning AI governance, explainable AI, algorithmic accountability, model risk management, catastrophe risk modelling, insurance regulation, operational resilience, digital trust, human oversight, and model lifecycle governance. Section III identifies the specific research gap this paper addresses. Section IV establishes the paper's conceptual foundation. Section V describes the qualitative research methodology employed. Section VI presents the TRUST Framework and the Catastrophe AI Governance Lifecycle. Section VII presents the AI Governance Maturity Model and the Executive Decision Model. Section VIII presents a comparative analysis of TRUST against six established governance and risk-management frameworks. Section IX presents an expanded case study analysis grounded in Hurricane Ian. Section X discusses results, Section XI addresses practical implications, Section XII discusses limitations, Section XIII proposes future research directions, and Section XIV concludes.

II. LITERATURE REVIEW

A. AI Governance and Responsible AI

Scholarship on AI governance has evolved from principle-level ethics toward increasingly operational concerns. Mittelstadt et al. [1] provide a foundational map of the epistemic and normative concerns raised by algorithmic decision-making. Jobin, Ienca, and Vayena [9], in a systematic review of AI ethics guidelines, found convergence around five principles, transparency, justice and fairness, non-maleficence, responsibility, and privacy, but limited consensus on operationalization. Floridi et al.'s AI4People framework [8] similarly proposes unifying ethical principles at a level of abstraction requiring substantial translation before enterprise application. Mittelstadt [32], in a direct follow-up critique, argues that principles alone cannot guarantee ethical AI outcomes because AI development lacks the common aims, professional history, and enforcement mechanisms of more mature professions, a critique with particular force in insurance, an industry that already possesses mature regulatory and actuarial professional infrastructure into which sector-specific AI governance can, in principle, be embedded, provided a suitable framework exists.

Regulatory codification has progressed furthest through the NIST AI RMF [18], which organizes governance into four functions (Govern, Map, Measure, Manage); ISO/IEC 42001 [19], a certifiable AI management system standard; the OECD Principles [20], values-level intergovernmental guidance; and the European Union's Artificial Intelligence Act [33], which entered into force in 2024 and imposes binding, risk-tiered obligations on high-risk AI systems. None of these instruments, however, specifies how its general provisions should be adapted to catastrophe-exposed insurance operations specifically, a limitation elaborated in Section III.



B. Explainable AI and Algorithmic Accountability

A parallel literature addresses how algorithmic systems can be made explainable and held accountable in practice. Model-agnostic interpretability techniques including LIME [5] and SHAP [4] have become standard tools for post hoc explanation, while Doshi-Velez and Kim [6] caution that interpretability is not a single well-defined property. Rudin [13] argues that for high-stakes decisions—a category plainly including insurance underwriting and claims determinations—organizations should prefer inherently interpretable models over post hoc explanation wherever performance parity allows, a position in productive tension with prevailing industry practice. Ananny and Crawford [2] critique the "transparency ideal" itself, arguing that accountability depends on surrounding organizational structures rather than disclosure of algorithmic logic alone—a position that anticipates this paper's argument that governance must be embedded structurally (through TRUST) rather than achieved through explainability tooling in isolation. Raji et al. [3] propose SMACTR, an end-to-end internal algorithmic auditing framework drawing on aerospace and safety-critical systems engineering, directly analogous to TRUST's Traceability and Auditability pillar. Kroll et al. [12] argue for algorithms verifiable through technical and procedural means, and Diakopoulos [14] frames accountability as organizational and investigative practice. Wachter, Mittelstadt, and Floridi [7] and Selbst and Barocas [11] examine the legal limits of explainability rights, while Barocas and Selbst [10] and Veale and Binns [15] demonstrate that facially neutral models can reproduce discriminatory outcomes through correlated proxy variables—directly relevant to insurance underwriting, where variables such as property age or geography can proxy for protected characteristics.

C. Model Risk Management and Enterprise Risk Governance

The financial services sector developed a mature model risk management (MRM) discipline well before AI-specific governance discourse emerged. The Federal Reserve and Office of the Comptroller of the Currency's Supervisory Guidance on Model Risk Management (SR 11-7) [25] established model validation, independent review, and documentation as supervisory expectations in 2011. The COSO Enterprise Risk Management framework [26] provides a widely adopted structure integrating risk management with organizational strategy, and the Institute of Internal Auditors' Three Lines Model [34], critically examined by Eulerich and Eulerich [35], offers an accountability architecture separating operational, oversight, and assurance functions. These frameworks anticipate elements later echoed in AI-specific guidance, but were developed against a different failure-mode profile: traditional financial models are typically retrained infrequently against comparatively stationary relationships, whereas AI systems trained on evolving data exhibit drift and emergent bias patterns these frameworks do not explicitly contemplate.

D. Catastrophe Risk Modelling and Climate Risk Intelligence

Catastrophe risk modelling has an established scholarly foundation. Grossi and Kunreuther [17] provide a foundational treatment of hazard-exposure-vulnerability modelling methodology; Woo [36] examines the deeper probabilistic and epistemic foundations of catastrophe risk estimation, including model uncertainty under conditions of event rarity; and Kunreuther and Michel-Kerjan [16] examine the institutional and insurance-market response to large-scale catastrophic risk. This literature establishes, with unusual consensus, that catastrophe risk is fundamentally non-stationary: historical loss data is an imperfect guide to future risk given evolving climate patterns, land use, and construction practice. This property has direct, underexamined implications for AI models trained on historical catastrophe data, since such models may embed the same epistemic uncertainty implicitly, without the tradition of independent scrutiny applied to classical catastrophe models, a concern this paper's TRUST Framework addresses directly through its Surveillance and Monitoring pillar.

