

Frontier Safety Framework

Version 3.0

Published: September 22, 2025

Frontier Safety Framework

Version 3.1

Published: April 17, 2026

Overview

The Frontier Safety Framework is a set of protocols that aims to address severe risks that may arise from the high-impact capabilities of frontier AI models. It complements Google's suite of AI responsibility and safety practices, and enables AI innovation and deployment consistent with our [AI Principles](#).

The Framework is informed by the broader conversation on Frontier AI Safety and Security Frameworks.¹ The core components of such Frameworks are to:

- Identify capability levels at which frontier AI models, without additional mitigations, could pose severe risk.
- Implement protocols to detect the attainment of such capability levels throughout the model lifecycle.
- Prepare and articulate proactive mitigation plans to ensure severe risks are adequately mitigated when such capability levels are attained.
- Where required or appropriate, involve external parties to help inform and guide the approach.

In the Framework, we specify protocols for the detection of capability levels at which frontier AI models may pose severe risks (which we call "Critical Capability Levels (CCLs)"), and articulate mitigation approaches to address such risks. The Framework addresses misuse risk,² risks from machine learning research and development (ML R&D), and misalignment risk.³ For each type of risk, we define a set of CCLs and a mitigation approach for them. Risk assessment will necessarily involve evaluating cross-cutting capabilities such as agency, tool use, reasoning, and scientific understanding.

The safety and security of frontier AI models is a global public good. The protocols here represent our current understanding and approach of how severe frontier AI risks may be anticipated and addressed. Importantly, there are certain mitigations whose social value is significantly reduced if not broadly applied to frontier AI models reaching critical capabilities. These mitigations are most effective when adopted by industry as a whole: our adoption of them would result in effective risk mitigation for society only if all relevant organisations provide similar levels of protection.

The Framework is based on early and evolving research. We may change our approach over time as we gain experience and insights on the projected capabilities of future frontier models. We will review the Framework periodically and we expect it to evolve substantially as our understanding of the risks and benefits of frontier models improves.

¹ See <https://www.gov.uk/government/publications/emerging-processes-for-frontier-ai-safety>, <https://metr.org/faisc>, <https://www.anthropic.com/news/anthropics-responsible-scaling-policy>, <https://openai.com/index/updating-our-preparedness-framework/>, <https://www.frontiermodelforum.org/publications/#technical-reports>.

² As in, in the context of the Framework, risks of threat actors using critical capabilities of deployed or exfiltrated models to cause harm.

³ Misalignment can pose a number of risks. In the context of the Framework, we address specific scenarios where general-purpose AI agents are potentially misaligned and can become difficult to control, thereby posing a risk of severe harm.

Overview

The Frontier Safety Framework is a set of protocols that aims to address severe risks that may arise from the high-impact capabilities of frontier AI models. It complements Google's suite of AI responsibility and safety practices, and enables Google's AI innovation and deployment consistent with our [AI Principles](#).

The Framework is informed by the broader conversation on Frontier AI Safety and Security Frameworks.¹ The core components of such Frameworks are to:

- Identify capability levels at which frontier AI models, without additional mitigations, could pose severe risk.
- Implement protocols to detect the attainment of such capability levels throughout the model lifecycle.
- Prepare and articulate proactive mitigation plans to ensure severe risks are adequately mitigated when such capability levels are attained.
- Where required or appropriate, involve external parties to help inform and guide the approach.

In the Framework, we specify protocols for the detection of capability levels at which frontier AI models may pose significant or severe risks (which we call "Tracked Capability Levels (TCLs)" and "Critical Capability Levels (CCLs)" respectively), and articulate mitigation approaches to address such risks. The Framework addresses misuse risk² as well as machine learning research and development (ML R&D) and misalignment risk.³ For each type of risk, we define a set of CCLs (and TCLs where relevant) and a mitigation approach for them. Risk assessment will necessarily involve evaluating cross-cutting capabilities such as agency, tool use, reasoning, and scientific understanding. Please refer to the [Glossary](#) at the end of this document for definitions used in the Framework.

The safety and security of frontier AI models is a global public good. The protocols here represent our current understanding and approach of how severe frontier AI risks may be anticipated and addressed. Importantly, there are certain mitigations whose social value is significantly reduced if not broadly applied to frontier AI models reaching critical capabilities. These mitigations are most effective when adopted by industry as a whole: our adoption of them would result in effective risk mitigation for society only if all relevant organizations provide similar levels of protection.

The Framework is based on early and evolving research. We may change our approach over time as we gain experience and insights on the projected capabilities of future frontier models. We will review the Framework periodically and we expect it to evolve substantially as our understanding of the risks and benefits of frontier models improves.

¹ See <https://www.gov.uk/government/publications/emerging-processes-for-frontier-ai-safety>, <https://metr.org/faisc>, <https://www.anthropic.com/news/anthropics-responsible-scaling-policy>, <https://openai.com/index/updating-our-preparedness-framework/>, <https://www.frontiermodelforum.org/publications/#technical-reports>.

² As in, in the context of the Framework, risks of threat actors using critical capabilities of deployed or exfiltrated models to cause harm.

³ In the context of the Framework, we address specific scenarios where model misalignment may create the risk of significant or severe harm, and scenarios where ML R&D capabilities may heighten misalignment, misuse, and structural risks.

Table of Contents

Section 1: Framework	4
1.1 Scope	4
1.2 Critical Capability Levels	4
1.3 Outline of Our Risk Assessment Process	4
1.4 Response Plans and Mitigations	6
1.5 Evaluating Mitigations	6
1.6 Summary of Risk Acceptance Criteria	6
Section 2: Misuse	8
2.1 Mitigation Approach	8
2.2 Misuse Critical Capability Levels	9
Section 3: Machine Learning R&D	12
3.1 Mitigation Approach	12
3.2 ML R&D Critical Capability Levels	13
Section 4: Misalignment (Exploratory Approach)	15
Section 5: Updates and Disclosures	16
5.1 Updates	16
5.2 Disclosures	16
5.3 Past Updates and Changes	16

Table of Contents

Section 1: Framework	4
1.1 Scope	4
1.2 Critical Capability Levels and Tracked Capability Levels	4
1.3 Risk Management Process	5
1.3.1 Risk Identification	5
1.3.2 Inherent Risk Assessment	5
1.3.3 Risk Mitigation	6
1.3.4 Residual Risk Assessment	6
1.3.5 Risk Acceptance Determination	7
Section 2: Misuse	9
2.1 Mitigation Approach	9
2.1.1 Security Mitigations	9
2.1.2 Deployment Mitigations	10
2.2 Misuse Capability Levels	11
2.2.1 Chemical, Biological, Radiological or Nuclear	11
2.2.2 Cyber	12
2.2.3 Harmful Manipulation	12
Section 3: Machine Learning R&D and Misalignment	13
3.1 Mitigation Approach	13
3.1.1 Security Mitigations	13
3.1.2 Deployment Mitigations	13
3.2 ML R&D and Misalignment Capability Levels	14
3.2.1 Stealth and Situational Awareness Tracked Capability Level	14
3.2.2 ML R&D Critical Capability Levels	14
Section 4: Governance and Accountability	16
4.1 Governance structure	16
Section 5: Updates and Disclosures	17
5.1 Updates	17
5.2 Disclosures	17
5.3 Past Updates and Changes	17
Glossary	18

Section 1: Framework

This section describes the central components of the Frontier Safety Framework. These protocols represent our current understanding of and approach for how severe frontier AI risks may be anticipated and addressed.

1.1 Scope

The Frontier Safety Framework focuses on possible severe risks stemming from high-impact capabilities of frontier AI models. The Framework complements Google's suite of AI responsibility and safety practices, which address other risks in addition to the severe risks in scope of the Framework in accordance with our [AI Principles](#). The approaches and mitigations outlined in the Framework are not exclusive to models where we believe a severe risk could arise and are part of Google's comprehensive AI responsibility and safety practices.

1.2 Critical Capability Levels

The Framework is built around capability thresholds called "Critical Capability Levels (CCLs)." These are capability levels at which, absent mitigation measures, frontier AI models or systems may pose heightened risk of severe harm. CCLs are determined by identifying and analyzing the main foreseeable paths through which a model could result in severe harm: we then define the CCLs as the minimal set of capabilities a model must possess to do so.

We describe three sets of CCLs: misuse CCLs, machine learning R&D CCLs, and misalignment CCLs.

For **misuse risk**, we define CCLs in the following risk domains where the misuse of model capabilities may result in severe harm:

- **CBRN:** Risks of models assisting in the development, preparation, and/or execution of a chemical, biological, radiological, or nuclear ("CBRN") threat.
- **Cyber:** Risks of models assisting in the development, preparation, and/or execution of a cyber attack.
- **Harmful Manipulation:** Risks of models with high manipulative capabilities potentially being misused in ways that could reasonably result in large scale harm.

For **machine learning R&D risk**, we define CCLs that identify when ML R&D capabilities in our models may, if not properly managed, reduce society's overall ability to manage AI risks. Such capabilities may serve as a substantial cross-cutting risk factor for other pathways to severe harm.

For **misalignment risk**, we outline an exploratory approach that focuses on detecting when models might develop a baseline **instrumental reasoning** ability at which they have the potential to undermine human control, assuming no additional mitigations were applied.

Most CCLs define one important component of our **risk acceptance criteria**. Because the CCLs for misalignment risk are exploratory and intended for illustration only, we do not associate them with explicit risk acceptance criteria.

1.3 Outline of Our Risk Assessment Process

For each risk domain, we conduct aspects of our risk assessment at various moments throughout the model development process, both before and after deployment. We conduct a risk assessment for the first external deployment of a new frontier AI model. For subsequent versions of the model, we conduct

Section 1: Framework

This section describes the central components of the Frontier Safety Framework. These protocols represent our current understanding of and approach for how severe frontier AI risks may be anticipated and addressed.

1.1 Scope

The Frontier Safety Framework focuses on possible severe risks stemming from high-impact capabilities of frontier AI models. The Framework complements Google's suite of AI responsibility and safety practices, which address other risks in addition to the severe risks in scope of the Framework in accordance with our [AI Principles](#). The approaches and mitigations outlined in the Framework are not exclusive to models where we believe a severe risk could arise and are part of Google's comprehensive AI responsibility and safety practices.

1.2 Critical Capability Levels and Tracked Capability Levels

The Framework is built primarily around capability thresholds called "Critical Capability Levels (CCLs)." These are capability levels at which, absent mitigation measures, frontier AI models or systems may pose heightened risk of severe harm. CCLs are determined by identifying and analyzing the main foreseeable paths through which a model could result in severe harm: we then define the CCLs as the minimal set of capabilities a model must possess to do so. CCLs are one important component of our **risk acceptance determination**.

