

PUBLIC SUBMISSION

As of: 9/16/21 9:52 AM Received: September 07, 2021 Status: Pending_Post Tracking No. kta-deow-byuc Comments Due: September 15, 2021 Submission Type: API
--

Docket: NIST-2021-0004
Artificial Intelligence Risk Management Framework

Comment On: NIST-2021-0004-0001
Artificial Intelligence Risk Management Framework

Document: NIST-2021-0004-DRAFT-0065
Comment on FR Doc # 2021-16176

Submitter Information

Name: Anonymous Anonymous

General Comment

Consider aligning with CAN/CIOSC 101:2019, Ethical design and use of automated decision systems.

In October 2019, the CIO Strategy Council published a National Standard of Canada for automated decision systems — the world's first consensus-based standard helping organizations design and implement responsible artificial intelligence (AI) solutions.

The CIO Strategy Council engaged over 100 thought leaders and experts from governments, industry, academia and civil society groups from coast-to-coast-to-coast in developing the National Standard of Canada, CAN/CIOSC 101:2019, Ethical design and use of automated decision systems.

CAN/CIOSC 101:2019 specifies the minimum requirements in protecting human values and incorporating ethics in the design and use of artificial intelligence (AI) using machine learning for automated decisions.

This national standard goes beyond a common set of aspirational principles. It provides a framework and process that can be both measured and tested for conformity, providing consumers with confidence in the technologies that are providing information, recommendations, or making decisions using AI and machine learning.

The national standard applies to all organizations including public and private companies, government entities, and not-for-profit organizations, and helps organizations address AI ethics principles, such as those described by the OECD:

AI should benefit people and the planet by driving inclusive growth, sustainable development and well-being.

AI systems should be designed in a way that respects the rule of law, human rights, democratic values and diversity, and they should include appropriate safeguards – for example, enabling human intervention

where necessary – to ensure a fair and just society.

There should be transparency and responsible disclosure around AI systems to ensure that people understand AI-based outcomes and can challenge them.

AI systems must function in a robust, secure and safe way throughout their life cycles and potential risks should be continually assessed and managed.

Organizations and individuals developing, deploying or operating AI systems should be held accountable for their proper functioning in line with the above principles.