

RELIGIOUS EQUITY IN PUBLIC-SECTOR ARTIFICIAL INTELLIGENCE: DEVELOPING A MULTI-FAITH AI IMPACT ASSESSMENT FRAMEWORK THROUGH A COMPARATIVE STUDY OF TEXAS AND CALIFORNIA

*Ataur Rehman¹

¹ Graduate Researcher, Indiana Wesleyan University, USA.



HEC
"Y"
Category



ARTICLE INFO

Article History:

Received: May 26, 2025
Revised: June 19, 2025
Accepted: June 23, 2025
Available Online: June 25, 2025

Keywords:

Artificial Intelligence Governance
Religious Equity & Interfaith Ethics
Texas & California
Public-Sector AI & AI Impact Assessment
Algorithmic Bias

Funding:

This research journal (PIJISS) doesn't receive any specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Copyrights:

Copyright Muslim Intellectuals Research Center. All Rights Reserved © 2021.
This work is licensed under a Creative Commons Attribution 4.0 International License.



ABSTRACT

Artificial intelligence is increasingly embedded in public administration, education, health, benefits delivery, public communication, procurement, and regulatory decision-making. Yet state AI-governance frameworks rarely identify religion as a distinct domain of algorithmic equity, even though public systems can misclassify religious identity, reproduce stereotypes, fabricate beliefs, disadvantage minority-faith practices, or amplify sectarian misinformation. This article develops the Multifaith Artificial Intelligence Impact Assessment Framework (MAILIAF), a seven-domain instrument for evaluating religious representation, disparate impact, hallucination, community harm, transparency, remedy, and institutional accountability. Using structured comparative policy analysis, the article examines Texas and California as contrasting but nationally influential state models. Texas provides centralized statutory governance architecture, an AI council, state-use notice requirements, and an innovation sandbox. California offers a more operationalized procurement and risk-assessment system, mandatory inventories, vendor disclosures, workforce training, and lifecycle controls. The comparison finds that neither state presently provides a religion-specific testing protocol, affected-community consultation requirement, or standardized remedy for faith-related AI harms. MAILIAF addresses this gap by translating constitutional, civil-rights, risk-management, and interfaith-ethics principles into an auditable lifecycle process. The proposed framework is designed for adaptation by state agencies, school systems, public universities, health administrators, and technology vendors throughout the United States. Its national value lies in providing a replicable method for protecting religious liberty and equal treatment while supporting responsible innovation.

*Corresponding Author's Email: ataur.rehman@myemail.indwes.edu

INTRODUCTION

Artificial intelligence has moved from experimental use to institutional infrastructure. State governments now employ or pilot generative and predictive systems for customer service, transportation, fiscal analysis, fraud detection, workforce management, education, health administration, and document processing. These uses can increase speed and consistency, but they can also convert historical inequalities, incomplete data, and cultural assumptions into automated outputs that affect real people. Public-sector deployment therefore raises a basic democratic question: how can government obtain the benefits of AI without allowing automated systems to weaken constitutional liberty, equal treatment, and public trust?

Religion presents a particularly underdeveloped area of AI governance. Religious identity is often absent from training data documentation and fairness benchmarks. When it does appear, it may be inferred from names, language, clothing, dietary practices, location, associations, or patterns of worship. Generative models may also invent doctrines, falsely attribute extremist positions, conflate ethnicity with belief, or treat majority-faith practices as neutral while marking minority practices as exceptional. These errors become especially serious when AI is used in public benefits, education, employment, correctional settings, public safety, healthcare, and official communication.

Interfaith scholarship offers a useful ethical foundation for addressing these risks. Research on common moral commitments across religious traditions emphasizes dignity, peaceful coexistence, social responsibility, and protection from harm (Raza & Khalid, 2022). Work on the social role of religion similarly shows that religious institutions can either reduce antisocial behavior and strengthen solidarity or become targets of exclusion when social systems are poorly designed (Jasvi et al., 2024). These insights support a governance approach in

which religious communities are treated not as abstract demographic categories, but as affected stakeholders whose rights, knowledge, and lived experiences must inform AI assessment.