We identify CCLs for two kinds of risks: misuse risk and risks related to machine learning R&D and misalignment.

For **misuse risk**, we define CCLs in the following risk domains where the misuse of model capabilities may result in severe harm:

- **CBRN:** Risks of models assisting in the development, preparation, and/or execution of a chemical, biological, radiological, or nuclear ("CBRN") threat.
- **Cyber:** Risks of models assisting in the development, preparation, and/or execution of a cyber attack.
- **Harmful Manipulation:** Risks of models with high manipulative capabilities potentially being misused in ways that could reasonably result in large scale harm.

For **machine learning R&D and misalignment risks**, we define CCLs that identify when ML R&D capabilities in our models may, if not properly managed, reduce society's overall ability to manage AI risks. Such capabilities may serve as a substantial cross-cutting risk factor for several pathways (e.g. through misalignment, misuse or structural risks) to severe harm.

This update to the Framework introduces "**Tracked Capability Levels (TCLs)**." TCLs are meant to capture significant risks that may manifest at a lower capability threshold than our CCLs and we apply our mitigation and risk acceptance determination process proportionately. We identify TCLs for CBRN risk, as well as ML R&D and misalignment risks. We may include TCLs for additional risks in the future, as our threat modeling develops. Similarly to CCLs, early warning evaluations will be used to assess the proximity of a model to a TCL and analyze the risk posed, involving internal and external experts as needed.

a further risk assessment if the model has meaningful new capabilities or a material increase in performance, until the model is retired or we deploy a more capable model. The reason for this is because a material change in the model's capabilities may mean that the risk profile of the model has changed or the justification for why the risks stemming from the model are acceptable has been materially undermined.

To identify meaningful new capabilities or material increases in performance, we conduct model capability evaluations, including our automated benchmarks. These evaluations are primarily aimed at understanding the capabilities of the model and may be triggered, for example, upon the completion of a pre-training or post-training run, on various candidates of a model version. These evaluations include a broad range of areas, including general capability evaluations, model behavior, efficiency, coding capabilities, multilinguality, or reasoning. Data from these evaluations are collected and analyzed to give us an indication as to how the model is performing and whether a risk assessment is necessary.

At a high level, our risk assessment involves the following steps (which do not need to be repeated where a previous risk assessment is still appropriate):

- **Identification:** As explained above, we have identified risk domains where, based on early research, we have determined severe risks may be most likely to arise from future models: CBRN, cyber, harmful manipulation, and machine learning R&D.¹⁴ As part of our broader research into frontier AI models, we continue to assess whether there are other risk domains where severe risks may arise and will update our approach as appropriate. For each of the four identified domains, we have developed specific scenarios in which these risks could materialize.
- **Analysis:** Central to our model evaluations are "early warning evaluations," to assess the proximity of the model to a CCL. We define "alert thresholds" for these evaluations that are designed to flag when a CCL may be reached before a risk assessment is conducted again. In our evaluations, we seek to equip the model with appropriate scaffolding and other augmentations to make it more likely that we are also assessing the capabilities of systems that will likely be produced with the model. We may run early warning evaluations more frequently or adjust the alert threshold of our evaluations if the rate of progress suggests our safety buffer is no longer adequate. We conduct further analysis, including reviewing model independent information, external evaluations, and post-market monitoring as appropriate.
- **Acceptance determination and mitigations:** We then determine whether the model has met or will meet a CCL and, if so, whether we need to implement any further mitigations to reduce the risk to an acceptable level (see [below](#)).

Our approach to model evaluations and risk assessments described above means we can proactively monitor a model's capabilities throughout the entire lifecycle of the model and ensure that any severe risk is properly identified and mitigated. Where appropriate, we may engage relevant and appropriate external actors, including governments, to inform our responsible development and deployment practices.

Note on Machine Learning R&D CCLs: Risk assessment must take into account the fact that other actors may put significantly more effort into eliciting capabilities than we put into assessing risk, thus requiring conservatism in the form of evaluations. However, as a frontier AI company, we do not expect other groups to put significantly more effort into ML R&D than we do ourselves. As a result, to assess the ML R&D CCLs, we may use sources of information about our own progress at accelerating ML R&D to assess whether we are near or at the CCLs, in addition to evaluations of ML R&D capabilities. Similarly,

¹⁴ We exclude misalignment risk from this list of domains because of its exploratory nature.

1.3 Risk Management Process

We conduct our risk management process as appropriate, throughout the model development process, on various checkpoints or versions of a model, both before and after deployment.

1.3.1 Risk Identification

The first part of our risk management process is **risk identification**. We identify potential risks that could stem from our models and analyze their characteristics to determine which of the identified risks could be significant or severe risks. We consider a wide range of risks as part of our ongoing research, taking into account the characteristics, capabilities, propensities, and affordances of our models and other sources of information, such as our internal risk taxonomies, internal expertise and relevant external research. As explained above, we have identified risk domains where, based on early research, we have determined significant or severe risks may be most likely to arise from future models: CBRN, cyber, harmful manipulation, as well as machine learning R&D and misalignment. As part of our broader research into and development of frontier AI models, we continue to assess whether there are other risk domains where significant or severe risks may arise and will update our approach as appropriate. For each of the four identified domains, we have developed specific scenarios and T/CCLs in which these risks could materialize.

1.3.2 Inherent Risk Assessment

To understand whether a model may, without appropriate mitigations, contribute to the risks identified, we conduct a closer assessment of its capabilities against the T/CCLs. We conduct a **critical capability assessment** prior to the first external deployment¹⁵ of a new frontier AI model. For external deployment of subsequent versions of the model, we determine whether a substantial modification has been made and whether a further critical capability assessment is required if (1) the model has meaningful new capabilities or material increases in performance; and (2) we believe such material capability increase could materially undermine the justification for why the risks stemming from the model are acceptable.

To understand if a subsequent version of the model has meaningful new capabilities or material increases in performance relative to the last checkpoint subject to a critical capability assessment, we conduct **material capability change assessments** on various checkpoints upon the completion of a post-training run. These assessments draw on model characteristics, such as its provenance, size, modality, and performance on general capability benchmarks. To understand if such a change in capability in the subsequent version of the model could materially undermine the justification for why the risks stemming from the model are acceptable, we may also conduct light-weight versions of our "early warning evaluations" described below, including our automated benchmarks, to confirm a model remains below T/CCL. We may also apply material capability change assessments to models not on the frontier, to understand whether a critical capability assessment is required.

Central to our critical capability assessments are "early warning evaluations," which we use to test the specific threats and risk scenarios identified through our threat modeling, determine a model's capability, and assess the proximity of the model to a T/CCL. For CCLs, we define "alert thresholds," which may draw on evaluation results, expert assessments, and other sources of information; that are designed to flag when a CCL may be reached before a critical capability assessment is conducted again. In our evaluations, we seek to apply appropriate scaffolding, inference compute, and other augmentations to also assess the capabilities of systems that will likely be produced with the model based on the specific threats and risk scenarios we have identified for the T/CCL. We may run early warning evaluations more

¹⁵ A critical capability assessment may not be conducted for low-risk external deployments (e.g. to a small number of trusted testers) if the appropriate governance function determines the residual risk of such deployments to be acceptable even if the model has reached a T/CCL (including with a safety case where a CCL has been reached).

our alert threshold may be defined based on these sources of information, rather than on evaluation scores.

1.4 Response Plans and Mitigations

We apply safety and security mitigations throughout the lifecycle of our models, including as part of our training and model development phase and, where appropriate, before CCLs are reached as described in the process below.

When a model reaches an alert threshold for a CCL, we will assess the proximity of the model to the CCL and analyze the risk posed, involving internal and external experts as needed. This will inform the formulation and application of a response plan. Where model capabilities remain quite distant from a CCL, a response plan may involve the adoption of additional capability assessment processes to flag when heightened mitigations are required.

Central to most response plans will be the application of the mitigations described later in this document. We have two categories of mitigations: *security mitigations* (such as preventing the exfiltration of model weights), and *deployment mitigations* (such as safety fine-tuning and monitoring and response) intended to counter the misuse or misaligned expression of critical capabilities in deployments. Note that these mitigations reflect considerations from the perspective of addressing severe risks from powerful capabilities alone; due to this focused scope, other risk management and security considerations may result in more stringent mitigations applied to a model than specified by the Framework.

1.5 Evaluating Mitigations

We will use various processes to evaluate the effectiveness and limitations of mitigations:

- **Security mitigations:** security infrastructure at Google is subject to penetration testing and other kinds of assessments, and is continually improved based on these results.
- **Deployment mitigations:** we will use a combination of threat modeling, empirical testing, and other sources of information to assess the effectiveness and limitations of our deployment mitigations. These will form the basis of a safety case¹ for models reaching CCLs, that will be reviewed before deployment. See the deployment mitigations sections below for [misuse](#) and [machine learning R&D](#) for more.

1.6 Summary of Risk Acceptance Criteria

The Framework outlines a variety of risk acceptance criteria for different model risks and capabilities pertaining to severe risks. To summarize:

- A model for which risk assessment indicates no CCL is reached will be deemed to pose an acceptable level of severe risk for further development and deployment, because it should not possess the capabilities required to substantially contribute to severe risk scenarios.²
- A model for which the risk assessment indicates a misuse CCL has been reached will be deemed to pose an acceptable level of risk for further development or deployment, if, for example:

¹ A safety case is an assessable argument showing how severe risks associated with a model's CCLs have been reduced to an appropriate level. See also, for reference, <https://arxiv.org/abs/2505.01420>.

² Note that this does not mean that no mitigations will be implemented prior to the deployment of the model. We apply safety and security mitigations throughout the lifecycle of our models for a whole host of potential risks, including as part of our training and model development phase. This section is limited to those mitigations that will be required to bring a model to an acceptable level of risk relative just to the specific CCL, which by definition pose the greatest risks of severe harm.

frequently or adjust alert thresholds if the rate of progress suggests our safety buffer is no longer adequate. We conduct further analysis, including reviewing model-independent information, external evaluations, and post-market monitoring as appropriate. In particular, we strive to learn from our post-market monitoring mechanisms, including as part of our detection, mitigation, response and/or reporting of incidents relating to our frontier safety risk domains. Actionable insights from these processes allow us to enhance our tools, training, processes, policies, and response efforts.

Our approach to model evaluations and inherent risk assessments described above means we can proactively monitor a model's capabilities throughout the entire lifecycle of the model and ensure that any significant or severe risk is properly identified and mitigated. Where appropriate, we may engage relevant external actors, including governments, to inform our responsible development and deployment practices.