E. Insurance-Specific AI Regulation

Insurance-specific AI regulation is comparatively recent. The NAIC Model Bulletin [21], adopted in December 2023, requires insurers to maintain a written AI governance program, and had been adopted by roughly half of U.S. states by 2025 and 2026. Legal commentary [22] examines its practical implications, and industry reporting [23, 24] documents the NAIC's development of a multistate AI Systems Evaluation Tool for structured examiner assessment. The Geneva Association has separately tracked the international



regulatory landscape [37, 38], while Deloitte's industry survey evidence [39] documents a persistent gap between AI adoption and governance maturity among insurers. This literature remains predominantly descriptive and practitioner-oriented rather than peer-reviewed and conceptual.

F. Operational Resilience, Digital Trust, and Human Oversight

Operational resilience scholarship, exemplified by the Basel Committee on Banking Supervision's Principles for Operational Resilience [40], frames resilience as an outcome of governance over critical operations, third-party dependency, and incident management, concepts directly transferable to catastrophe-exposed insurers whose AI systems depend heavily on external data and modelling vendors, though the Basel principles, like SR 11-7, were developed for banking rather than insurance. On human oversight specifically, Green and Chen [41] provide empirical evidence that human decision-makers interacting with algorithmic recommendations ("algorithm-in-the-loop" processes) do not straightforwardly correct algorithmic errors, and that the design of the human-algorithm interaction materially affects outcomes, a finding with direct relevance to the human oversight expectations embedded in TRUST's Risk Ownership and Accountability pillar, and a caution against treating human review as an automatic safeguard rather than a designed control requiring its own governance.

G. Model Lifecycle Governance

A distinct engineering literature addresses the lifecycle practices through which machine learning systems are built and maintained. Amershi et al. [42], in an empirical study of Microsoft engineering teams, document a nine-stage machine learning workflow and identify data management, model versioning, and non-monotonic error behaviour as distinguishing engineering challenges relative to traditional software. This literature substantiates, from a software engineering perspective, the same lifecycle logic this paper's Catastrophe AI Governance Lifecycle formalizes for governance purposes: that AI systems require continuous, stage-specific oversight across their operational life rather than a single point-in-time validation, and that the transition points between lifecycle stages, particularly between deployment and ongoing monitoring, are where governance most commonly fails in practice.

III. RESEARCH GAP

The literature reviewed in Section II, read cumulatively, reveals six specific, interrelated gaps.

First, existing AI governance frameworks, NIST AI RMF [18], ISO/IEC 42001 [19], OECD Principles [20], and the EU AI Act [33], are sector-agnostic by design and do not specify adaptation to catastrophe-exposed insurance operations.

Second, the catastrophe risk modelling literature [16, 17, 36] predates the current AI governance discourse and does not engage directly with algorithmic accountability, explainability, or auditing frameworks developed in the intervening decade, leaving the intersection of catastrophe modelling and AI governance comparatively unexamined.

Third, model risk management and enterprise risk frameworks developed for banking [25, 26, 34, 35] and operational resilience frameworks developed for banking supervision [40] provide structural analogues but were not designed for insurance's distinctive non-stationary catastrophe risk environment or its third-party catastrophe-modelling data dependency.

Fourth, regulatory readiness and examination preparedness have received considerably less scholarly attention than AI ethics and algorithmic fairness. The bulk of the algorithmic accountability literature [2, 3, 12, 14] addresses accountability as a property of systems or organizations in the abstract, rather than as a maturity-assessable operational capability, as this paper's AI Governance Maturity Model proposes.

Fifth, human oversight research such as Green and Chen [41] demonstrates empirically that human review is not an automatic safeguard, yet this finding has not been systematically incorporated into insurance-specific AI governance guidance, which tends to treat human oversight as a stated requirement rather than a control requiring its own design and validation.

Sixth, and most fundamentally, no existing framework identified in this review integrates system-level AI governance principles, the specific non-stationarity of catastrophe risk, and the trajectory of insurance regulatory examination within a single, sector-specific model connecting governance quality to enterprise and



national resilience outcomes. This compound gap is what motivates the TRUST Framework, the AI Governance Maturity Model, the Catastrophe AI Governance Lifecycle, and the Executive Decision Model as this paper's original contributions, developed to extend—not replace—the frameworks reviewed above.

IV. CONCEPTUAL FOUNDATION

This paper's conceptual foundation rests on three propositions established by the literature reviewed in Section II. First, from the responsible AI literature [1, 2, 9, 32], principles alone do not produce governance outcomes; governance requires structural, organizational embedding. Second, from the catastrophe modelling literature [16, 17, 36], catastrophe risk is fundamentally non-stationary, meaning AI systems trained on historical catastrophe data face a validity-decay problem not adequately addressed by point-in-time validation practices common in general-purpose model governance. Third, from the model risk management and operational resilience literatures [25, 26, 40], mature accountability architectures already exist in adjacent regulated industries and provide a structural template, but not a substantive solution, for AI-specific governance in insurance.

