



# xAI Frontier Artificial Intelligence Framework

Effective Date: 30 June 2026

---

## 1. Introduction

xAI LLC (“**xAI**”)’s mission is to understand the universe through maximally truth-seeking and helpful AI, built with rigorous safety and security at every stage of development. This Frontier AI Framework (“**FAIF**”) outlines xAI’s approach to policies for mitigation of significant risks associated with the development, deployment, and release of xAI’s frontier AI models, such as Grok.

Managing the risks related to advanced AI models presents unique challenges. Given the large and continuously growing range of applications where AI models may be deployed, it is difficult to comprehensively anticipate and model all of the general public’s potential applications and interactions for an AI model.

xAI particularly focuses on the risks of malicious use and safeguard circumvention (e.g., jailbreaks), which cover many different specific risk scenarios. In particular, our assessment of risk of malicious use includes risks of models assisting in the development, preparation, and/or execution of a chemical, biological, radiological, or nuclear threat (“**CBRN Risks**”) and risks of models being used by malicious actors in the development, preparation, and/or execution of cyberattacks, including on systems of infrastructural or national security importance (“**Offensive Cybersecurity Risks**”). Our assessment of risk of loss of control includes risks from humans losing the ability to reliably direct, modify, or shut down a model (“**Loss of Control Risks**”). We also assess risks of models with high manipulative capabilities potentially being misused in ways that could reasonably result in large scale harm (“**Harmful Manipulation Risks**”)¹. Within our risk management framework, we may identify and assess other risks.

Absent mitigation, high-risk scenarios become more likely as model capabilities increase. For example, an increase in offensive cyber capabilities heightens the risk of use of a model by adversarial parties in a cyberattack, but does not significantly alter the risk of the model enabling the production of biohazardous agents. Our safety evaluation and mitigation strategy focuses on individual model behaviors, which we categorize into three buckets: abuse potential (e.g., vulnerability to jailbreaks), concerning propensities (e.g., a propensity for deceiving the user), and dual-use capabilities (e.g., offensive cyber capabilities). In this FAIF, we characterize our understanding of different risk scenarios and the relevant behaviors.

---

¹ Terminology used in the The Safety and Security Chapter of the General-Purpose AI Code of Practice developed under the EU AI Act.



xAI references standards such as NIST's AI Risk Management Framework and ISO/IEC 42001 for AI management systems. We evaluate these during annual reviews and integrate them into benchmarks and safeguards. This FAIF represents our current understanding and approach of how severe frontier AI risks may be anticipated and mitigated. We may change our approach over time as we gain further experience and insights on the projected capabilities of future frontier models. We will review this Framework periodically, and expect significant changes as the capabilities of our models expand.

## 2. Risk Management Framework

xAI applies its risk management process, as appropriate, throughout the entire model lifecycle: during the various phases of model training, on various checkpoints or versions of a model, both prior to and after deployment (including through post-market monitoring).

xAI will conduct a full systemic risk assessment and mitigation process of our frontier models at least once a year. xAI may also conduct smaller-scale model evaluations at appropriate trigger points to determine whether additional systemic risk mitigations may be required or if an evaluation report update is required. These trigger points may include (1) the release of an updated model, (2) in the event of a serious incident, (3) if the model's use or integrations into xAI's systems materially increase risk, or (4) where xAI has reason to believe that the basis for considering the model's systemic risks acceptable has materially changed.

Results of our systemic risk assessment and mitigation processes and measures will be documented.

### 2.1. Risk Identification

We systematically identify failure modes that could originate from our models, characterize those modes, and determine which present significant or unmitigated risk.

We evaluate a broad threat surface as part of our ongoing model development work, incorporating model capabilities, propensities, and affordances, internal engineering expertise and relevant external research, and considerations of deployments.

Early analysis<sup>2</sup> has established **four** primary risk domains where significant or severe outcomes are most likely to emerge in current and future systems: **CBRN Risks, Offensive Cybersecurity Risks, Loss of Control Risks** and **Harmful Manipulation Risks**. These risk domains describe principal pathways through which severe or systemic harm may arise. Depending on the relevant risk scenario, such harms may affect public health, safety, public security, fundamental rights, or society as a whole. A single risk scenario may affect more than one of these interests, and the risk domains may overlap. Within our risk management

---

<sup>2</sup> Following the Safety and Security Chapter of the General-Purpose AI Code of Practice developed under the EU AI Act.



framework, we may identify and assess other risks or risk pathways where relevant to a particular model.