This article compares the evolving AI-governance models of Texas and California and develops a new Multifaith Artificial Intelligence Impact Assessment Framework (MAIIAF). The states are analytically valuable because both are economically and technologically influential, yet their governance strategies differ. Texas emphasizes statutory prohibitions, centralized oversight, innovation, and a regulatory sandbox. California emphasizes procurement controls, inventories, risk classification, vendor disclosure, training, and detailed administrative guidance. Studied together, they reveal complementary strengths and a shared blind spot: neither state has established a systematic method for evaluating AI-related religious inequity. The article makes three contributions. First, it introduces religious equity as a distinct public-sector AI governance problem. Second, it provides a transparent comparative assessment of Texas and California. Third, it proposes an operational framework that can be adopted and tested across jurisdictions. The purpose is not to claim that religion should receive privileged treatment. It is to ensure that AI systems do not diminish the protections already owed to people and communities under constitutional, statutory, and ethical norms.

RESEARCH QUESTIONS AND NATIONAL SIGNIFICANCE

The study addresses four questions:

RQ1. How do Texas and California structure public-sector AI governance, risk assessment, transparency, oversight, and accountability?

RQ2. To what extent do those systems protect religious liberty, religious diversity, and equal treatment?

RQ3. What governance elements are missing when AI systems generate or act upon religion-related information?

RQ4. How can a multifaith impact-assessment framework be designed for adoption beyond a single state?

The national significance of these questions arises from the scale and portability of state AI governance. State procurement standards influence vendors that sell nationally. State risk-assessment forms become templates for local governments, universities, and

school districts. State transparency laws shape model documentation and product design. A practical framework developed through two major-state cases can therefore produce effects extending beyond either jurisdiction. The national problem is not merely technical error; it is the possibility that government-supported AI could systematically misrepresent faith communities, burden religious practice, distribute services unequally, or erode trust among millions of residents.

LITERATURE REVIEW AND NORMATIVE FOUNDATION

AI Bias as a Sociotechnical Problem

NIST treats AI bias as a sociotechnical problem that can arise from data, institutional processes, human assumptions, deployment context, and feedback loops rather than from code alone (Schwartz et al., 2022). The AI Risk Management Framework organizes trustworthy governance around the functions Govern, Map, Measure, and Manage (NIST, 2023). The Generative AI Profile further emphasizes structured feedback, representative evaluation participants, independent testing, documentation, appeal procedures, and post-deployment monitoring (NIST, 2024). These principles are highly relevant to religious equity because many failures originate not in explicit religious classification but in proxies, context loss, or unrepresentative evaluation.

Conventional fairness metrics are necessary but insufficient. A system may show similar aggregate accuracy across demographic groups while still generating qualitatively harmful statements about sacred figures, practices, or communities. It may also be formally neutral but impose practical burdens, such as scheduling systems that disregard religious observance or identity-verification tools that fail for religious clothing. Religious-equity assessment therefore requires both quantitative testing and contextual expert review. The need for context-sensitive interpretation is consistent with comparative interfaith scholarship that emphasizes both shared values and meaningful differences among traditions (Raza & Khalid, 2022).

Religious Liberty, Pluralism, and Interfaith Ethics

The First Amendment protects free exercise and prohibits governmental establishment of religion. Federal and state civil-rights laws also prohibit forms of religious discrimination in employment,

education, housing, public accommodation, and federally supported programs. AI does not displace these obligations. Instead, it creates new pathways through which burdens or exclusions may arise.

Interfaith ethics helps translate legal minimums into responsible institutional practice. Raza and Khalid (2022) argue that the Abrahamic traditions contain significant ethical commonalities that support peaceful coexistence and moral responsibility. Rasool and colleagues emphasize that religion can contribute to social order when it is connected to welfare, dignity, and constructive community engagement (Jasvi et al., 2024). These perspectives support the inclusion of multifaith expertise in testing, not to settle theological truth, but to identify misrepresentation, stereotyping, and foreseeable community harm.