These three propositions jointly imply that a fit-for-purpose AI governance framework for catastrophe risk intelligence must (a) translate abstract governance principles into system-level, verifiable pillars; (b) treat continuous monitoring as structurally necessary rather than optional, given catastrophe risk's non-stationarity; and (c) borrow the accountability architecture of adjacent regulated industries while adapting it to insurance-specific regulatory examination practice. The TRUST Framework, introduced in Section VI, is constructed directly from this three-part foundation.

V. RESEARCH METHODOLOGY

This paper employs a qualitative, conceptual research methodology appropriate to framework development in an applied governance domain. No survey, experiment, or interview-based primary data collection was conducted, and no such claim is made anywhere in this paper.

The methodology comprises four components. First, literature synthesis: peer-reviewed and authoritative gray literature spanning the seven domains reviewed in Section II was analysed to identify convergent themes and the specific research gap addressed in Section III. Second, comparative regulatory analysis: primary regulatory and standards texts, the NAIC Model Bulletin, NIST AI RMF, ISO/IEC 42001, OECD Principles, COSO ERM, and SR 11-7, were analysed directly to produce the comparative framework analysis in Section VIII. Third, enterprise governance analysis: the author's applied professional experience in enterprise technology governance for catastrophe risk intelligence platforms informed the structural design of TRUST, the Maturity Model, the Lifecycle, and the Executive Decision Model, consistent with established practice in applied conceptual framework development, where practitioner expertise legitimately informs framework design subject to grounding in the broader literature. Fourth, case-study illustration: the market disruption following Hurricane Ian is analysed in Section IX using exclusively public, cited sources, as an illustrative device for the governance questions this paper's frameworks are designed to answer—not as an empirical test of those frameworks, which have not yet been subjected to empirical validation.

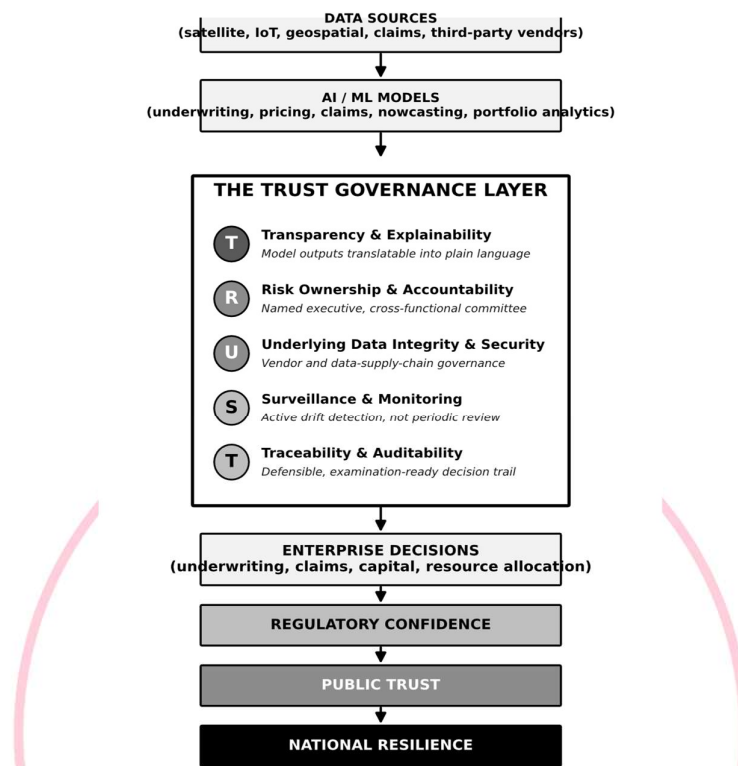
Consistent with this qualitative methodology, all statistics referenced throughout this paper are drawn from published, cited, third-party sources. No dataset, survey finding, or empirical result has been fabricated or estimated without citation.

VI. THE TRUST FRAMEWORK

The TRUST Framework is this paper's central original contribution: a five-pillar, system-level governance standard for AI deployed in catastrophe-exposed insurance operations. TRUST synthesizes the general risk-management functions of the NIST AI RMF with the documentation and accountability expectations of the NAIC Model Bulletin, sequenced specifically around the failure modes catastrophe risk intelligence faces, as established in Section II-D.



FIGURE 1
THE TRUST FRAMEWORK FOR AI GOVERNANCE IN CATASTROPHE RISK INTELLIGENCE



Transparency and Explainability requires model outputs to be translatable into terms usable by a regulator, claims examiner, or policyholder, drawing directly on the explainability literature reviewed in Section II-B while resolving its central tension (interpretable-by-design versus post hoc explanation) into a documentation requirement: insurers must specify, per model, which approach is taken and why.

Risk Ownership and Accountability requires a named executive and cross-functional governance committee for every AI system, addressing the human oversight concern raised by Green and Chen [41] by requiring that oversight be a designed, accountable structure rather than an assumed safeguard.

Underlying Data Integrity and Security extends governance to the data supply chain, satellite and geospatial vendors, IoT networks, third-party catastrophe modelling providers, recognizing that catastrophe risk intelligence's third-party dependency, discussed in Section II-F, is itself a governance surface.