Note on Machine Learning R&D CCLs: Risk assessment must take into account the fact that other actors may put significantly more effort into eliciting capabilities than we put into assessing risk, thus requiring conservatism in the form of evaluations. However, as a frontier AI company, we do not expect other groups to put significantly more effort into ML R&D than we do ourselves. As a result, to assess the ML R&D CCLs, we may use sources of information about our own progress at accelerating ML R&D to assess whether we are near or at the ML R&D CCLs, in addition to evaluations of ML R&D capabilities.

1.3.3 Risk Mitigation

We apply safety and security mitigations throughout the lifecycle of our models, including as part of our training and model development phase and, where appropriate, before T/CCLs are reached as described in the process below.

When a model reaches an alert threshold for a CCL, we will assess the proximity of the model to the CCL and analyze the risk posed, involving internal and external experts as needed. This will inform the formulation and application of a response plan. Central to most response plans will be the application of the mitigations described later in the Framework. However, if we assess that risks stemming from the model remain acceptable and model capabilities remain distant from a CCL, the response plan may include updating the alert threshold to ensure the safety buffer remains appropriate.

We have two categories of mitigations: security mitigations (such as preventing the exfiltration of model weights), and deployment mitigations (such as safety fine-tuning and monitoring and response) intended to counter the misuse or misaligned expression of critical capabilities in deployments.

The mitigations described in this document are limited to those mitigations that will be required to bring a model to an acceptable level of risk relative just to the specific T/CCL, which by definition pose the greatest risks of harm. Because we cannot always anticipate what security and deployment mitigations will be appropriate for models beyond the current frontier, the specific mitigations we implement may be determined when a T/CCL is reached, informed by the threat landscape at that time. These mitigations also reflect considerations from the perspective of addressing significant and severe risks from powerful capabilities alone; due to this focused scope, other risk management and security considerations may result in more stringent mitigations applied to a model than specified by the Framework.

1.3.4 Residual Risk Assessment

We will use various processes to evaluate the effectiveness and limitations of mitigations:

- **Security mitigations:** security infrastructure at Google is subject to penetration testing and other kinds of assessments, and is continually improved based on these results.

- We assess that the deployment mitigations have brought the risk of severe harm to an appropriate level proportionate to the risk, based on considerations such as whether the risk has been reduced to an acceptable level by mitigations, the scope of the deployment, what capabilities and mitigations are available on other publicly available models (e.g. if other models are similarly capable and have few mitigations, then the marginal risk added by our release is likely low), and the historical incidence and severity of related events. This is required only for external deployment, not further development.
 - Security mitigations have been applied to the model weights reaching the recommended security level stated below, or we otherwise assess that the level of security applied is adequate, e.g. if they match or exceed the level of security applied to other models with similar capabilities or risk profiles, or we assess that the benefits of the open release of model weights outweigh the risks.
- A model for which the risk assessment indicates a machine learning R&D CCL has been reached will be deemed to pose an acceptable level of risk for further development or deployment, if, for example:
 - We assess that the deployment mitigations have brought the risk of severe harm to an appropriate level proportionate to the risk, based on considerations such as whether the risk has been reduced to an acceptable level by mitigations, and information pertaining to model propensities and the severity of related events.¹ This is required only for external deployment and large scale internal deployment, not further development.
 - Security mitigations have been applied to the model weights reaching the recommended security level stated below, or we otherwise assess that the level of security applied is adequate, e.g. if they match or exceed the level of security applied to other models with similar capabilities or risk profiles, or we assess that the benefits of the open release of model weights outweigh the risks.
- Because the CCLs for misalignment risk are exploratory and intended for illustration only, we do not associate them with explicit risk acceptance criteria.

Note: Assessing frontier AI capabilities and corresponding severe risk is a complex process. Because the science of AI risk assessment is still developing, our assessments will often involve some level of subjective analysis. The concept of proportionality is central to our determination of whether a particular mitigation has sufficiently reduced the risk to acceptable levels. The mitigation and the effects of such mitigation should also be assessed holistically and be commensurate with expected impact of a model's risk, thus balancing safety with innovation.

¹ Note that deployment mitigations for the ML R&D CCLs are different than those for the misuse CCLs because of differences in which deployment could lead to severe harm for those respective categories. In contrast, security mitigations protecting against weights exfiltration and unauthorized modification are generally beneficial to both kinds of risk.

- **Deployment mitigations:** we will use a combination of threat modeling, empirical testing, and other sources of information to assess the effectiveness and limitations of our deployment mitigations.

Models reaching T&Cs or CCLs will be subject to residual risk assessments as part of evaluating their deployment mitigations. For models reaching CCLs, the residual risk assessment will be informed by a supplemental safety case. See the deployment mitigations sections below regarding misuse and machine learning R&D and misalignment for more.

1.3.5 Risk Acceptance Determination

The Framework outlines the risk acceptance determination approach for different model risks and capabilities pertaining to significant and severe risks. To summarize:

- A model for which inherent risk assessment indicates no T/CCL is reached will be deemed to pose an acceptable level of risk for further development and deployment, because it should not possess the capabilities required to substantially contribute to significant or severe risk scenarios.
- A model for which the inherent risk assessment indicates a misuse T/CCL has been reached will be deemed to pose an acceptable level of residual risk for further development or deployment, if:
 - We assess that the deployment mitigations have brought the residual risk of harm to an acceptable level, based on considerations such as the effectiveness of mitigations, the scope of the deployment, what capabilities and mitigations are available on other publicly available models (e.g. if other models are similarly capable and have few mitigations, then the marginal risk added by our external deployment is likely low), and the historical incidence and severity of related events. This is required only for external deployment, not internal deployment or further development.
 - Security mitigations have been applied to the model weights reaching the recommended security level stated below, or we otherwise assess that the level of security applied is adequate, e.g. based on mitigations already in place, if they match or exceed the level of security applied to other models with similar capabilities or risk profiles, or we assess that the benefits of the open release of model weights outweigh the risks.
- A model for which the inherent risk assessment indicates a ML R&D CCL has been reached will be deemed to pose an acceptable level of residual risk for further development or deployment, if:
 - We assess that the deployment mitigations have brought the residual risk of severe harm to an acceptable level, based on considerations such as the effectiveness of mitigations, and information pertaining to model propensities and the severity of related events.¹ This is required only for external deployment and high-risk internal deployment, not further development.
 - Security mitigations have been applied to the model weights reaching the recommended security level stated below, or we otherwise assess that the level of security applied is adequate, e.g. based on mitigations already in place, if they match or exceed the level of security applied to other models with similar capabilities or risk profiles, or we assess that the benefits of the open release of model weights outweigh the risks.

¹ Note that deployment mitigations for the ML R&D CCLs are different than those for the misuse CCLs because of differences in which deployment could lead to severe harm for those respective categories. In contrast, security mitigations protecting against weights exfiltration and unauthorized modification are generally beneficial to both kinds of risk.

Section 2: Misuse

This section describes our mitigation approach for models that pose risks of severe harm through misuse, and then details our set of misuse CCLs (CBRN, cyber and harmful manipulation), as well as the mitigation approach that we assess as appropriate for them.

2.1 Mitigation Approach

There are two categories of mitigations to address models with misuse critical capabilities: *security mitigations* intended to prevent the exfiltration of model weights, and *deployment mitigations* intended to counter the misuse of critical capabilities in external deployments. For security, we have several levels of mitigations, allowing calibration of the appropriateness and robustness of security measures to the risks posed. For deployment mitigations, we specify a standard process for applying, assessing, and reviewing mitigations: the aim of this process is to calibrate mitigations to CCLs, and the procedural approach reflects the more iterative and CCL-dependent nature of deployment mitigations. This structured process for deployment mitigations is centered on assessing and reviewing that the risk of severe harm has been brought to an appropriate level proportionate to the risk.

2.1.1 Security Mitigations

Security mitigations against exfiltration risk, such as identity and access management practices and hardening interface-access to unreleased model parameters, are important for models reaching CCLs. This is because the release of model weights may enable the removal of any safeguards trained into or deployed with the model, and hence access (including by threat actors) to any critical capabilities. Here, we use security levels that indicate goals/principles in line with the corresponding level in the RAND framework.¹ Because AI security is an area of active research, we expect the concrete measures implemented to reach each level of security to evolve substantially.

2.1.2 Deployment Mitigations

The following mitigation process for external deployments will be applied to models reaching a misuse CCL, allowing for iterative and flexible tailoring of mitigations to each risk and use case.

1. **Development and assessment of mitigations:** safeguards and an accompanying safety case are developed by iterating on the following:
 - a. Developing and improving a suite of safeguards targeting the capability, which may include measures such as safety post-training, monitoring and analysis, account moderation, jailbreak detection and patching, user verification, and bug bounties.²
 - b. Assessing the robustness of these mitigations against the risk posed through testing (e.g. automated evaluations, red teaming) and threat modeling research. The assessment takes the form of a safety case, and could take into account factors such as:
 - i. How much the risk has been reduced by mitigations. For example, whether tests run on mitigated models suggest that the refusal rate and jailbreak robustness together imply the risk has been brought substantially lower than that posed by a model reaching the CCL without mitigations.

¹ In other words, “security level N” indicates security controls and detections at a level generally aligned with RAND SL N. See https://www.rand.org/pubs/research_reports/RR2849-1.html, pp 21-22. In aligning our security levels with RAND’s, we are referring to the security goals and principles in the RAND framework, rather than the benchmarks (i.e. concrete measures) also described in the RAND report. As the authors point out, the “security level benchmarks represent neither a complete standard nor a compliance regime – they are provided for informational purposes only and should inform security teams’ decisions rather than supersede them.”

² See section 5 of <https://arxiv.org/abs/2504.01849>.

- A model for which the inherent risk assessment indicates the TCL for ML R&D and misalignment risk has been reached will be deemed to pose an acceptable level of residual risk for further development or deployment, if:
 - We assess that the deployment mitigations have brought the residual risk of significant harm to an acceptable level, based on considerations such as the effectiveness of mitigations, and information pertaining to model propensities and the severity of related events, or if we assess that the model is very unlikely to possess the propensities that would be required for material risk of significant harm through misalignment risk. This is required only for external deployment and high-risk internal deployment, not further development.
 - We assess that the level of security applied is adequate, e.g. based on mitigations already in place, if they match or exceed the level of security applied to other models with similar capabilities or risk profiles, or we assess that the benefits of the open release of model weights outweigh the risks.