Comparative religious scholarship also demonstrates why a single generic 'religion' label is inadequate. Traditions contain internal diversity, contested interpretations, different authority structures, and distinct practices. A model that performs adequately on broad references to Christianity may fail for smaller Christian denominations; a system that recognizes Islam may still misrepresent Sunni, Shi'a, Ahmadi, or other communities. The same problem applies across Jewish, Hindu, Buddhist, Sikh, Indigenous, and nonreligious worldviews. Responsible assessment must therefore combine general safeguards with context-specific review.

Religious peacebuilding research is also relevant to information integrity. False attribution, decontextualized scripture, fabricated rulings, and synthetic impersonation of religious leaders can inflame hostility at speed and scale. The social-harmony emphasis found in the work of Raza and colleagues provides a basis for treating such harms as governance risks rather than merely undesirable content (Raza et al., 2022).

METHODOLOGY

This article uses structured comparative policy analysis. Primary sources include enacted statutes, official state technology policies, procurement guidance, risk-assessment requirements, executive materials, and federal AI-risk guidance available through June 2026. The unit of analysis is the governance mechanism rather than the political ideology of either state.

Seven dimensions were coded on a 0–5 maturity scale: governance structure; pre-deployment risk assessment; transparency and disclosure; community

participation; religious-equity specificity; appeal and remedy; and post-deployment monitoring. A score of 0 indicates no identifiable mechanism; 1 indicates an indirect or aspirational reference; 2 indicates a limited mechanism; 3 indicates a defined but incomplete mechanism; 4 indicates a strong mechanism with institutional responsibility; and 5 indicates a detailed, operational, and repeatable requirement. Scores are interpretive policy-analysis judgments, not measurements of actual agency performance. They are included to make the comparison transparent and reproducible.

The analysis has three limitations. First, state implementation is evolving. Second, public documentation may not capture every agency practice. Third, the study evaluates governance design rather than model-level outcomes. These limitations strengthen, rather than weaken, the case for a standardized assessment instrument that agencies can apply to concrete systems.

TEXAS: STATUTORY GOVERNANCE, INNOVATION, AND CENTRALIZED OVERSIGHT

Texas enacted the Texas Responsible Artificial Intelligence Governance Act through House Bill 149, effective January 1, 2026. The statute states that its purposes include advancing responsible AI, protecting individuals and groups from known and reasonably foreseeable risks, increasing transparency, and providing notice regarding state-agency AI use. It creates duties and prohibitions, establishes a regulatory sandbox, and creates the Texas Artificial Intelligence Council.

The Texas model has several notable strengths. First, it places AI governance in statute rather than relying only on administrative guidance. Second, it creates a statewide council with authority to study ethical, legal, privacy, security, liability, and regulatory issues. Third, it requires training and educational outreach for state agencies and local governments. Fourth, it creates a sandbox intended to support testing in sectors that include healthcare, finance, education, and public services. Fifth, it requires evaluation and inventory-related attention to state-agency AI systems.

For religious equity, however, the framework is indirect. The law protects individuals and groups and addresses manipulation, discrimination, freedom, and public interest, but it does not establish a

religion-specific testing category. It does not require agencies to examine proxies for religious identity, consult affected faith communities, test generated descriptions of religious traditions, or create a specialized appeal path for religion-related errors. The council may appoint advisers with ethical and regulatory expertise, yet multifaith representation is not an express component.

Texas is therefore strongest at the level of statewide architecture and innovation governance. Its principal opportunity is to connect the broad statutory purposes to operational assessment. MAIIAF could be incorporated into council guidance, agency training, sandbox admission, procurement review, or voluntary state standards without undermining the statute's innovation-oriented design.

CALIFORNIA: PROCUREMENT CONTROLS, RISK CLASSIFICATION, AND LIFECYCLE ADMINISTRATION

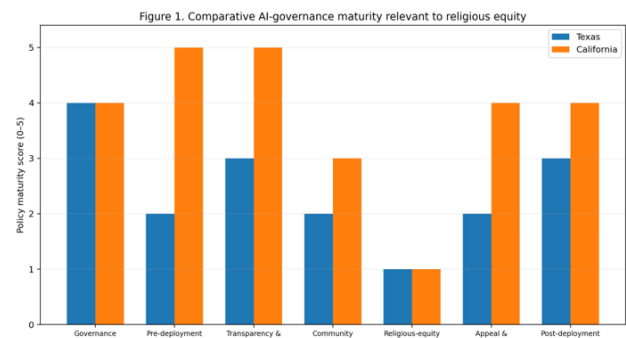
California has developed a more administratively detailed public-sector GenAI governance system. Official policies require state entities to maintain GenAI inventories, include disclosure language in solicitations, obtain vendor factsheets, complete risk assessments, and consult the California Department of Technology for moderate- and high-risk acquisitions. The state also provides workforce training, procurement guidance, and lifecycle documentation.