Surveillance and Monitoring requires continuous, rather than point-in-time, model validation, directly operationalizing the non-stationarity finding from the catastrophe modelling literature (Section II-D) into a specific, ongoing governance obligation.

Traceability and Auditability requires a defensible evidence trail, model version, decision inputs, reviewer identity—for every consequential decision, aligning with the algorithmic auditing literature [3, 12] and directly supporting the examination-readiness discussed in Sections II-E and IX.

TRUST is designed to operate continuously across an AI system's operational life, not as a one-time certification. Figure 2 presents the Catastrophe AI Governance Lifecycle, a second original contribution of this paper, formalizing the eleven-stage operational cycle through which TRUST is applied and re-applied to a given AI system, from initial data collection through continuous improvement feeding back into subsequent iterations.



FIGURE 2
THE CATASTROPHE AI GOVERNANCE LIFECYCLE

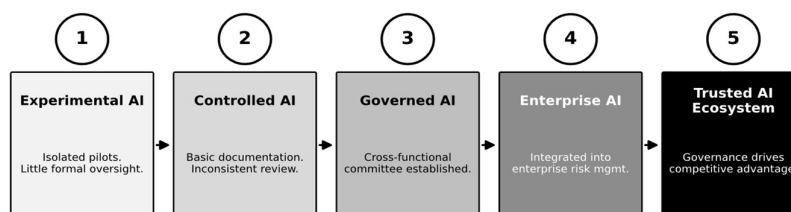


The lifecycle formalization is directly consistent with the model lifecycle engineering literature reviewed in Section II-G. Amershi et al. [42] find, empirically, that the transitions between lifecycle stages, rather than any single stage in isolation—are where AI engineering practice most commonly encounters difficulty; the Catastrophe AI Governance Lifecycle applies the same logic to governance, treating the transitions between Production Deployment, Continuous Monitoring, Model Drift Detection, and Regulatory Reporting as the specific points at which governance evidence must be actively produced, not merely assumed to carry over from the preceding stage.

VII. AI GOVERNANCE MATURITY MODEL

Where the TRUST Framework specifies governance content, the AI Governance Maturity Model, this paper's third original contribution—specifies organizational progression, allowing an insurer to assess its own governance sophistication against five discrete stages: Experimental AI, Controlled AI, Governed AI, Enterprise AI, and Trusted AI Ecosystem.

FIGURE 3
THE AI GOVERNANCE MATURITY MODEL FOR CATASTROPHE RISK INTELLIGENCE

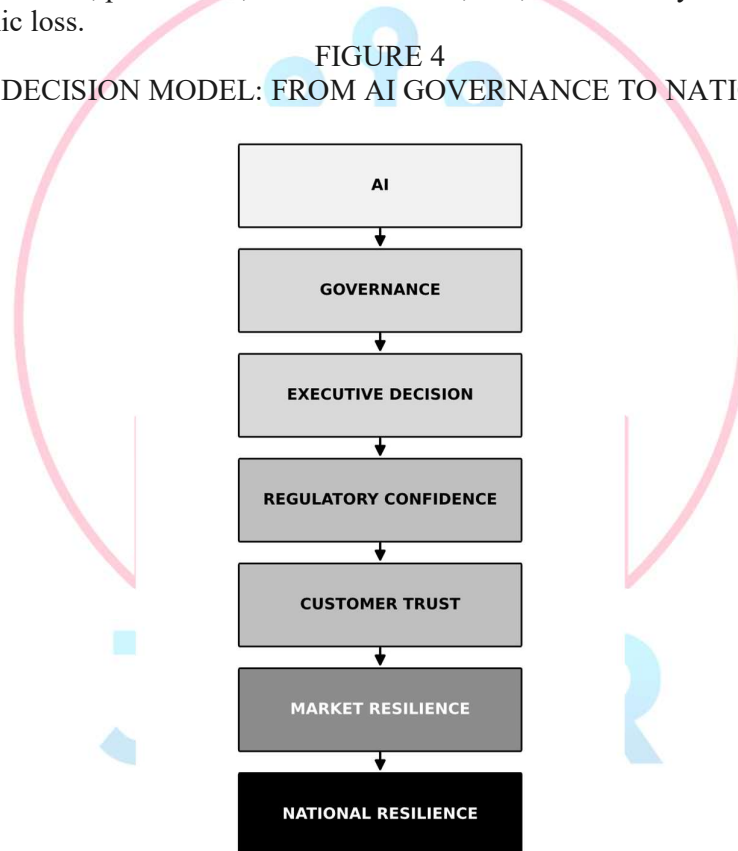




At the Experimental AI stage, pilots proceed with little formal oversight and no TRUST pillar is systematically applied; this is a normal starting condition, not itself a governance failure, unless an experimental system moves into production without corresponding oversight. At the Controlled AI stage, documentation exists but is inconsistently applied across use cases. At the Governed AI stage, a cross-functional governance committee exists with defined accountability and a model inventory. At the Enterprise AI stage, governance is integrated into enterprise risk management, with continuous monitoring operational rather than aspirational. At the Trusted AI Ecosystem stage, governance functions as a recognized differentiator to capital providers, reinsurers, and rating agencies, and the organization contributes to shaping emerging regulatory norms rather than only responding to them.