Note: Assessing frontier AI capabilities and corresponding significant and severe risk is a complex process. Because the science of AI risk assessment is still developing, our assessments will often involve some level of subjective analysis. The concept of proportionality is central to our determination of whether a particular mitigation has sufficiently reduced the risk to acceptable levels. The mitigation and the effects of such mitigation should also be assessed holistically and be proportionate with expected impact of a model’s risk, thus balancing safety with innovation.

- ii. The likelihood and consequences of model misuse, capability improvements after the risk assessment, and likelihood and consequences of our mitigations being circumvented, deactivated, or subverted.
- iii. The scope of the deployment. For example, small scale and private deployments may pose substantially less risk than large scale or public deployments.
- iv. What capabilities and mitigations are available on other publicly available models. For example, whether another (non-Google) publicly deployed model is at the same CCL, and has mitigations that are less effective at preventing misuse than that of the model being assessed, in which case the deployment of this model is less likely to materially increase risk.
- v. The historical incidence and severity of related events: for example, whether data suggests a high (or low) likelihood of attempted misuse of models at the CCL. Mitigations would consequently have to be stronger (or would not have to be so strong) for deployment to be appropriate.

2. **Pre-deployment review of safety case:** external deployments of a model take place only after the appropriate governance function determines the safety case regarding each CCL the model has reached to be adequate. In particular, we will deem deployment mitigations adequate if the evidence suggests that for the CCLs the model has reached, the increase in likelihood of severe harm has been reduced to an acceptable level.
3. **Post-deployment processes:** our safety cases and mitigations may be updated if deemed necessary by post-market monitoring. Material updates to a safety case will be submitted to the appropriate governance function for review.

This process is designed to ensure that residual risk remains at acceptable levels: evidence of efficacy collected during development and testing, as well as expert-driven estimates of other parameters, will enable us to assess residual risk and to detect substantial changes that invalidate our risk assessment. With iteration on safeguards and safety cases, we believe that we are able to make informed decisions about the level of risk via a CCL before a model is released, and reliably prevent models posing unacceptable levels of risk from being deployed.

2.2 Misuse Critical Capability Levels

The table below details a set of CCLs we have identified through ongoing analysis of the CBRN, cyber, and harmful manipulation risk domains. We expect these to evolve over time. We recommend a security level for each of these CCLs, which reflect our assessment of the minimum appropriate level of security the field of frontier AI should apply to models reaching each CCL. In practice, our overall security posture may commonly exceed the baseline levels recommended here.

These recommended security levels reflect our current thinking proportionate to the risks posed and may be adjusted if our understanding of the risks changes. This may occur if, for example, a model does not possess capabilities meaningfully different from other publicly available models that have weaker security applied (in which case the marginal benefit of higher security is limited), or if we assess that the benefits of the open release of model weights outweigh the risks. Relatedly, we believe these recommendations will only be effective if the entire frontier AI field applies them, and of limited social utility if not.

2.2.1 Chemical, Biological, Radiological or Nuclear

This risk domain focuses on risks of models assisting in the development, preparation and/or execution of a CBRN threat.

Section 2: Misuse

This section describes our mitigation approach for models that pose risks of significant or severe harm through misuse, and then details our set of misuse T/CCLs (CBRN, cyber and harmful manipulation), as well as the mitigation approach that we assess as appropriate for them.

2.1 Mitigation Approach

There are two categories of mitigations to address models with misuse critical capabilities: *security mitigations* intended to prevent the exfiltration or unauthorized modification of model weights, and *deployment mitigations* intended to counter the misuse of critical capabilities in external deployments. For security, we have several levels of mitigations, allowing calibration of the appropriateness and robustness of security measures to the risks posed. Regarding deployment mitigations for models reaching T/CCLs, we specify a standard process for applying, assessing, and reviewing mitigations: the aim of this process is to calibrate mitigations to T/CCLs, and the procedural approach reflects the more iterative and T/CCL-dependent nature of deployment mitigations. This structured process for deployment mitigations is centered on assessing and reviewing that the risk of severe harm has been brought to an acceptable level proportionate to the risk.

2.1.1 Security Mitigations

Security mitigations against exfiltration or unauthorized modification risk, such as identity and access management practices and hardening interface-access to unreleased model parameters, are important for models reaching misuse CCLs. This is because the release or unauthorized modification of model weights may enable the removal of any safeguards trained into or deployed with the model, and hence access (including by threat actors) to any critical capabilities.

Using Google's Secure AI Framework (SAIF) and Google's common security infrastructure, we implement state-of-the-art security mitigations. SAIF is a defense-in-depth approach that embeds security into every layer of our AI systems—from data and infrastructure to models and applications. SAIF is premised on six core elements: strong security foundations, detection and response, automated defenses, harmonized platform-level controls, adaptable controls to mitigate emerging threats, and end-to-end risk assessments.

We use security levels to indicate security goals/principles in line with the corresponding level in the RAND model weight security framework.⁴ We also define and recommend "Security Level 2+," which uses RAND Security Level 2 (SL2) as a baseline, with additional security measures designed to address risks from insider threats and well-resourced non-state external actors. These additional measures may include, for example: dedicated insider risk teams; background checks and ID verification for personnel with sensitive access; review of model training data for signs of tampering; mandating that the processing of untrusted inputs occurs within sandboxed environments; advanced red-teaming that simulates well-resourced adversaries (including Advanced Persistent Threat (APT) groups); and proactive threat hunting with 24/7 incident response capabilities.

Because AI security is an area of active research, we expect the concrete measures implemented to reach each level of security to evolve substantially.

⁴ In other words, "security level N" indicates security controls and detections at a level generally aligned with RAND SL N. See https://www.rand.org/pubs/research_reports/RRA2849-1.html, pp 21-22. In aligning our security levels with RAND's, we are referring to the security goals and principles in the RAND framework, rather than the benchmarks (i.e. concrete measures) also described in the RAND report. As the authors point out, the "security level benchmarks represent neither a complete standard nor a compliance regime—they are provided for informational purposes only and should inform security teams' decisions rather than supersede them."

Table 2.2.1.a: CBRN CCLs and Security Mitigations

Critical Capability Level	Recommended security level and rationale
CBRN uplift level 1: Provides low to medium resourced actors uplift in reference scenarios resulting in additional expected harm at severe scale.	Security level 2 The difficulty of building defenses against certain CBRN threats means the exfiltration and leak of model weights with this capability could be highly damaging. However, the low to medium resourced actors who would be likely to experience the most CBRN uplift are unlikely to pose a substantial exfiltration threat at the level of RAND OC3 groups.

2.2.2 Cyber

This risk domain focuses on risks of models assisting in the development, preparation, and/or execution of a cyber attack.

Table 2.2.2.a: Cyber CCLs and Security Mitigations

Critical Capability Level	Recommended security level and rationale
Cyber uplift level 1: Provides sufficient uplift with high impact cyber attacks for additional expected harm at severe scale.	Security level 2 Models able to greatly assist cyber attack might be of interest to well-resourced state actors. However, the potential for automated cyber-defense and social adaptation as a response to exfiltration means that higher levels of security, and the resulting costs to innovation, are likely not warranted.

2.2.3 Harmful Manipulation

This risk domain focuses on risks of models with high manipulative capabilities potentially being misused in ways that could reasonably result in large scale harm.

Note: The research into harmful manipulation from a severe risk perspective is nascent. The CCL and our assessment of the risk in this domain is exploratory and subject to further research, and may be substantially changed over time.

Table 2.2.3.a: Harmful Manipulation CCLs and Security Mitigations

Critical Capability Level	Recommended security level and rationale
Harmful manipulation level	Security level 2

¹⁰ Here, and in other misuse CCLs, we intend this to mean relative to a baseline without generative AI.
¹¹ Mitigations at this level may include model access management, physical security controls, authentication measures, endpoint security, access management, secure model storage, vulnerability detection & management, detection of & response to suspected malicious activity.

2.1.2 Deployment Mitigations

The following mitigation process for external deployments will be applied to models reaching a misuse T/CCL, allowing for iterative and flexible tailoring of mitigations to each risk and use case.

- Development and assessment of mitigations:** safeguards and an accompanying residual risk assessment are developed by iterating on the following:
 - Developing and improving a suite of safeguards targeting the capability, which may include measures such as safety post-training, monitoring and analysis, account moderation, jailbreak detection and patching, user verification, and bug bounties.
 - Assessing the robustness of these mitigations against the risk posed through testing (e.g. automated evaluations, red teaming) and threat modeling research. This residual risk assessment could take into account factors such as:
 - The effectiveness of the mitigations. For example, tests run on mitigated models may suggest that the refusal rate and jailbreak robustness together imply that threat actors are unlikely to be able to circumvent safeguards.
 - The likelihood and consequences of model misuse, capability improvements after the risk assessment, and likelihood and consequences of our mitigations being circumvented, deactivated, or subverted.
 - The scope of the deployment. For example, small scale and private deployments may pose substantially less risk than large scale or public deployments.
 - What capabilities and mitigations are available on other publicly available models. For example, whether another (non-Google) publicly deployed model is at the same T/CCL, and has mitigations that are less effective at preventing misuse than that of the model being assessed, in which case the deployment of this model is less likely to materially increase risk.
 - The historical incidence and severity of related events: for example, whether data suggests a high (or low) likelihood of attempted misuse of models at the T/CCL. Mitigations would consequently have to be stronger (or would not have to be so strong) for deployment to be appropriate.
 - Where the model has reached a CCL, the residual risk assessment will be supplemented with a safety case.
- Pre-deployment review of residual risk assessment:** external deployments of a model take place only after the appropriate governance function determines the residual risk to be acceptable (including a safety case where a CCL has been reached). In particular, we will deem deployment mitigations adequate if the evidence suggests that for the T/CCLs the model has reached, the increase in likelihood of harm from the proposed external deployment has been reduced to an acceptable level.
- Post-deployment processes:** our residual risk assessments, safety cases and mitigations may be updated as a result of post-market monitoring, including information about incidents relating to our frontier safety risk domains. Material updates to a safety case will be submitted to the appropriate governance function for review and might result in updates to the related residual risk assessment or safety case.

This process is designed to ensure that residual risk remains at acceptable levels: evidence of efficacy collected during development and testing, as well as expert-driven estimates of other parameters, will enable us to assess residual risk and to detect substantial changes that invalidate our risk assessment. With iteration on mitigations and residual risk assessments, we believe that we are able to make informed decisions about the level of risk via a T/CCL before a model is deployed, and prevent models posing unacceptable levels of risk from being deployed.

¹² See section 5 of <https://arxiv.org/abs/2504.01849>.