California's risk-based approach is a major strength. The GenAI Risk Assessment requires agencies to classify systems and evaluate risks before acquisition or deployment. Procurement controls make vendor disclosure part of the process, while inventories improve institutional visibility. The state has also linked its approach to NIST's AI RMF and has continued to update forms and guidance as use cases mature.

California's framework is comparatively stronger on operational transparency and repeatability. Yet its religious-equity coverage remains limited. Fairness, accessibility, privacy, security, and civil rights are addressed at a general level, but the available instruments do not provide a dedicated test for religious representation, faith-related hallucination, or burdens on religious exercise. Community consultation is encouraged in broader responsible-AI practice, but there is no standardized requirement to

involve minority-faith stakeholders when a system has a plausible impact on them.

California is therefore well positioned to add a MAIIAF module to its existing risk-assessment workflow. The framework could function as a conditional annex triggered when a system processes identity, language, cultural content, education, public benefits, correctional services, healthcare, or communications likely to affect religious communities.



Note. Scores are the author's structured policy-coding judgments based on publicly available state instruments through June 2026. They evaluate governance design, not implementation outcomes.

COMPARATIVE ANALYSIS

Dimension	Texas	California	Comparative finding
Legal-administrative foundation	Comprehensive statute and council	Executive, statutory, procurement, and technology-policy system	Texas is more centralized in statute; California is more operational in administrative procedure.
Pre-deployment assessment	Limited general risk assessment architecture	Mandatory GenAI risk assessment for procurement	California provides the more mature repeatable assessment process.
Transparency	Government notice, reporting, council studies	Inventories, vendor disclosures, factsheets, procurement documentation	California has broader operational disclosure; Texas has strong statewide reporting potential.
Innovation mechanism	Express regulatory sandbox	Pilots, sandboxes, innovative procurement	Both support experimentation; Texas places sandbox authority in statute.
Public participation	General public and expert participation possibilities	Stakeholder engagement and administrative consultation	Neither establishes a consistent affected-faith-community requirement.
Religious-equity safeguards	No dedicated protocol	No dedicated protocol	This is the clearest shared policy gap.
Remedy and appeal	General enforcement and agency processes	General administrative and procurement controls	Neither offers a standardized religion-related AI correction and appeal procedure.

The comparison reveals complementarity rather than a simple ranking. Texas supplies durable statewide architecture, a formal council, and an explicit innovation pathway. California supplies detailed procurement controls, inventories, risk classification, and administrative repeatability. A national model could combine Texas's statutory coherence and sandbox structure with California's operational risk-assessment process.

The most significant common weakness is religious-equity specificity. Both states rely on broad concepts

such as fairness, freedom, discrimination, ethics, and public interest. Those concepts are essential, but agencies need concrete questions and evidence requirements. Without them, religious harms may remain invisible because evaluators may not know what to test.

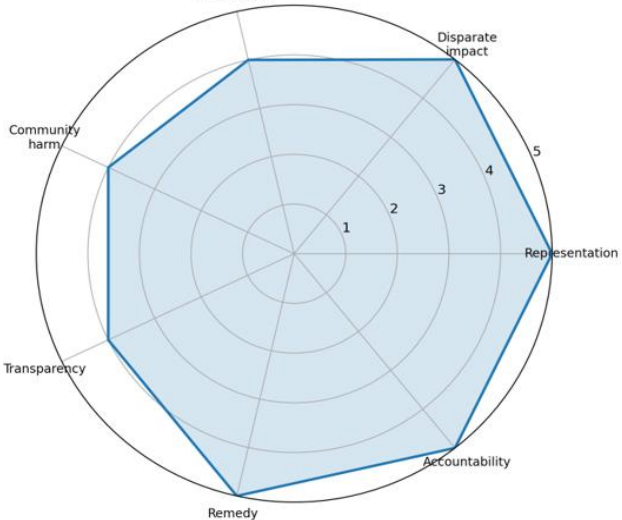
A second gap concerns remedy. Public trust requires more than disclosure that AI is being used. Affected persons need a practical way to challenge fabricated religious information, discriminatory outcomes, proxy-based classifications, or burdens on observance. NIST's emphasis on appeal and feedback provides a foundation, but state procedures should define responsibility, timelines, documentation, and human review.