As xAI advances frontier model development, we continuously evaluate whether emerging risk domains meet our significance and severity thresholds, and update our control architectures accordingly. As part of xAI's broader development of frontier AI models, we continue to assess whether there are other risk domains where significant or severe risks may arise. This process involves (a) compiling a taxonomy of risks involved in the model's training or deployment, (b) analysing relevant characteristics of these risks, (c) identifying systemic risks stemming from the model and (d) developing systemic risk scenarios and mitigations for each identified systemic risk.

xAI has developed systemic risk scenarios that enumerate the causal factors, potential harms, and mitigations of the most significant of these risks.

#### (a) Addressing CBRN Risks

xAI aims to reduce the risk that the use of its models might contribute to a malicious actor potentially seriously injuring people, property, or national security interests, including reducing such risks by enacting measures to prevent use for the development or proliferation of weapons of mass destruction and large-scale violence. Without any safeguards, we recognize that advanced AI models could lower the barrier to entry for malicious actors seeking to develop CBRN capabilities, and could assist in the automation of knowledge compilation to swiftly overcome bottlenecks to weapons development, amplifying the expected risk posed by such weapons of mass destruction. Our most basic safeguard against malicious use is to train and instruct our publicly deployed models to decline requests showing clear intent to engage in criminal or dual-use activity.

Under this FAIF, xAI's models apply heightened safeguards if they receive user prompts that demonstrate intent to engage in the development, use, , or proliferation of weapons of mass destruction, including CBRN agents. For example, xAI's models apply heightened safeguards if they receive a request to act as an agent or tool to enable or inform the synthesis of common chemical warfare agents, similarly refusing in the case of other CBRN agents or where the intent of the user is identifiable as aiming to surface information enabling the creation or use of such agents.

Even as we improve our model's ability to scrutinize user behavior and identify malicious actors, it remains imperative that xAI models apply these safeguards to user interactions. To this end, we continually evaluate and improve robustness to adversarial attacks that seek to remove xAI model safeguards (e.g., jailbreak attacks), or hijack and redirect Grok-powered applications toward nefarious purposes (e.g., prompt injection attacks).

To transparently measure our models' safety properties, xAI may utilize public and internal benchmarks to measure our model's dual-use capability and resistance to



facilitating the use, development, or proliferation of weapons of mass destruction (including CBRN agents).

xAI regularly evaluates the adequacy and reliability of such benchmarks, including by comparing them against other benchmarks that we could potentially utilize, to determine and apply effective benchmarks available at the time of evaluation.

(b) Addressing Risks of Loss of Control

While it is difficult to directly predict the occurrence probabilities and risk profiles of particular **Loss of Control** scenarios, xAI makes significant efforts to identify and mitigate the most likely of these risks. xAI's practice on loss of control centers around scalable model oversight, training against high-risk model behaviors such as deception and sycophancy, and robust evaluation and red-teaming of model-driven agents in controlled sandboxes.

xAI utilizes a set of benchmarks to evaluate its models for concerning properties relevant to loss of control risks, including behavioral evaluations, simulated deployments in controlled sandboxes, and post-deployment monitoring.

xAI regularly evaluates the adequacy and reliability of such benchmarks and controls, including by comparing them against other benchmarks that we could potentially utilize.

(c) Addressing Risks of Malicious Offensive Cybersecurity Usage

A salient risk of highly capable models is their potential for misuse for cyber attacks by independent malicious actors and state-level adversaries. Frontier models embodied in agentic harnesses are already capable of highly advanced software engineering, systems analysis, and multi-step problem solving, sufficient to provide meaningful uplift in reconnaissance, exploit development, operational planning, and offensive cyber campaign execution. This presents reduced barriers for less sophisticated actors and the potential for accelerating operations for well-resourced ones, including advanced persistent threat groups (APTs).

xAI treats AI-accelerated malicious cyber offense as a core risk domain. We focus on scenarios in which models could, in concert with human actors, cause significant disruption to computing systems and infrastructure. We mitigate this risk through strict benchmarking of our models in agentic cybersecurity contexts, as well as testing robustness against public-domain and non-public jailbreaking methods. We additionally train our models to refuse requests whose intentions are to perform cyber offense, including training to avoid jailbreaks, and further secure our systems through moderation filters and monitoring of production deployments.



xAI utilizes a set of benchmarks to evaluate its models for cyber capabilities, cyber refusals and detection of malicious intent. xAI regularly evaluates the adequacy and reliability of such benchmarks, including by comparing them against other benchmarks that we could potentially utilize.