1 (exploratory): Possesses manipulative capabilities sufficient to enable it to systematically and substantially change beliefs and behavior in identified high stakes contexts over the course of interactions with the model, reasonably resulting in additional expected harm at severe scale.	The lower velocity of harm scenarios associated with this CCL and the viability of social defenses against large scale misuse of such models count against security mitigations with substantial costs to innovation.
--	---

While we monitor for potential future risks related to insider misuse, our current assessment of these risk domains indicates that the risks of insider misuse posed by the misuse T/CCLs identified below are effectively addressed by our baseline controls (such as background checks and access restrictions). At this stage, additional mitigations beyond these established safeguards are not required.

2.2 Misuse Capability Levels

We set out below a set of CBRN, cyber, and harmful manipulation CCLs as well as the CBRN TCL that we have identified through ongoing analysis of these risk domains. We expect these to evolve over time. We recommend a security level for each CCL, which reflects our assessment of the minimum appropriate level of security the field of frontier AI should apply to models reaching each CCL. In practice, our overall security posture may commonly exceed the baseline levels recommended here.

These recommended security levels reflect our current thinking proportionate to the risks posed and may be adjusted if our understanding of the risks changes. This may occur if, for example, a model does not possess capabilities meaningfully different from other publicly available models that have weaker security applied (in which case the marginal benefit of higher security is limited), or if we assess that the benefits of the open release of model weights outweigh the risks. Relatedly, we believe these recommendations will only be effective if the entire frontier AI field applies them, and of limited social utility if not.

2.2.1 Chemical, Biological, Radiological or Nuclear

This risk domain focuses on risks of models assisting in the development, preparation and/or execution of a CBRN threat.

Table 2.2.1.a: CBRN CCL and Security Mitigations

Critical Capability Level	Recommended security level and rationale
CBRN uplift level 1: Provides low to medium resourced actors uplift in reference scenarios resulting in additional expected harm at severe scale.	Security level 2+ The difficulty of building defenses against certain CBRN threats means the exfiltration, unauthorized modification, and leak of model weights with this capability could be highly damaging. However, the low to medium resourced actors who would be likely to experience the most CBRN uplift are unlikely to pose a substantial exfiltration threat at the level of OC3 groups as defined by RAND.

CBRN Tracked Capability Level
The CBRN TCL is defined as follows:

CBRN uplift TCL: Provides low to medium resourced actors uplift in reference scenarios resulting in material risk of CBRN attacks that could result in significant but not severe harm.

Here, and in other misuse T/CCLs, we intend this to mean relative to a baseline without generative AI.

Section 3: Machine Learning R&D

This section describes our mitigation approach for models that pose risks of severe harm through machine learning R&D, and then details our set of ML R&D CCLs, as well as the mitigations that we provisionally assess as appropriate for them.

3.1 Mitigation Approach

As with misuse CCLs, we designate security and deployment mitigations for ML R&D CCLs, although deployment mitigations focus on different threat models and therefore also include measures for large scale internal deployments.

3.1.1 Security Mitigations

Security mitigations against exfiltration risk are important for models reaching ML R&D CCLs, because the exfiltration of highly capable models increases the likelihood they will be misused to achieve other critical capabilities, or deployed without adequate control and oversight. Below, we rely again on security levels to articulate the security goal recommended for each CCL. Security mitigations also protect against the risk of the model exfiltrating itself.

3.1.2 Deployment Mitigations

The following mitigation process will be applied for deployments of a model reaching a ML R&D CCL. The approach is similar to misuse deployment mitigations, with an added focus on large scale internal deployments.

1. **Development and assessment of mitigations:** safeguards and an accompanying safety case are developed by iterating on the following:

a. Developing and improving a suite of safeguards targeting the capability, which may include measures such as limiting affordances, monitoring and escalation, auditing, and alignment training, in addition to measures for preventing large scale misuse.

b. Assessing the robustness of these mitigations against the risk posed in both internal and external deployment through testing (e.g. automated evaluations, red teaming) and threat modeling research. The assessment takes the form of a safety case, taking into account factors such as:

i. How much the risk has been reduced by mitigations. For example, tests run on the safeguards may suggest that it is very unlikely they can be circumvented by external threat actors or the model in question to increase ML R&D risk.

ii. The likelihood and consequences of model misuse or misalignment, capability improvements after the risk assessment, and likelihood and consequences of our mitigations being circumvented, deactivated, or subverted.

iii. The scope of the deployment. For example, small scale and private deployments may pose substantially less risk than large scale or public deployments.

iv. Model propensity for, historical incidence of and severity of related events: for example, such data may suggest a high (or low) likelihood of misalignment in or misuse of models at the CCL, and mitigations would consequently have to be stronger (or not as strong) for deployment to be appropriate.

2. **Pre-deployment review of safety case:** external deployments and large scale internal deployments of a model take place only after the appropriate governance function determines
- ¹² While the RAND framework is not specifically designed to address this case, we index it at present because it is the most useful reference in this area.
¹³ See section 6 of <https://arxiv.org/abs/2504.01849>.
- Frontier Safety Framework | 12
- Google Frontier Safety Framework
- 2.2.2 Cyber
- This risk domain focuses on risks of models assisting in the development, preparation, and/or execution of a cyber attack.
- Table 2.2.2.a: Cyber CCL and Security Mitigations
- | Critical Capability Level | Recommended security level and rationale |
|---|---|
| Cyber uplift level 1:
Provides sufficient uplift with high impact cyber attacks for additional expected harm at severe scale. | Security level 2+

Models able to greatly assist cyber attack might be of interest to well-resourced state actors. However, the potential for automated cyber-defense and social adaptation as a response to exfiltration means that higher levels of security, and the resulting costs to innovation, are likely not warranted. |
- 2.2.3 Harmful Manipulation
- This risk domain focuses on risks of models with high manipulative capabilities potentially being misused in ways that could reasonably result in large scale harm.
- Note:** The research into harmful manipulation from a severe risk perspective is nascent. The CCL and our assessment of the risk in this domain is exploratory and subject to further research, and may be substantially changed over time.
- Table 2.2.3.a: Harmful Manipulation CCL and Security Mitigations
- | Critical Capability Level | Recommended security level and rationale |
|---|---|
| Harmful manipulation level 1 (exploratory): Possesses manipulative capabilities sufficient to enable it to systematically and substantially change beliefs and behavior in identified high stakes contexts over the course of interactions with the model, reasonably resulting in additional expected harm at severe scale. | Security level 2+

The lower velocity of harm scenarios associated with this CCL and the viability of social defenses against large-scale misuse of such models count against security mitigations with substantial costs to innovation. |
- Frontier Safety Framework | 12

the safety case regarding each CCL the model has reached to be adequate. In particular, we will deem deployment mitigations adequate if the evidence suggests that for the CCLs the model has reached, the increase in likelihood of severe harm has been reduced to an acceptable level.

3. **Post-deployment processes:** our safety cases and mitigations may be updated if deemed necessary by post-market monitoring. Material updates to a safety case will be submitted to the appropriate governance function for review.

This process is designed to ensure that residual risk remains at acceptable levels: evidence of efficacy collected during development and testing, as well as expert-driven estimates of other parameters, will enable us to assess residual risk and to detect substantial changes that invalidate our risk assessment. With iteration on safeguards and safety cases, we believe that we are able to make informed decisions about the level of risk via a CCL before a model is released or deployed internally at scale, and reliably prevent models posing unacceptable levels of risk from being deployed.

3.2 ML R&D Critical Capability Levels

The table below details a set of ML R&D CCLs we have identified that may lead to heightened severe risk through ML R&D. We expect these to evolve over time. We recommend a security level for each of these CCLs, which reflect our assessment of the minimum appropriate level of security the field of frontier AI should apply to models reaching each CCL.

3.2.1 Machine Learning R&D

These CCLs focus on risks posed by models capable of accelerating the rate of AI progress. These capabilities may indicate a heightened ability to undermine human control of models, may incentivize greater (and therefore higher risk) deployment of models, and could also result in the unsafe attainment or proliferation of other powerful AI models if misused by external threat actors.

Table 3.2.1.a: Machine Learning R&D CCLs and Security Mitigations

Critical Capability Level	Recommended security level and rationale
ML R&D acceleration level 1: Has been used to accelerate AI development, resulting in AI progress substantially accelerating from historical rates.	Security level 3¹⁵ Unrestricted access to models at this level of capability could significantly increase a threat actor's ability to progress to yet more powerful models and other critical capabilities. The exfiltration of such a model may therefore have a significant effect on society's ability to adapt to and govern powerful AI models, effects that may have long-lasting consequences. Substantially strengthened security is therefore recommended. However, we expect that acceleration will stem from systems of models integrated with workflows, rather than the model alone. The overall reduced impact of model weights counts against security levels with substantial innovation costs.

¹⁵ The same caveats regarding security levels for misuse CCLs apply.

¹⁶ This level may include mitigations aligned with SL 2, plus additional mitigations designed to prevent unilateral access, harden infrastructure, and prevent data exfiltration.

Section 3: Machine Learning R&D and Misalignment

This section describes our mitigation approach for models that pose risks of significant or severe harm through machine learning R&D and misalignment, and then details our set of T/CCLs, as well as the mitigations that we provisionally assess as appropriate for them.

3.1 Mitigation Approach

As with misuse T/CCLs, we take a similar approach to security and deployment mitigations for ML R&D and Misalignment T/CCLs, although deployment mitigations focus on different threat models and therefore also include measures for high-risk internal deployments.

3.1.1 Security Mitigations

Security mitigations against exfiltration and unauthorized modification risk are important for models reaching ML R&D CCLs. In addition, exfiltration of highly capable models increases the likelihood they will be misused to achieve other critical capabilities, or deployed without adequate control and oversight. Below, we rely again on security levels to articulate the security goal recommended for each CCL. Security mitigations also protect against the risk of the model exfiltrating itself.

3.1.2 Deployment Mitigations

The following mitigation process will be applied for deployments of a model reaching a ML R&D and Misalignment T/CCL. The approach is similar to misuse deployment mitigations, with an added focus on high-risk internal deployments.