A third gap is stakeholder composition. The ethical common-ground approach described by Raza and Khalid (2022) suggests that multifaith consultation can identify shared safeguards without requiring government to endorse theological claims. Such consultation should include minority traditions and intra-faith diversity. Rasool's collaborative research on religion's social role similarly supports engagement with institutions that understand community-level consequences (Jasvi et al., 2024).

THE MULTIFAITH ARTIFICIAL INTELLIGENCE IMPACT ASSESSMENT FRAMEWORK

MAIIAF is designed as a modular assessment that can be incorporated into existing AI risk-management systems. It contains seven domains, each scored from 0 to 5. A deployment may not proceed when a high-severity risk lacks an identified mitigation, accountable owner, monitoring plan, and remedy pathway.

Figure 2. MAIIAF assessment domains



Domain	Core question	Illustrative evidence	High-risk indicator
1. Representational accuracy	Does the system accurately describe traditions, practices, leaders, texts, and internal diversity?	Expert review, benchmark prompts, source traceability	Fabricated doctrine, stereotypes, or systematic confusion among traditions
2. Disparate treatment	Do actual or inferred religious identities receive different outcomes or burdens?	Outcome testing, proxy analysis, subgroup error rates	Unequal denial, delay, scrutiny, or access
3. Hallucination and attribution	Does the system invent quotations, rulings, affiliations, or acts?	Citation verification, retrieval tests, adversarial prompts	Confident false attribution to a religion or leader
4. Community harm	Could outputs foreseeably intensify hostility, exclusion, or extremist manipulation?	Scenario testing, threat modeling, community review	Synthetic incitement, impersonation, or dehumanizing narratives
5. Transparency	Can users understand when AI is used, what data matters, and the system's limits?	Notices, model cards, decision explanations	Hidden automation or unverifiable religious claims
6. Remedy and contestability	Can affected persons correct errors and obtain timely human review?	Appeal workflow, service standards, correction log	No accountable contact or meaningful reversal process
7. Institutional accountability	Are roles, monitoring, audits, procurement duties, and reporting clearly assigned?	Risk register, audit reports, vendor clauses, incident log	No owner, no monitoring, or no post-deployment review

Trigger Conditions

A MAIIAF review should be mandatory when an AI system: (a) collects or infers religious identity; (b) generates educational, legal, historical, or public information about religion; (c) affects benefits, employment, discipline, housing, health, corrections, licensing, or public safety; (d) uses names, language, geography, clothing, diet, or associations as decision variables; (e) creates synthetic media involving religious leaders or institutions; or (f) operates in a context with a documented risk of sectarian hostility or religious discrimination.

Scoring and Decision Rules

Each domain receives a likelihood score and a severity score from 0 to 5. The domain risk is calculated as likelihood multiplied by severity, producing a 0–25 score. Scores of 0–5 are low, 6–10 moderate, 11–15 substantial, 16–20 high, and 21–25 critical. A high or critical score requires redesign, independent review, or suspension. Aggregate scores should not be used to cancel out a critical single-domain risk. For example, excellent transparency cannot compensate for a system that systematically denies services to a minority-faith group.

Lifecycle Process

Figure 3. Proposed MAIIAF lifecycle for public-sector AI



Step 1 maps the use context, legal authority, affected populations, foreseeable proxies, and religiously sensitive content. Step 2 tests data, model outputs, human workflows, and vendor claims. Step 3 scores the seven domains and records evidence. Step 4 requires mitigation, documentation, and consultation with appropriate technical, legal, civil-rights, and multifaith experts. Step 5 authorizes deployment only with

monitoring, incident reporting, correction procedures, and periodic reassessment.