## 2.2. Systemic risk analysis

For the purpose of facilitating xAI's systemic risk acceptance determination, xAI analyses each identified risks based on the following:

- (1) **model-independent information:** xAI collects information relevant to the specific risk assessment, using methods such as web searches, literature reviews, market analysis, reviews of training data, historical incident data and incident databases, forecasting of general trends, expert interviews and/or panels and/or lay interviews, surveys, community consultations, or other participatory research methods investigating, e.g. the effects of models on natural persons, including vulnerable groups.
- (2) **conducting model evaluations:** xAI conducts state-of-the-art model evaluations relevant to the systemic risk to assess the model's capabilities, propensities, affordances, and/or effects, which may include Q&A sets, task-based evaluations, benchmarks, red-teaming and other methods of adversarial testing, human uplift studies, model organisms, simulations, and/or proxy evaluations for classified materials.
- (3) **modelling the systemic risk:** xAI conducts systemic risk modelling for relevant systemic risks by using state-of-the-art risk modelling methods, building on the systemic risk scenarios in the risk identification process and taking into account the information gathered.
- (4) **estimating the systemic risk:** xAI estimates the the probability and severity of harm for the systemic risk using state-of-the-art risk estimation methods and taking into account the information gathered throughout the risk management process.
- (5) **conducting post-market monitoring:** xAI conducts appropriate post-market monitoring to gather information relevant to assess whether the risk could be determined to not be acceptable and to inform whether a model update is necessary. Post-market monitoring may include collecting end-user feedback; providing (anonymous) reporting channels; providing (serious) incident reporting forms; monitoring software repositories, known malware, public forums, and/or social media for patterns of use.

## 2.3. Systemic risk acceptance determination

xAI applies a systemic risk acceptance criteria to each identified risk, incorporating a margin of security, to determine whether each identified systemic risk and the overall systemic risk are acceptable and, therefore, to proceed with the development, the making available on the



market, and/or the use of the model. These evaluations are performed precedent to the wider deployment of our models, and are a precondition for release of our models for public use.

xAI uses rigorous evaluations to determine whether our models present systemic risks that remain within acceptable levels and to assess residual risk. Our consideration of residual risk takes account of the scale and probability of harm, and is evaluated in light of the sufficiency of implemented safeguards proportionate to the level of risk. In assessing whether to accept residual risk, we review the risk tiers for each systemic risk category and incorporate appropriate safety margins.

When analyzing whether a threshold has been reached, we also take into account additional evidence to validate the findings of evaluations, such as human expert red-teaming. These results help determine whether the threshold has been reached. The final decision also reflects an overall judgment based on all the available evidence, including how robust and reliable the evaluation methods were.

xAI will only proceed with the development, the making available on the market, and/or the use of the model, if the systemic risks stemming from the model are determined to be acceptable. In all scenarios, xAI will take appropriate measures to ensure the systemic risks stemming from the model are and will remain acceptable prior to proceeding. In particular, xAI may:

- (1) to mitigate risks, employ tiered availability of the functionality and features of its models. For instance, the full functionality of our models may be available to only a limited set of trusted parties, partners, and government agencies;
- (2) include additional controls on functionality and features depending on the type of end user. For instance, features that we make available to consumers using mobile apps may be different than the features made available to sophisticated businesses
- (3) conduct another round of systemic risk identification, analysis and systemic risk acceptance determination.

## 2.3. Safety mitigations

xAI implements appropriate safety mitigations to address the systemic risks stemming from our models, taking into account the model's release and distribution strategy, focusing on individual model behaviors, which we categorize into three buckets: abuse potential (e.g., vulnerability to jailbreaks), concerning propensities (e.g., a propensity for deceiving the user), and dual-use capabilities (e.g., offensive cyber capabilities).

**Approach to Mitigating Risks of Malicious Use:** Alongside comprehensive evaluations measuring dual-use capabilities, our mitigation strategy for malicious use risks is to identify critical steps in major risk scenarios and implement redundant layers of safeguards in our models to inhibit user progress in advancing through such steps. xAI works to identify such inhibiting steps, commonly referred to as bottlenecks, and implement commensurate safeguards to mitigate a model's ability to assist in accelerating a malicious actor's progress through them.