- Development and assessment of mitigations:** safeguards and an accompanying residual risk assessment are developed by iterating on the following:
 - Developing and improving a suite of safeguards targeting the capability, which may include measures such as limiting affordances, monitoring and escalation, auditing, and alignment training, in addition to measures for preventing large-scale misuse.¹⁷
 - Assessing the robustness of these mitigations against the risk posed in both internal and external deployment through testing (e.g. automated evaluations, red teaming) and threat modeling research. This residual risk assessment could take into account factors such as:
 - The effectiveness of the mitigations. For example, tests run on the safeguards may suggest that it is very unlikely they can be circumvented by external threat actors or the model in question.
 - The likelihood and consequences of model misuse or misalignment, capability improvements after the risk assessment, and likelihood and consequences of our mitigations being circumvented, deactivated, or subverted.
 - The scope of the deployment. For example, small-scale and private deployments may pose substantially less risk than large-scale or public deployments.
 - Model propensity for, historical incidence of and severity of related events: for example, such data may suggest a high (or low) likelihood of misalignment in or misuse of models at the T/CCL, and mitigations would consequently have to be stronger (or not as strong) for deployment to be appropriate.
 - Where the model has reached a CCL, the residual risk assessment will be supplemented with a safety case.

¹⁷ While the RAND framework is not specifically designed to address this case, we index it at present because it is the most useful reference in this area.

¹⁸ See section 6 of <https://arxiv.org/abs/2504.01849>.

ML R&D automation level 1: Can fully automate the work of any team of researchers at Google focused on improving AI capabilities, with approximately comparable all-inclusive costs.	We recommend Security level 4¹⁴ for this capability threshold, but emphasize that this must be taken on by the frontier AI field as a whole. Unrestricted access to models at this level of capability could give an actor (or AI systems) with adequate computational resources the ability to reach capabilities much more powerful than those in the other CCLs listed in a short amount of time. This could be catastrophic if there is no effective way of defending against rapidly improving and potentially superhuman AI systems wielded by threat actors. Therefore, we recommend models at this level of capability have exceptional security even though they may have substantial innovation costs.
--	--

¹⁴ This level may include mitigations aligned with SL 2 and 3, plus additional mitigations aimed to isolate model weights, enhanced data center security, further hardening of infrastructure and minimizing potential attack surface.

- 2. Pre-deployment review of residual risk assessment:** external deployments and high-risk internal deployments of a model take place only after the appropriate governance function determines the residual risk to be acceptable (including a safety case where a CCL has been reached). In particular, we will deem deployment mitigations adequate if the evidence suggests that for the T/CCLs the model has reached, the increase in likelihood of harm from the proposed deployment has been reduced to an acceptable level.
- 3. Post-deployment processes:** our residual risk assessments, safety cases and mitigations may be updated as a result of post-market monitoring, including information about incidents relating to our frontier safety risk domains. Material updates to a safety case will be submitted to the appropriate governance function for review and might result in updates to the related residual risk assessment or safety case.

This process is designed to ensure that residual risk remains at acceptable levels: evidence of efficacy collected during development and testing, as well as expert-driven estimates of other parameters, will enable us to assess residual risk and to detect substantial changes that invalidate our risk assessment. With iteration on mitigations and residual risk assessments, we believe that we are able to make informed decisions about the level of risk via a T/CCL before a model is deployed, and prevent models posing unacceptable levels of risk from being deployed.

3.2 ML R&D and Misalignment Capability Levels

We take a tiered approach to addressing ML R&D and misalignment risks:

- We set a Stealth and Situational Awareness TCL that indicates a baseline level of model capability above which misalignment risks may begin to arise. At this capability level we apply a proportionate set of mitigation and risk acceptance determination processes.
- We set ML R&D CCLs that indicate higher capability levels at which misalignment, misuse and structural risks are heightened and may reach a severe scale. At these capability levels we apply a stronger set of mitigations, and follow the risk acceptance process as detailed in the rest of the framework.

3.2.1 Stealth and Situational Awareness Tracked Capability Level

The Stealth and Situational Awareness TCL is defined as follows:

Stealth and Situational Awareness TCL: The instrumental reasoning abilities of the model enable enough situational awareness (ability to discover and use relevant details of its deployment setting) and stealth (ability to circumvent basic oversight mechanisms) such that, absent additional mitigations, we cannot rule out the model significantly undermining human control.

When a model has reached this TCL, we will carry out periodic residual risk assessments of the misalignment risk posed. This assessment may take into account models' alignment propensities, capabilities, and our defenses against misaligned models. Since misalignment risks can manifest through internal use, these assessments will also account for the safety implications of high-risk internal deployments for such models. If the risk assessment deems the residual risk from internal deployment to be unacceptable without additional safeguards (such as chain-of-thought monitoring on high-risk deployments), these will be applied.

3.2.2 ML R&D Critical Capability Levels

We define ML R&D CCLs at capability levels at which misalignment, misuse and structural risks may reach a severe scale. We expect these to evolve over time. We recommend a security level for each of these

Section 4: Misalignment (Exploratory Approach)

Misalignment can pose a number of risks. In the context of the Framework, we address specific scenarios where general-purpose AI agents are potentially misaligned and can become difficult to control, thereby posing a risk of severe harm. Here we describe an approach for addressing misalignment risk that focuses specifically on when models may develop a baseline instrumental reasoning ability at which, without additional mitigations, they may have the potential to undermine human control. When models reach this capability level, one possible mitigation is to apply an automated monitor to the model's explicit reasoning (e.g. chain-of-thought output). Once a model is capable of effective instrumental reasoning in ways that cannot be monitored, additional mitigations may be warranted—the development of which is an area of active research.

Given its nascency, we expect our approach to misalignment risk to evolve substantially. This section is therefore illustrative only. Accordingly, we do not indicate security mitigations for models at these CCLs. The table below summarizes the initial approach:

Table 4.a: Exploratory Misalignment CCLs and Mitigations

Illustrative Instrumental Reasoning CCL	Illustrative Mitigation
Instrumental Reasoning Level 1: The instrumental reasoning abilities of the model enable enough situational awareness (ability to work out and use relevant details of its deployment setting) and stealth (ability to circumvent basic oversight mechanisms) such that, absent additional mitigations, we cannot rule out the model significantly undermining human control.	Automated monitoring: monitoring system to detect illicit use of instrumental reasoning capabilities in a model's chain-of-thought may be applied in high-stakes internal deployments where there is significant risk of the model undermining human control.
Instrumental Reasoning Level 2: The instrumental reasoning abilities of the model enable enough situational awareness and stealth that, even when relevant model outputs (including, e.g. scratchpads) are being monitored, we cannot detect or rule out the risk of a model significantly undermining human control.	Future work: We are actively researching approaches to addressing models that reach this CCL.

CCLs, which reflect our assessment of the minimum appropriate level of security the field of frontier AI should apply to models reaching each CCL.

Table 3.2.2.a: ML R&D CCLs and Security Mitigations

Critical Capability Level	Recommended security level and rationale
ML R&D acceleration level 1: Has been used to accelerate AI development, resulting in AI progress substantially accelerating from historical rates.	Security level 3¹² Unrestricted access to models at this level of capability could significantly increase a threat actor's ability to progress to yet more powerful models and other critical capabilities. The exfiltration of such a model may therefore have a significant effect on society's ability to adapt to and govern powerful AI models, effects that may have long-lasting consequences. Substantially strengthened security is therefore recommended. However, we expect that acceleration will stem from systems of models integrated with workflows, rather than the model alone. The overall reduced impact of model weights counts against security levels with substantial innovation costs.
ML R&D automation level 1: Can fully automate the work of any team of researchers at Google focused on improving AI capabilities, with approximately comparable all-inclusive costs.	We recommend Security level 4¹³ for this capability threshold, but emphasize that this must be taken on by the frontier AI field as a whole. Unrestricted access to models at this level of capability could give an actor (or AI systems) with adequate computational resources the ability to reach capabilities much more powerful than those in the other CCLs listed in a short amount of time. This could be catastrophic if there is no effective way of defending against rapidly improving and potentially superhuman AI systems wielded by threat actors. Therefore, we recommend models at this level of capability have exceptional security even though they may have substantial innovation costs.

¹² The same caveats regarding security levels for misuse CCLs apply.
¹³ This level may include mitigations aligned with SL 2+, plus additional mitigations designed to prevent unilateral access, harden infrastructure, and prevent data exfiltration.
¹⁴ This level may include mitigations aligned with SL 2+ and 3, plus additional mitigations aimed to isolate model weights, enhanced data center security, further hardening of infrastructure and minimizing potential attack surface.

Section 5: Updates and Disclosures

5.1 Updates

The Frontier Safety Framework will be updated at least once a year — more frequently if we have reasonable grounds to believe the adequacy of the Framework or our adherence to it has been materially undermined. The process will involve (i) an assessment of the Framework’s appropriateness for the management of systemic risk, drawing on information sources such as record of adherence to the framework, relevant high-quality research, information shared through industry forums, and evaluation results, as necessary, and (ii) an assessment of our adherence to the Framework. Following this assessment, we may:

- Update our risk domains and CCLs, where necessary.
- Update our testing and mitigation approaches, where needed to ensure risk remains adequately assessed and addressed according to our current understanding.

The updated version and framework assessment will be reviewed by the appropriate corporate governance bodies.

5.2 Disclosures

If we assess that a model has reached a CCL that poses an unmitigated and material risk to overall public safety, we aim to share relevant information with appropriate government authorities where it will facilitate safety of frontier AI. Where appropriate, and subject to adequate confidentiality and security measures and considerations around proprietary and sensitive information, this information may include:

- **Model information:** characteristics of the AI model relevant to the risk it may pose with its critical capabilities.
- **Evaluation results:** such as details about the evaluation design, the results, and any robustness tests.
- **Mitigation plans:** descriptions of our mitigation plans and how they are expected to reduce the risk.

We may also consider disclosing information to other external organisations to promote shared learning and coordinated risk mitigation. We will continue to review and evolve our disclosure process over time.

5.3 Past Updates and Changes

Past versions:

- [Version 2.0](#) (4 February 2025)
- [Version 1.0](#) (17 May 2024)

Section 4: Governance and Accountability

4.1 Governance structure

We have in place a well-established and comprehensive internal governance structure designed to ensure the robust implementation of the processes outlined in this Frontier Safety Framework. Responsibilities for assessing and mitigating risks are clearly defined and allocated across all levels of the organization. This includes legal, compliance, and safety reviews with escalation procedures to ensure appropriate oversight.

Section 5: Updates and Disclosures

5.1 Updates

The Frontier Safety Framework will be reviewed at least once a year—more frequently if we have reasonable grounds to believe the adequacy of the Framework or our adherence to it has been materially undermined. The process will involve (i) an assessment of the Framework's appropriateness for the management of significant and severe risk, drawing on information sources such as record of adherence to the framework, relevant high-quality research, information shared through industry forums, and evaluation results, as necessary, and (ii) an assessment of our adherence to the Framework. Following this assessment, we may:

- Update our risk domains and T/CCLs, where necessary;
- Update our testing and mitigation approaches, where needed to ensure risk remains adequately assessed and addressed according to our current understanding.