The Maturity Model's practical value depends on its final original contribution, the Executive Decision Model, which specifies the causal chain through which governance quality is hypothesized to translate into enterprise and societal outcomes: from the underlying AI system, through governance, to executive decision-making, regulatory confidence, public trust, market resilience, and, cumulatively across the industry, national resilience to catastrophic loss.

FIGURE 4
THE EXECUTIVE DECISION MODEL: FROM AI GOVERNANCE TO NATIONAL RESILIENCE



The Executive Decision Model is offered as a conceptual, not empirically validated, causal chain; Section XII addresses this limitation directly. Its scholarly function is to make explicit an assumption implicit throughout the AI governance literature reviewed in Section II—that governance quality has downstream consequences for trust and resilience—and to render that assumption specific and falsifiable for future empirical research, discussed in Section XIII.

VIII. COMPARATIVE FRAMEWORK ANALYSIS

TRUST is designed to complement, not replace, the six established governance and risk-management frameworks discussed throughout this paper. Table I presents a comparative analysis across dimensions relevant to catastrophe-exposed insurance: sector specificity, treatment of non-stationary risk, third-party data governance, and alignment with insurance regulatory examination practice.



TABLE I
COMPARATIVE FRAMEWORK ANALYSIS

Framework	Sector-Specific	Addresses Non-Stationary Risk	Third-Party Data Governance	Examination Alignment
NIST AI RMF [18]	No	No	Partial	Indirect
ISO/IEC 42001 [19]	No	No	Partial	Indirect
OECD Principles [20]	No	No	No	Indirect
NAIC Model Bulletin [21]	Yes	No	Partial	Direct
COSO ERM [26]	No	No	No	Indirect
SR 11-7 [25]	No (banking)	No	Partial	Indirect
TRUST (this paper)	Yes	Yes	Yes	Direct

NIST AI RMF excels at general-purpose risk-management structure but stops short of specifying sector adaptation. ISO/IEC 42001 excels at certifiable management-system rigor but, like NIST, remains sector-agnostic. The OECD Principles excel at values-level international consensus but require substantial translation before enterprise application. The NAIC Model Bulletin excels at insurance-specific regulatory precision but functions as a compliance requirement rather than a system-level governance standard, leaving open exactly how insurers should structure the underlying AI governance the bulletin requires them to document. COSO ERM excels at integrating risk management with organizational strategy but was not designed with AI-specific failure modes in view. SR 11-7 excels at model validation rigor but was designed for banking models that are typically retrained infrequently against comparatively stationary relationships, not the non-stationary, peril-driven risk environment catastrophe-exposed insurance AI systems must operate within.

TRUST's contribution, evident in Table I, is not superiority along any single dimension already served by an existing framework, but combined coverage of all four dimensions simultaneously—sector specificity, explicit treatment of non-stationary risk, third-party data governance, and direct regulatory examination alignment—none of which is offered by the other six frameworks in combination.

IX. CASE STUDY ANALYSIS

Published statistical data establish the broader trend against which the Hurricane Ian case study should be read. The National Centers for Environmental Information (NCEI) records 403 U.S. weather and climate disasters with individual losses exceeding USD 1 billion (Consumer Price Index-adjusted) between 1980 and 2024, with average annual frequency rising from 3.3 events per year in the 1980s to 23.0 events per year in the 2020–2024 period [30]—an increase consistent with the non-stationarity of catastrophe risk discussed in Section II-D. Global insured natural catastrophe losses recorded by Swiss Re Institute's sigma series exceeded USD 100 billion for six consecutive years through 2025, reaching USD 108 billion in 2023, USD 137 billion in 2024, and USD 107 billion in 2025 [27, 28, 29]. Over the same period, McKinsey & Company's global AI adoption survey found that 78 percent of organizations use AI in at least one business function, yet only approximately 6 percent qualify as "AI high performers" attributing more than 5 percent of enterprise earnings before interest and taxes to AI use [31]—an adoption-value gap consistent with this paper's premise that AI governance capability is lagging AI adoption itself. Hurricane Ian, analyzed below, is a single, well-documented instance of the broader pattern these statistics describe.

Hurricane Ian made landfall in southwest Florida on September 28, 2022, as a Category 4 storm, producing estimated insured losses of USD 50–65 billion and contributing to global insured natural catastrophe losses that reached USD 115 billion for the year [43]. This section analyzes the market disruption that followed Ian as an illustrative case for the governance questions the TRUST Framework is designed to answer, rather than as an empirical test of TRUST itself, since AI-assisted claims handling was not a documented feature of the Ian response specifically.

Florida's Office of Insurance Regulation recorded hundreds of thousands of residential claims following Ian, with a substantial share closed without payment and thousands of consumer complaints alleging



underpayment or denial [44, 46]. Several carriers active in the Florida market became insolvent in the storm's wake, and litigation costs on contested claims reached the hundreds of millions of dollars [45]. The broader market response was a sharp, sustained tightening of availability and a steep rise in premiums across hurricane-exposed counties, a pattern widely documented in subsequent industry and regulatory reporting.