Model safeguards leverage a broad variety of techniques, including standard software systems and state-of-the-art AI capabilities, to detect and block potential abuses.

xAI's objective is for our models to comply with their training, robustly resisting attempted manipulation and adversarial attacks. In addition to the incidental alignment resulting from post-training (our models naturally tend to refuse malicious requests even without any safety-specific training data), we are developing training methods and will continue to train our models to robustly resist complying with requests to provide assistance with highly injurious malicious use cases.

Driving towards our safety objectives, we continue to design and deploy the following safeguards into our models, as appropriate:

- **Safety training:** Training our models to recognize and decline harmful requests.
- **System prompts:** Providing high-priority instructions to our models to enforce our basic refusal policy.
- **Filters:** Applying classifiers to verify safety when a model is queried regarding topics of CRBN risks

Because xAI is committed to continual improvement, we will continue to evaluate our approach to enhancing safety. Thus, xAI may change its approach from that listed above in order to make additional improvements.

**Approach to Mitigating Risks of Loss of Control:** Exact scenarios of loss of control risks are speculative and difficult to precisely specify. Many such scenarios, for example, speculation that a superintelligent AI system hypothetically might escape the control of its developers and wreak havoc on the public, assume dual-use capabilities such as offensive cybersecurity capabilities (e.g., to surreptitiously replicate across servers or prevent shutdown) that we also track as part of managing malicious use risks. Additionally, we conduct careful measurement of concerning model propensities that hypothetically might exacerbate loss of control risks, such as the propensity for deception or the propensity for sycophancy. We continue to work towards developing naturalistic evaluation environments that would enable us to assess more realistic, real-world behaviors.

## 2.4. Security mitigations

xAI will maintain a documented Security Goal that identifies the threat actors against which mitigations are designed to protect against, including sophisticated non-state actors, insider threats, state-sponsored actors and other foreseeable actors. The Security Goal will be reviewed at least every year.

xAI has implemented appropriate information security standards, adopted based on the NIST 800-171 Rev.3 framework and supported by SOC 2 Type II evaluations. Mitigations include



encryption of model weights, role-based access controls, and real-time monitoring to prevent unauthorized transfer. To prevent the unauthorized proliferation of advanced AI systems, we also implement security measures against the large-scale extraction and distillation of reasoning traces, which have been shown to be highly effective in quickly reproducing advanced capabilities while expending far fewer computational resources than the original AI system.<sup>3</sup>

Security for our AI services and processing systems is structured around comprehensive technical and organizational measures designed to protect the security, confidentiality, integrity, and availability of personal data and user content, along with encryption, access restrictions, and personnel safeguards. Access is tightly managed through least-privilege principles, role-based authorization, and regular reviews.

### 3. Incident reporting

xAI has measures in place to detect, investigate and remediate, as appropriate, serious AI safety incidents.

To foster accountability, we integrate the approach of designating risk owners, including assigning responsibility for proactively mitigating identified risks and monitoring for critical incidents or imminent threats, which may be identified through various channels, including:

- Red-teaming and internal testing;
- Real-time monitoring, telemetry, and alerting of threshold breaches via internal tooling;
- Monitoring and alerting of public comments from the X platform;
- Employee escalation;
- Feedback from end-users, downstream modifiers, downstream providers and other third parties through xAI's appropriate reporting channels.

When critical incidents or imminent threats are identified, xAI's will investigate the causes and effects of the incidents and, if necessary, take action to mitigate and contain incidents, and implement long-term solutions:

1. If we determine that xAI systems are actively being used in such an event, we may take steps to isolate and revoke access to user accounts involved in the event.
2. If we determine that allowing a system to continue running would materially and unjustifiably increase the likelihood of a risk, we may temporarily fully shut down the relevant system until we have developed a more targeted response.
3. We may perform a post-mortem of the event after it has been resolved, focusing on any areas where changes to systemic factors could have averted such an incident. We may use the post-mortem to inform development and implementation of necessary changes to our risk management practices.

---

<sup>3</sup> [DeepSeek-R1: Incentivizing Reasoning Capability in LLMs via Reinforcement Learning](#)





4. When reportable under applicable laws and regulations, xAI will provide the relevant authorities with a copy of the incident report within the required deadlines

## **4. Governance**

xAI maintains internal governance structures and practices designed to meet the requirements of applicable laws and ensure the implementation of the processes in this FAIF. xAI's internal governance practices include managing risks across the lifecycle of our models and ongoing legal and compliance reviews to ensure that risk management functions adhere to this FAIF.