The updated version and framework assessment will be reviewed by the appropriate corporate governance bodies.

5.2 Disclosures

If we assess that a model has reached a CCL that poses an unmitigated and material risk to overall public safety, we aim to share relevant information with appropriate government authorities where it will facilitate safety of frontier AI. Where appropriate, and subject to adequate confidentiality and security measures and considerations around proprietary and sensitive information, this information may include:

- **Model information:** characteristics of the AI model relevant to the risk it may pose with its critical capabilities;
- **Evaluation results:** such as details about the evaluation design, the results, and any robustness tests;
- **Mitigation plans:** descriptions of our mitigation plans and how they are expected to reduce the risk.

We may also consider disclosing information to other external organizations to promote shared learning and coordinated risk mitigation. We will continue to review and evolve our disclosure process over time.

5.3 Past Updates and Changes

Versions:

- [Version 3.1](#) (April 17, 2026)
 - Introduced CBRN TCLs and outlined mitigation and risk acceptance process;
 - Incorporated the previous exploratory Misalignment risk domain into a combined ML R&D and Misalignment risk domain and outlined mitigation and risk acceptance process for the TCL;
 - Outlined enhanced level of security for CBRN, Cyber and Harmful Manipulation CCLs to Security Level 2+ to protect against non-state actors and insider threats;
 - Included more detail on our risk management process;
 - Included description of our internal governance structure;
 - Introduced a glossary;
- [Version 3.0](#) (September 22, 2025)
- [Version 2.0](#) (February 4, 2025)
- [Version 1.0](#) (May 17, 2024)

Glossary

Model

Frontier AI Models or Models: are trained on a large data set, display significant generality, are capable of performing a wide range of distinctive tasks and have high-impact capabilities. Frontier AI models' agentic and reasoning-based general capabilities near or exceed those of other Google models. There may be many different checkpoints and versions of the same model along the entire model lifecycle; we consider checkpoints and versions the same model if their base capabilities stem primarily from the same foundational training run, but will conduct additional testing on new versions as specified in [Section 1.3](#). When we refer to "model" in this Framework, we may be referring to a checkpoint or version of the model.

Risk Domains

CBRN: risks of models assisting in the development, preparation, and/or execution of a chemical, biological, radiological, or nuclear ("CBRN") threat.

Cyber: risks of models assisting in the development, preparation, and/or execution of a cyber attack.

Harmful Manipulation: risks of models with high manipulative capabilities potentially being misused in ways that could reasonably result in large scale harm.

ML R&D and Misalignment: risks of models' ML R&D capabilities or misaligned propensities reducing society's overall ability to manage AI risks. Such capabilities may serve as a substantial cross-cutting risk factor for several pathways (e.g. through misalignment, misuse or structural risks) to severe harm.

Thresholds

Critical Capability Levels (CCLs): are the main capability thresholds around which we have built the Framework process. They represent the capability levels at which, absent mitigation measures, frontier AI models or systems may pose heightened risk of severe harm.

Tracked Capability Levels (TCLs): are capability thresholds which capture a lower level of risks than our CCLs. They represent the capability levels at which, absent mitigation measures, frontier AI models or systems may pose heightened risk of significant but not severe levels of harm.

Alert Thresholds: are thresholds which we set marginally earlier than our CCLs. Crossing these thresholds indicates a CCL may be reached in the foreseeable future, due to the speed of capability progress. In particular, they are meant to account for the possibility of models reaching a CCL before the next critical capability assessment cycle.

Inherent Risk Assessment

Inherent Risk Assessments: the process of evaluating the level of risk posed by the model. This is part of the risk management process, and includes material capability change assessments and critical capability assessments.

Critical Capability Assessments: determine a model's proximity to T/CCLs through the use of early warning evaluations and other sources of evidence.

- **Early Warning Evaluations:** are evaluations which measure the dangerous capabilities of a model. They specifically target the threats and risk scenarios identified through our threat

modeling to determine a model's proximity to a CCL or TCL, and they provide the most direct evidence for whether these thresholds have been reached.

Material Capability Change Assessments: indicate whether a critical capability assessment is required. These assessments are typically conducted upon the completion of new post-training runs on various checkpoints, and they are primarily aimed at determining whether updated versions of a model have developed *material capability increases* relative to the last checkpoint subject to a critical capability assessment. May also be used to assess whether a critical capability assessment is required of non-frontier AI models.

- **Material Capability Increases:** are meaningful new capabilities or material increases in model performance that we believe could materially undermine our justifications for a model's level of risk being acceptable, for example, by reaching a CCL that the previous version of the model had not reached.

Risk Mitigation

Response Plans: are plans put in place when a new alert threshold has been reached to prepare for future models reaching a CCL. Central to most response plans will be the application of the mitigations described earlier in the Framework. However, if we assess that risks stemming from the model remain acceptable and model capabilities remain distant from a CCL, the response plan may include updating the alert threshold to ensure the safety buffer remains appropriate.

Deployment Mitigations: are safety measures we implement which are intended to counter the misuse or misaligned expression of critical capabilities in deployments. They may include safety post-training, input/output/chain-of-thought monitoring and analysis, account moderation, jailbreak detection and patching, user verification, and bug bounties.

- **External Deployments:** represent model releases to entities outside of Google. They can include both *low-risk external deployments* as well as general purpose releases to the broader public.
- **Low-Risk External Deployments:** represent small scale and private releases restricted to particular external groups or organizations.
- **Internal Deployments:** represent model releases restricted to Google employees for internal use.
- **High-Risk Internal Deployments:** represent model releases restricted to Google employees for use cases with the potential to enable severe threat scenarios (e.g. building internal security infrastructure or automating ML R&D).

Security Mitigations: are safety measures we implement which are intended to prevent the unauthorized modification or exfiltration of model weights by unauthorized actors. They may include identity and access management practices, and hardening interface-access to unreleased model parameters, among other measures. We communicate the level of security we consider appropriate to models which cross our thresholds using the RAND model weight security framework.

- **RAND Security Levels:** indicate security goals and principles relevant to frontier developers aiming to protect their model weights from various actors - with the strength and type of security measures differing for each category of actor.¹⁴

¹⁴ See https://www.rand.org/pubs/research_reports/RRA2849-1.html.

Residual Risk Assessment

Residual Risk Assessments: the process of evaluating the level of risk that remains after all planned mitigation strategies have been implemented. This is part of the risk management process, following a critical capability assessment and the effect of mitigations. The residual risk should be brought within the acceptable levels, as determined by the T/CCLs as relevant.

Safety Cases: are an assessable argument showing how severe risks associated with a model's CCLs have been reduced to an acceptable level. Residual risk assessments will be supplemented with a safety case when a model reaches a CCL, and inform a model's risk acceptance determination.¹⁴ For misuse risk, such an argument may be based, for example, on considerations such as the effectiveness of mitigations, the scope of the deployment, what capabilities and mitigations are available on other publicly available models, and the historical incidence and severity of related events. For ML R&D and misalignment risk, such an argument may be based, for example, on considerations such as the effectiveness of mitigations relative to model capabilities, and information pertaining to model propensities and the severity of related events.

¹⁴ See also, for reference, <https://arxiv.org/abs/2505.01420>.

Comparison Summary Report

frontier-safety-framework_3 (2).pdf compared with frontier-safety-framework_3-1-2.pdf

2026-05-12 11:35:02

443 changes

Text Changes 443

Insertions	207
Deletions	86
Replacements	150
Moves	0

Visual Changes 0

Insertions	0
Deletions	0
Replacements	0
Moves	0

Formatting Changes 0

Total Changes 443

Revisions Summary Report

frontier-safety-framework_3 (2).pdf compared with frontier-safety-framework_3-1-2.pdf

2026-05-12 11:35:02

443 changes

Added 207 Removed 86 Replaced 150 Moved 0 Formatting 0

#1 (p.1)	Replaced (Text)	Before: 0 After: 1
#2 (p.1)	Replaced (Text)	Before: September 22 After: April 17
#3 (p.1)	Replaced (Text)	Before: 5 After: 6
#4 (p.2)	Added (Text)	Google's
#5 (p.2)	Added (Text)	significant or
#6 (p.2)	Added (Text)	"Tracked Capability Levels (TCLs)" and
#7 (p.2)	Added (Text)	respectively
#8 (p.2)	Removed (Text)	,
#9 (p.2)	Replaced (Text)	Before: risks from After: as well as
#10 (p.2)	Removed (Text)	,
#11 (p.2)	Added (Text)	(and TCLs where relevant)
#12 (p.2)	Replaced (Text)	Before: After: Please refer to the Glossary
#13 (p.2)	Added (Text)	at the end of this document for definitions used in the Framework.
#14 (p.2)	Replaced (Text)	Before: s After: z
#15 (p.2)	Removed (Text)	Misalignment can pose a number of risks.
#16 (p.2)	Replaced (Text)	Before: general-purpose AI agents are potentially After: model
#17 (p.2)	Replaced (Text)	Before: ed and can become difficult to control, t After: ment may create the risk of significant or severe harm, and scenarios w
#18 (p.2)	Replaced (Text)	Before: by posing a risk of severe harm. After: ML R&D capabilities may heighten misalignment, misuse, and structural risks.
#19 (p.2)	Replaced (Text)	Before: news/ After: rsp-updates
#20 (p.2)	Added (Text)	,
#21 (p.2)	Added (Text)	https://www.