This case study does not claim that AI caused, worsened, or was meaningfully present in Florida's post-Ian market disruption; Ian-era claims volume was handled predominantly through conventional adjustment processes, and the disruption reflected structural issues in Florida's insurance market—litigation costs and long-standing underinsurance—well documented independent of any AI system. The analytical value of the case lies elsewhere: Ian illustrates, in a well-documented public record, precisely the conditions under which AI-enabled catastrophe response is now being deployed, and precisely the governance failures the TRUST Framework is designed to prevent as AI's role in catastrophe response expands.

Four governance-relevant analytical points follow from the Ian record, mapped directly to TRUST's five pillars. First, on model drift: a claims triage or damage-estimation model trained on pre-2022 loss patterns would have faced a storm whose combined wind and storm-surge profile, in a market already under litigation strain, sat outside its training distribution—precisely the failure mode TRUST's Surveillance and Monitoring pillar is designed to catch through continuous, rather than point-in-time, validation. Second, on explainability under dispute: with claims volume in the hundreds of thousands and a meaningful share contested, any AI-assisted determination would face intense simultaneous scrutiny from policyholders, plaintiffs' counsel, and regulators, precisely the scenario TRUST's Transparency and Explainability pillar anticipates by requiring documented clarity on interpretability approach in advance of dispute. Third, on auditability under regulatory inquiry: Florida's OIR conducted structured data calls compelling insurers to report claims data during the Ian response [44], a preview of the documentation burden TRUST's Traceability and Auditability pillar is designed to make a matter of routine practice rather than crisis-driven reconstruction. Fourth, on enterprise resilience: the market-wide destabilization following Ian illustrates, at industry scale, the Executive Decision Model's central claim that governance failures at the individual-decision level aggregate into market-level and, cumulatively, national resilience consequences.

Ian is treated in this paper as a preview rather than an outlier: as AI's role in catastrophe response expands from nowcasting toward automated damage estimation and claims triage, the next comparably severe storm will test not only model accuracy but whether governance was mature enough, under the Maturity Model's terms, to keep that accuracy trustworthy under simultaneous operational, litigation, and regulatory stress.

X. DISCUSSION

Three principal findings emerge from the literature synthesis, comparative analysis, and case study presented above.

First, the comparative analysis in Table I confirms the research gap identified in Section III: none of the six established frameworks reviewed combines sector specificity, explicit treatment of non-stationary catastrophe risk, third-party data governance, and direct regulatory examination alignment. TRUST is the only framework in the comparison satisfying all four criteria simultaneously, supporting its status as a genuine extension of, rather than a restatement of, the existing literature.

Second, the Hurricane Ian case study substantiates, through a well-documented public record, the practical consequences of the non-stationarity concern raised in the catastrophe modeling literature [16, 17, 36]. The scale of market disruption following a single severe event illustrates why point-in-time model validation is structurally insufficient for catastrophe-exposed AI systems, reinforcing the necessity of TRUST's Surveillance and Monitoring pillar and the Catastrophe AI Governance Lifecycle's treatment of monitoring as continuous rather than periodic.

Third, the human oversight literature [41] and the model lifecycle engineering literature [42] jointly suggest that governance failures cluster at specific, identifiable points—the human-algorithm interaction boundary and the transitions between lifecycle stages—rather than being uniformly distributed across a system's operation. This finding is consistent with, and lends empirical support to, TRUST's decision to



specify five distinct pillars rather than a single undifferentiated governance requirement, since each pillar targets a different point at which governance is empirically most likely to fail.

XI. PRACTICAL IMPLICATIONS

For insurance executives, this analysis suggests that AI governance investment should be structured around the five TRUST pillars individually, since the literature and case study evidence indicate that governance failures are pillar-specific rather than general, and that a maturity assessment conducted pillar by pillar will surface gaps a single aggregate assessment would obscure.

For boards and risk committees, the Executive Decision Model offers a structured basis for board-level reporting on AI governance, translating technical governance activity into the resilience and trust outcomes a board is ultimately accountable for.

For regulators, the comparative analysis in Section VIII suggests that sector-specific system-level standards such as TRUST may usefully complement, rather than substitute for, principle-level and compliance-level guidance such as the NAIC Model Bulletin [21], particularly as multistate examination tooling continues to mature [23].

XII. LIMITATIONS

This paper is subject to several limitations. First, the research methodology is conceptual and qualitative; TRUST, the AI Governance Maturity Model, the Catastrophe AI Governance Lifecycle, and the Executive Decision Model have not been subjected to empirical validation through survey, case study of a documented implementation, or controlled comparison, and no such claim is made.

Second, the Executive Decision Model in particular proposes a causal chain—from governance quality to enterprise decision-making, regulatory confidence, public trust, market resilience, and national resilience—that is conceptually motivated by the literature reviewed in Section II but has not been empirically tested; the model should be read as a structured hypothesis for future research rather than a demonstrated causal finding.

Third, the Hurricane Ian case study is illustrative rather than confirmatory: as stated in Section IX, AI-assisted claims handling was not a documented feature of the actual Ian response, and the case study's analytical value lies in the governance questions it raises rather than in any claim that TRUST would have altered the actual outcome.

Fourth, insurance sector-specific AI governance maturity data are not, to the author's knowledge, available in sufficiently granular public form to permit direct quantitative benchmarking of the Maturity Model against observed industry practice.

XIII. FUTURE RESEARCH

Four directions for future research follow from the limitations identified in Section XII.