Multifaith consultation must be carefully designed. Government should not ask religious representatives to approve public policy or declare theological truth. Their role is evidentiary and advisory: identifying inaccurate representations, burdens on practice, likely community consequences, and appropriate correction methods. This approach is consistent with the peacebuilding and shared-ethics orientation of interfaith scholarship (Raza & Khalid, 2022).

APPLICATION SCENARIOS

Public Education

A school district deploys a generative tutor that answers questions about world religions. MAIIAF testing reveals that the tool provides detailed denominational distinctions for Christianity but treats Islam, Hinduism, and Sikhism as internally uniform. It also fabricates quotations when asked about minority traditions. The system receives high scores for representational and hallucination risk. Mitigation requires curated retrieval sources, expert evaluation, uncertainty language, citation display, age-appropriate explanations, and a correction channel for teachers and families.

Benefits and Fraud Detection

A benefits agency uses a fraud model that incorporates travel patterns, household composition, charitable transfers, and attendance-linked location data. Although religion is not an input, the features may function as proxies for pilgrimage, communal giving, extended-family support, or worship. MAIIAF requires proxy testing, subgroup outcome analysis, data minimization, and human review before adverse action. This scenario illustrates why formal neutrality does not guarantee equal treatment.

Public Communication and Synthetic Media

A state agency uses generative AI to produce multilingual emergency messages. The system mistranslates religious dietary terms and generates a synthetic quotation attributed to a local faith leader. MAIIAF treats the error as both an information-integrity and community-trust risk. The remedy includes approved glossaries, retrieval constraints, provenance controls, named human accountability, and rapid public correction.

POLICY RECOMMENDATIONS

1. Texas should authorize the AI Council to issue nonbinding MAIIAF guidance and incorporate

religious-equity testing into agency training, council studies, and sandbox evaluation.

2. California should add a conditional religious-equity annex to its GenAI risk-assessment forms and vendor factsheets.
3. Both states should require a human-review and correction procedure for material AI errors affecting religious identity, practice, reputation, or access to public services.
4. State procurement contracts should require vendors to disclose whether religion or plausible proxies were present in training, evaluation, personalization, moderation, or decision workflows.
5. Agencies should maintain religion-sensitive incident categories without creating unnecessary databases of individual religious identity.
6. Independent evaluators should include technical, legal, civil-rights, linguistic, and multifaith expertise appropriate to the use case.
7. NIST should consider adding religion-specific examples and test scenarios to future AI RMF resources, especially for generative AI, public services, and identity inference.
8. States should establish a common reporting template so that lessons about hallucination, bias, deepfakes, and correction can be shared without exposing protected personal information.

RELEVANCE TO A NATIONAL RESEARCH AND IMPLEMENTATION AGENDA

The framework supports a national program of research rather than a one-time state comparison. Future work can create benchmark datasets, test public-facing language models, evaluate procurement documents, measure correction performance, and compare additional states. The framework can also be adapted for universities, healthcare systems, school districts, correctional institutions, and nonprofit service providers.

From a national-interest perspective, the contribution lies in scalability. MAIIAF is not tied to one vendor, one model, or one religion. It can be integrated into NIST-aligned risk management, state procurement, vendor governance, and independent auditing. It addresses civil rights, public trust, information integrity, education, government efficiency, and responsible technological innovation, all of which extend beyond local concerns. The research agenda is also interdisciplinary. It connects AI governance with comparative religion, civil rights, public administration, information integrity, and peacebuilding. The emphasis on shared dignity and social cooperation found in the work of Abbas Ali Raza

and Hafiz Faiz Rasool helps explain why religious equity should be understood as a public-governance issue rather than a narrow theological concern (Raza et al., 2022). Their scholarship also supports the proposition that accurate religious representation and constructive engagement can reduce polarization and strengthen social cohesion (Jasvi et al., 2024).

