

All comments will be made public as-is, with no edits or redactions. Please be careful to not include confidential business or personal information, otherwise sensitive or protected information, or any information you do not wish to be posted.

Comment Template for
Responses to NIST
Artificial Intelligence Risk
Management Framework

Submit comments by August 19, 2021:

General RFI Topics (Use as	Response #	Responding	Responder's	Paper	Response/Comment (Include rationale)	Suggested change
Responses to Specific Request						
1. The greatest challenges in improving how AI actors manage AI-related risks – where “manage” means identify, assess, prioritize, respond to, or communicate those risks;	1	Preamble	Jeremy McHugh		The lack of regulation in tech and especially in the AI space, allows AI service companies to only worry about mitigating enough AI risk to maintain a fair image for public relations. In order to manage risk, standards need to be developed to identify, measure, then make adjustments. The greatest challenge is capturing the societal level risks from the large scale deployment of AI algorithms for big tech companies. The cumulative risk over an extended period of time can turn into a massive problem for society. The risk where a 1 in 1,000 chance that an individual is radicalized by YouTube or TikTok videos, can aggregate into a massive problem for society. One component to assess AI risk could be to measure unintended time on a platform due to recommender systems.	
					AI alignment and measures of alignment could be a good starting point for measuring AI risk as described in one of our papers - https://arxiv.org/abs/2107.10939 .	
2. How organizations currently						
3. How organizations currently						

4. The extent to which AI risks are incorporated into different organizations' overarching enterprise risk management – including, but not limited to, the management of risks related to cybersecurity, privacy, and safety;		Preamble	Jeremy McHugh		In terms of safety, the risk unmeasured by tech companies that use dangerous algorithms to recommend content on social media platforms can lead to negative personal behaviors from the consistent reinforcement of content that was intended to increase user engagement. Unfortunately, one of the most efficient ways to increase user engagement is to recommend content that is morally outrageous. Another issue with content recommendation algorithms is that training from user engagement and personal information can lead to more privacy and safety concerns when personal information is collected from individuals at a large scale. If users could otherwise select their preferences from an independent party, then their personal information would be isolated, while curating more quality content to present the user.	
5. Standards, frameworks,						
6. How current regulatory or						
7. AI risk management						
8. How organizations take into						
9. The appropriateness of the						
10. Effective ways to structure						
11. How the Framework could						

12. The extent to which the Framework should include governance issues, including but not limited to make up of design and development teams, monitoring and evaluation, and grievance and redress.	12	Preamble	Jeremy McHugh		The framework should include governance that outlines how transparent models are and how much control users have over the content recommendation algorithms through a middleware service. There also needs to be a way to measure influence from the algorithms. This measurement could then help characterize the decisions that were made on the platform. Once there are tools and methods to measure the influence and choices for algorithms, then standards for these measurements can be set. An independent certifying party could monitor and evaluate AI services. The main caveat to any monitoring solution is having access to the data from big technology companies pre-deployment. It is hard to assess what could happen in a real world deployment until it happens, unless a governing body has early access for auditing. A market solution could be created that monitors AI systems to assess the inputs and compare them to the outputs in order to determine the appropriate settings that can be made to realign the system.	
---	----	----------	---------------	--	---	--