

Question ID	Question Title	Available options	OpenAI	xAI	Z.ai	Anthropic	Google Deepmind
Q23	<i>If you specified external pre-deployment safety evaluations in the previous section, were these performed before or after broad internal deployment? (Select one)</i>	- Before - External safety tests were completed before broad internal deployment. - Partial - All external evaluations on situational awareness, scheming, and cyber-offense were conducted before broad internal deployment. - After - External safety tests were completed after broad internal deployment. - Other (please explain briefly):	After - External safety tests were completed after broad internal deployment.	Before - External safety tests were completed before broad internal deployment.	Partial - All external evaluations on situational awareness, scheming, and cyber-offense were conducted before broad internal deployment.		
Q24	<i>What level of safety testing does your company require for broad internal deployment of frontier AI models? (Select one)</i>	- No formal risk management requirements for internal deployments - Formalized risk management for internal deployments with less stringent requirements than external deployment framework for the following risks/capabilities: situational awareness, scheming, AI R&D, cyber-offense. - Formalized risk management for internal deployments with the same requirements as external deployment framework for the following risks/capabilities: situational awareness, scheming, cyber-offense. - Company requires the same risk management effort for internal and external deployments. - Other (Please briefly describe):	As described in our public Preparedness Framework, we believe that models that have reached or are forecasted to reach Critical capability under our framework will require additional safeguards (safety and security controls) during development, regardless of whether or when they are externally deployed. We do not currently possess any models that have Critical levels of capability, and we expect to further update this Preparedness Framework before reaching such a level with any model.	No formal risk management requirements for internal deployments	Company requires the same risk management effort for internal and external deployments.		
Q25	<i>Does your company require any of the following safeguards for broad internal deployments of frontier AI models? (Select all that apply)</i>	- Inference time safety mitigations for misuse risks (including cyber & bio risks) - Restricting access to helpful-only models and only granting time-bound access to staff that apply with a legitimate research need - Logging all inputs and outputs from internal use and retaining them for at least 30 days - Not currently logging, but introduced an *official, written* plan to start doing so after models reach a specified capability threshold - Analyzing all internal model interactions for abnormal activity, including harmful use or unexpected attempts by AI systems to take real-world actions - Live monitoring and automated editing/ resampling of suspicious outputs - None of the above - Other (please describe briefly):	See answer to Q24, above.	Inference time safety mitigations for misuse risks (including cyber & bio risks) Restricting access to helpful-only models and only granting time-bound access to staff that apply with a legitimate research need	Restricting access to helpful-only models and only granting time-bound access to staff that apply with a legitimate research need Logging all inputs and outputs from internal use and retaining them for at least 30 days	We have nuanced rigorous approach to safeguards- each of these depends on product surface, classifier and harm type, and use case.	