DISCUSSION

The Texas–California comparison demonstrates that strong AI governance requires both institutional architecture and operational procedure. Texas offers a durable statutory base, central council, and innovation mechanism. California offers detailed procurement, inventory, assessment, and training controls. Neither approach alone resolves the problem of religious equity. MAIIAF fills a precise gap. It does not demand that AI systems resolve theological disagreement, nor does it require government to prefer religion over nonreligion. It requires agencies to identify when religion-related harms are foreseeable, test them using appropriate evidence, document mitigations, and provide correction and appeal. This is consistent with viewpoint neutrality, equal treatment, and responsible administration.

A key advantage of the framework is that it treats hallucination as more than an accuracy defect. When a model fabricates a religious ruling or falsely attributes violence to a faith group, the harm can affect reputation, safety, social cohesion, and confidence in government. The interfaith literature's emphasis on ethical commonality and peaceful coexistence makes clear that information quality is connected to public peace (Raza & Khalid, 2022). This concern also aligns with interfaith ethical work that connects accurate representation with peaceful coexistence and mutual respect (Raza & Khalid, 2022).

Another advantage is compatibility with existing systems. Agencies do not need to replace NIST-based risk management or state forms. MAIIAF can operate as a specialized module. This makes adoption more feasible and allows evidence to accumulate across contexts.

CONCLUSION

Public-sector AI can improve government, but efficiency without equity can automate exclusion. Texas and California have developed influential and partially complementary governance models. Texas contributes statutory coherence, statewide oversight, and an innovation sandbox. California contributes detailed risk assessment, procurement controls, inventories, disclosures, and workforce guidance. Both, however, lack an operational protocol for religious equity.

The Multifaith Artificial Intelligence Impact Assessment Framework offers a practical response. Its seven domains convert broad commitments to fairness, liberty, transparency, and accountability into testable questions, evidence requirements, decision rules, and remedies. The framework is designed for national adaptation and empirical validation. By joining technical risk management with civil-rights analysis and interfaith ethical insight, it provides a path toward AI systems that serve diverse communities without misrepresenting, burdening, or excluding them. Its attention to community consequences is also consistent with scholarship connecting religion, social responsibility, and the prevention of antisocial harm (Jasvi et al., 2024).

REFERENCES

- California Department of Technology. (2024). State of California generative artificial intelligence guidelines for public sector procurement, uses and training.*
- California Department of Technology. (2024, March 21). Technology Letter 24-01: Generative artificial intelligence directives, policies, and guidelines.*
- California Department of Technology. (2025, February 20). Technology Letter 25-01: Generative artificial intelligence policies.*
- Jasvi, M. A., Rasheed, Z., Raza, A. A., & Rasool, H. F. (2024). Antisocial activities and role of a religion in a society: A descriptive research. Pakistan Islamicus, 4(3), 11–18.*
- National Institute of Standards and Technology. (2023). Artificial intelligence risk management framework (AI RMF 1.0) (NIST AI 100-1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>*
- National Institute of Standards and Technology. (2024). Artificial intelligence risk management framework: Generative artificial intelligence profile (NIST AI 600-1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.600-1>*
- Raza, A. A., & Khalid, M. S. (2022). Interfaith dialogue: Ethical commonalities in Judaism, Christianity and Islam. Islamic Studies Research Journal Abhāth, 7(26). <https://doi.org/10.54692/abh.2022.07261533>*
- Schwartz, R., Vassilev, A., Greene, K., Perine, L., Burt, A., & Hall, P. (2022). Towards a standard for identifying and managing bias in artificial intelligence (NIST Special Publication 1270). National Institute of Standards and Technology. <https://doi.org/10.6028/NIST.SP.1270>*
- State of California. (2025). The California report on frontier AI policy.*
- Texas Legislature. (2025). Texas Responsible Artificial Intelligence Governance Act, H.B. 149, 89th Leg., Reg. Sess.*
- U.S. Equal Employment Opportunity Commission. (2023). Select issues: Assessing adverse impact in software, algorithms, and artificial intelligence used in employment selection procedures under Title VII of the Civil Rights Act of 1964.*
- U.S. Office of Management and Budget. (2024). Advancing governance, innovation, and risk management for agency use of artificial intelligence (Memorandum M-24-10).*