First, empirical case study research examining actual insurer implementations of TRUST or comparable system-level AI governance frameworks would test the structural claims made in this paper against observed organizational outcomes.

Second, survey-based research empirically testing the Executive Decision Model's proposed causal chain—for example, examining whether documented governance maturity is associated with more favorable regulatory examination outcomes or lower loss volatility—would convert this paper's conceptual contribution into a testable empirical research program.

Third, research quantifying insurance-specific AI governance maturity, ideally conducted in partnership with regulatory bodies such as the NAIC, would address the benchmarking data gap identified in Section XII.

Fourth, research extending TRUST and the Catastrophe AI Governance Lifecycle to emerging AI categories—particularly autonomous, agentic AI systems capable of taking bounded actions in catastrophe response without direct human sign-off—would test the framework's durability as the underlying technology evolves beyond the recommendation-generating systems that motivated its initial design.

XIV. CONCLUSION

This paper has introduced the TRUST Framework as an original, system-level AI governance model purpose-built for catastrophe-exposed Property and Casualty insurance, together with three complementary



original contributions—the AI Governance Maturity Model, the Catastrophe AI Governance Lifecycle, and the Executive Decision Model—that extend TRUST from a governance standard into a practically applicable, conceptually coherent research program. Drawing on literature spanning AI governance, explainable AI, algorithmic accountability, model risk management, catastrophe risk modeling, insurance regulation, operational resilience, and model lifecycle engineering, this paper identified a specific, underexamined research gap: no existing framework integrates system-level AI governance with catastrophe risk's distinctive non-stationarity and the trajectory of insurance regulatory examination. The comparative analysis in Section VIII and the Hurricane Ian case study in Section IX together demonstrate that this gap has practical, not merely theoretical, consequence. While the frameworks proposed here remain conceptual and have not yet been empirically validated, they are offered as a practically grounded, academically defensible foundation for the empirical research program outlined in Section XIII, and as a structural resource for insurance executives, boards, and regulators navigating the operational demands of AI governance in an increasingly examination-intensive regulatory environment.

REFERENCES

- [1] B. D. Mittelstadt, P. Allo, M. Taddeo, S. Wachter, and L. Floridi, “The ethics of algorithms: Mapping the debate,” *Big Data & Society*, vol. 3, no. 2, 2016.
- [2] M. Ananny and K. Crawford, “Seeing without knowing: Limitations of the transparency ideal and its application to algorithmic accountability,” *New Media & Society*, vol. 20, no. 3, pp. 973–989, 2018.
- [3] I. D. Raji, A. Smart, R. N. White, M. Mitchell, T. Gebru, B. Hutchinson, J. Smith-Loud, D. Theron, and P. Barnes, “Closing the AI accountability gap: Defining an end-to-end framework for internal algorithmic auditing,” in *Proc. 2020 Conf. Fairness, Accountability, and Transparency (FAT* '20)*, Barcelona, Spain, 2020, pp. 33–44.
- [4] S. M. Lundberg and S.-I. Lee, “A unified approach to interpreting model predictions,” in *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 30, 2017.
- [5] M. T. Ribeiro, S. Singh, and C. Guestrin, “‘Why should I trust you?’: Explaining the predictions of any classifier,” in *Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining (KDD)*, 2016, pp. 1135–1144.
- [6] F. Doshi-Velez and B. Kim, “Towards a rigorous science of interpretable machine learning,” *arXiv preprint arXiv:1702.08608*, 2017.
- [7] S. Wachter, B. Mittelstadt, and L. Floridi, “Why a right to explanation of automated decision-making does not exist in the General Data Protection Regulation,” *International Data Privacy Law*, vol. 7, no. 2, pp. 76–99, 2017.
- [8] L. Floridi et al., “AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations,” *Minds and Machines*, vol. 28, no. 4, pp. 689–707, 2018.
- [9] A. Jobin, M. Ienca, and E. Vayena, “The global landscape of AI ethics guidelines,” *Nature Machine Intelligence*, vol. 1, pp. 389–399, 2019.
- [10] S. Barocas and A. D. Selbst, “Big data’s disparate impact,” *California Law Review*, vol. 104, pp. 671–732, 2016.
- [11] A. D. Selbst and S. Barocas, “The intuitive appeal of explainable machines,” *Fordham Law Review*, vol. 87, pp. 1085–1139, 2018.
- [12] J. A. Kroll, J. Huey, S. Barocas, E. W. Felten, J. R. Reidenberg, D. G. Robinson, and H. Yu, “Accountable algorithms,” *University of Pennsylvania Law Review*, vol. 165, no. 3, pp. 633–705, 2017.
- [13] C. Rudin, “Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead,” *Nature Machine Intelligence*, vol. 1, pp. 206–215, 2019.
- [14] N. Diakopoulos, “Accountability in algorithmic decision making,” *Communications of the ACM*, vol. 59, no. 2, pp. 56–62, 2016.
- [15] M. Veale and R. Binns, “Fairer machine learning in the real world: Mitigating discrimination without collecting sensitive data,” *Big Data & Society*, vol. 4, no. 2, 2017.



- [16] H. Kunreuther and E. Michel-Kerjan, *At War with the Weather: Managing Large-Scale Risks in a New Era of Catastrophes*. Cambridge, MA, USA: MIT Press, 2009.
- [17] P. Grossi and H. Kunreuther, Eds., *Catastrophe Modeling: A New Approach to Managing Risk*. New York, NY, USA: Springer, 2005.
- [18] National Institute of Standards and Technology, *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, U.S. Dept. of Commerce, Gaithersburg, MD, USA, 2023.
- [19] International Organization for Standardization, *ISO/IEC 42001:2023—Information Technology—Artificial Intelligence—Management System*, ISO, Geneva, Switzerland, 2023.
- [20] Organisation for Economic Co-operation and Development, *OECD Principles on Artificial Intelligence*, OECD, Paris, France, 2019.
- [21] National Association of Insurance Commissioners, *Model Bulletin: Use of Artificial Intelligence Systems by Insurers*, NAIC, Kansas City, MO, USA, Dec. 2023.
- [22] Kennedys Law LLP, “Understanding the NAIC model AI bulletin: What it means for insurers,” Jan. 21, 2025.
- [23] Fenwick, “NAIC expands AI Systems Evaluation Tool pilot program to 12 states: Key updates for insurers and AI vendors supporting insurers,” Mar. 9, 2026.
- [24] Crowell & Moring LLP, “NAIC intensifies AI regulatory focus: What health insurance payors need to know,” Mar. 25, 2026.
- [25] Board of Governors of the Federal Reserve System and Office of the Comptroller of the Currency, *Supervisory Guidance on Model Risk Management (SR 11-7)*, Washington, DC, USA, 2011.
- [26] Committee of Sponsoring Organizations of the Treadway Commission (COSO), *Enterprise Risk Management—Integrating with Strategy and Performance*, 2017.
- [27] Swiss Re Institute, “sigma 1/2024: Natural catastrophes in 2023—gearing up for today’s and tomorrow’s weather risks,” Zurich, Switzerland, Mar. 2024.
- [28] Swiss Re Institute, “sigma 1/2025: Natural catastrophes: insured losses on trend to USD 145 billion in 2025,” Zurich, Switzerland, Apr. 2025.
- [29] Swiss Re Institute, “sigma 1/2026: Natural catastrophes in 2025—the persistent rise of wildfire and storm risk,” Zurich, Switzerland, 2026.
- [30] National Centers for Environmental Information, “Billion-dollar weather and climate disasters,” National Oceanic and Atmospheric Administration (NOAA), Asheville, NC, USA, 2025.
- [31] McKinsey & Company, “The state of AI: Global survey,” QuantumBlack, AI by McKinsey, Nov. 2025.
- [32] B. Mittelstadt, “Principles alone cannot guarantee ethical AI,” *Nature Machine Intelligence*, vol. 1, no. 11, pp. 501–507, 2019.
- [33] European Parliament and Council of the European Union, *Regulation (EU) 2024/1689 Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act)*, Official Journal of the European Union, Jul. 2024.
- [34] The Institute of Internal Auditors, *The IIA’s Three Lines Model: An Update of the Three Lines of Defense*, IIA, Lake Mary, FL, USA, 2020.
- [35] M. Eulerich and A. Eulerich, “The new three lines model for structuring corporate governance—A critical discussion of similarities and differences,” *Corporate Ownership & Control*, vol. 18, no. 2, pp. 218–225, 2020.
- [36] G. Woo, *Calculating Catastrophe*. London, U.K.: Imperial College Press, 2011.
- [37] B. Keller, *Promoting Responsible Artificial Intelligence in Insurance*, The Geneva Association, Zurich, Switzerland, 2020.
- [38] I. Flückiger and K.-U. Schanz, *Regulation of Artificial Intelligence in Insurance: Balancing Consumer Protection and Innovation*, The Geneva Association, Zurich, Switzerland, 2023.
- [39] Deloitte, “The AI imperative in insurance,” presented to the Federal Advisory Committee on Insurance, U.S. Dept. of the Treasury, Washington, DC, USA, Mar. 20, 2024.



- [40] Basel Committee on Banking Supervision, Principles for Operational Resilience, Bank for International Settlements, Basel, Switzerland, Mar. 2021.
- [41] B. Green and Y. Chen, “The principles and limits of algorithm-in-the-loop decision making,” Proceedings of the ACM on Human-Computer Interaction, vol. 3, no. CSCW, Article 50, pp. 1–24, 2019.
- [42] S. Amershi, A. Begel, C. Bird, R. DeLine, H. Gall, E. Kamar, N. Nagappan, B. Nushi, and T. Zimmermann, “Software engineering for machine learning: A case study,” in Proc. 2019 IEEE/ACM 41st Int. Conf. Software Engineering: Software Engineering in Practice (ICSE-SEIP), Montreal, Canada, 2019, pp. 291–300.
- [43] Swiss Re Institute, “Hurricane Ian drives natural catastrophe year-to-date insured losses to USD 115 billion,” Zurich, Switzerland, Dec. 1, 2022.
- [44] State of Florida Office of Insurance Regulation, “Hurricane Ian catastrophe claims data,” [flor.gov](https://floridaregulation.com/), 2024.
- [45] Older & Lundy, “Hurricane Ian insurance claim deadline,” May 23, 2024.
- [46] United Policyholders, “Resolving Hurricane Ian insurance disputes,” Oct. 16, 2023.