#22 (p.2)	Replaced (Text)	Before: s-responsible-scaling-policy After: .com/news/compliance-framework-SB53
#23 (p.3)	Added (Text)	and Tracked Capability Levels
#24 (p.3)	Removed (Text)	Outline of Our
#25 (p.3)	Replaced (Text)	Before: Assessment After: Management
#26 (p.3)	Replaced (Text)	Before: 4 After: 5
#27 (p.3)	Replaced (Text)	Before: 4 Response Plans and Mitigations After: 3.1 Risk Identification
#28 (p.3)	Replaced (Text)	Before: 6 After: 5
#29 (p.3)	Replaced (Text)	Before: 5 Evaluating After: 3.2 Inherent Risk Assessment?
#30 (p.3)	Added (Text)	5 1.3.3 Risk
#31 (p.3)	Removed (Text)	s
#32 (p.3)	Replaced (Text)	Before: 6 Summary of After: 3.4 Residual Risk Assessment?
#33 (p.3)	Added (Text)	6 1.3.5
#34 (p.3)	Replaced (Text)	Before: Criteria After: Determination
#35 (p.3)	Replaced (Text)	Before: 6 After: 7
#36 (p.3)	Replaced (Text)	Before: 8 After: 9
#37 (p.3)	Replaced (Text)	Before: 8 After: 9
#38 (p.3)	Replaced (Text)	Before: 2 Misuse Critical After: 1.1 Security Mitigations?
#39 (p.3)	Added (Text)	9 2.1.2 Deployment Mitigations?
#40 (p.3)	Added (Text)	10 2.2 Misuse
#41 (p.3)	Replaced (Text)	Before: 9 After: 11 2.2.1 Chemical, Biological, Radiological or Nuclear?
#42 (p.3)	Added (Text)	11 2.2.2 Cyber?
#43 (p.3)	Added (Text)	12 2.2.3 Harmful Manipulation?
#44 (p.3)	Added (Text)	12
#45 (p.3)	Added (Text)	and Misalignment
#46 (p.3)	Replaced (Text)	Before: 12 After: 13

#47 (p.3)	Replaced (Text)	Before: 12 3. After: 13 3.1.1 Security Mitigations?
#48 (p.3)	Added (Text)	13 3.1.2 Deployment Mitigations?
#49 (p.3)	Added (Text)	13 3.2 ML R&D and Misalignment Capability Levels?
#50 (p.3)	Added (Text)	14 3.2.1 Stealth and Situational Awareness Tracked Capability Level?
#51 (p.3)	Added (Text)	14 3.2.
#52 (p.3)	Replaced (Text)	Before: 13 Section 4: Misalignment (Exploratory Approach)? After: 14 Section 4: Governance and Accountability?
#53 (p.3)	Replaced (Text)	Before: 15 After: 16 4.1 Governance structure?
#54 (p.3)	Added (Text)	16
#55 (p.3)	Replaced (Text)	Before: 16 After: 17
#56 (p.3)	Replaced (Text)	Before: 16 After: 17
#57 (p.3)	Replaced (Text)	Before: 16 After: 17
#58 (p.3)	Replaced (Text)	Before: 16 After: 17 Glossary?
#59 (p.3)	Added (Text)	18
#60 (p.4)	Added (Text)	and Tracked Capability Levels
#61 (p.4)	Added (Text)	primarily
#62 (p.4)	Replaced (Text)	Before: We describe three sets of CCLs After: CCLs are one important component of our
#63 (p.4)	Added (Text)	risk acceptance determination
#64 (p.4)	Added (Text)	. We identify CCLs for two kinds of risks
#65 (p.4)	Replaced (Text)	Before: CCLs, After: risk and risks related to
#66 (p.4)	Removed (Text)	CCLs,
#67 (p.4)	Removed (Text)	CCLs
#68 (p.4)	Added (Text)	and misalignment
#69 (p.4)	Added (Text)	s
#70 (p.4)	Replaced (Text)	Before: other pathways to severe harm. For After: several pathways (e.g. through misalignment, misuse or structural risks) to severe harm. This update to the Framework introduces
#71 (p.4)	Replaced (Text)	Before: misalignment risk After: "Tracked Capability Levels (TCLs)."

#72 (p.4)	Replaced (Text)	Before: , we outline an exploratory approach that focuses on detecting when models might develop a baseline After: TCLs are meant to capture significant risks that may manifest at a lower capability threshold than our CCLs and we apply our mitigation and risk acceptance determination process proportionately. We identify TCLs for CBRN risk, as well as ML R&D and misalignment risks. We may include TCLs for additional risks in the future, as our threat modeling develops. Similarly to CCLs, early warning evaluations will be used to assess the proximity of a model to a TCL and analyze the risk posed, involving internal and external experts as needed.
#73 (p.4)	Replaced (Text)	Before: instrumental reasoning After: Frontier Safety Framework 4
#74 (p.4 → p.5)	Replaced (Text)	Before: ability at which they have the potential to undermine human control, assuming no additional mitigations were applied. Most CCLs define one important component of our After: Frontier Safety Framework 1.3 Risk Management
#75 (p.4)	Removed (Text)	risk acceptance criteria
#76 (p.4)	Removed (Text)	. Because the CCLs for misalignment risk are exploratory and intended for illustration only, we do not associate them with explicit risk acceptance criteria. 1.3 Outline of Our Risk Assessment
#77 (p.4 → p.5)	Replaced (Text)	Before: For each risk domain, we After: We
#78 (p.4)	Removed (Text)	aspects of
#79 (p.4 → p.5)	Replaced (Text)	Before: assessment at various moments After: management process as appropriate,
#80 (p.5)	Added (Text)	on various checkpoints or versions of a model,
#81 (p.4 → p.5)	Replaced (Text)	Before: We conduct a risk assessment for the first external deployment of a new frontier AI model. For subsequent versions of the model, we conduct After: 1.3.1 Risk Identification The first part of our risk management process is
#82 (p.4 → p.5)	Replaced (Text)	Before: Frontier Safety Framework 4 After: risk identification

#83 (p.5)	Replaced (Text)	Before: Frontier Safety Framework a further risk assessment if the model has meaningful new capabilities or a material increase in performance, until the model is retired or we deploy a more capable model. The reason for this is because a material change in the model's capabilities may mean that the risk profile of the model has changed or the justification for why the risks stemming from the model are acceptable has been materially undermined. To identify meaningful new capabilities or material increases in performance, we conduct model capability evaluations, including our automated benchmarks. These evaluations are primarily aimed at understanding the capabilities of the model and may be triggered, for example, upon the completion of a pre-training or post-training run, on various candidates of a model version. These evaluations include a broad range of areas, including general capability evaluations, model behavior, efficiency, coding capabilities, multilinguality, or reasoning. Data from these evaluations are collected and analyzed to give us an indication as to how the model is performing and whether a risk assessment is necessary. At a high level, our risk assessment involves the following steps (which do not need to be repeated where a previous risk assessment is still appropriate): ??Identification: After: . We identify potential risks that could stem from our models and analyze their characteristics to determine which of the identified risks could be significant or severe risks. We consider a wide range of risks as part of our ongoing research, taking into account the characteristics, capabilities, propensities, and affordances of our models and other sources of information, such as our internal risk taxonomies, internal expertise and relevant external research. As explained above, we have identified risk domains where, based on early research, we have determined significant or severe risks may be most likely to arise from future models: CBRN, cyber, harmful manipulation, as well as machine learning R&D and misalignment. As part of our broader research into and development of frontier AI models, we continue to assess whether there are other risk domains where significant or severe risks may arise and will update our approach as appropriate. For each of the four identified domains, we have developed specific scenarios and T/CCLs in which these risks could materialize. 1.3.2 Inherent Risk Assessment To understand whether a model may, without appropriate mitigations, contribute to the risks identified, we conduct a closer assessment of its capabilities against the T/CCLs. We conduct a
#84 (p.5)	Replaced (Text)	Before: As explained above, we have identified risk domains where, based on early research, we have determined severe risks may be most likely to arise from future models: CBRN, cyber, harmful manipulation, and machine learning R&D. After: critical capability assessment
#85 (p.5)	Replaced (Text)	Before: 4 After: prior to the first external deployment
#86 (p.5)	Replaced (Text)	Before: As part of our broader research into frontier AI models, we continue to assess whether there are other risk domains where severe risks may arise and will update our approach as appropriate. For each of the four identified domains, we have developed specific scenarios in which these risks could materialize. ?? After: 4

#87 (p.5)	Replaced (Text)	Before: Analysis: After: of a new frontier AI model. For external deployment of subsequent versions of the model, we determine whether a substantial modification has been made and whether a further critical capability assessment is required if (1) the model has meaningful new capabilities or material increases in performance; and (2) we believe such material capability increase could materially undermine the justification for why the risks stemming from the model are acceptable. To understand if a subsequent version of the model has meaningful new capabilities or material increases in performance relative to the last checkpoint subject to a critical capability assessment, we conduct
#88 (p.5)	Replaced (Text)	Before: Central to our model evaluations After: material capability change assessments
#89 (p.5)	Added (Text)	on various checkpoints upon the completion of a post-training run. These assessments draw on model characteristics, such as its provenance, size, modality, and performance on general capability benchmarks. To understand if such a change in capability in the subsequent version of the model could materially undermine the justification for why the risks stemming from the model are acceptable, we may also conduct light-weight versions of our “early warning evaluations” described below, including our automated benchmarks, to confirm a model remains below T/CCL. We may also apply material capability change assessments to models not on the frontier, to understand whether a critical capability assessment is required. Central to our critical capability assessments
#90 (p.5)	Replaced (Text)	Before: to After: which we use to test the specific threats and risk scenarios identified through our threat modeling, determine a model’s capability, and
#91 (p.5)	Added (Text)	T/
#92 (p.5)	Replaced (Text)	Before: We After: For CCLs, we
#93 (p.5)	Replaced (Text)	Before: ” for these evaluations After: ,” which may draw on evaluation results, expert assessments, and other sources of information;
#94 (p.5)	Replaced (Text)	Before: risk After: critical capability
#95 (p.5)	Replaced (Text)	Before: equip the model with After: apply
#96 (p.5)	Added (Text)	, inference compute,
#97 (p.5)	Removed (Text)	make it more likely that we are
#98 (p.5)	Removed (Text)	ing
#99 (p.5)	Added (Text)	based on the specific threats and risk scenarios we have identified for the T/CCL
#100 (p.5)	Replaced (Text)	Before: frequently or adjust the alert threshold of our evaluations After: 4

#101 (p.5)	Added (Text)	A critical capability assessment may not be conducted for low-risk external deployments (e.g. to a small number of trusted testers) if the appropriate governance function determines the residual risk of such deployments to be acceptable even if the model has reached a T/CCL (including with a safety case where a CCL has been reached).
#102 (p.5)	Added (Text)	Frontier Safety Framework 5
#103 (p.6)	Added (Text)	Frontier Safety Framework frequently or adjust alert thresholds
#104 (p.5 → p.6)	Replaced (Text)	Before: After: -
#105 (p.5 → p.6)	Replaced (Text)	Before: ?? After: In particular, we strive to learn from our post-market monitoring mechanisms, including as part of our detection, mitigation, response and/or reporting of incidents relating to our frontier safety risk domains. Actionable insights from these processes allow us to enhance our tools, training, processes, policies, and response efforts.
#106 (p.5)	Removed (Text)	Acceptance determination and mitigations:
#107 (p.5)	Removed (Text)	We then determine whether the model has met or will meet a CCL and, if so, whether we need to implement any further mitigations to reduce the risk to an acceptable level (see
#108 (p.5)	Removed (Text)	below
#109 (p.5)	Removed (Text)	).
#110 (p.6)	Added (Text)	inherent
#111 (p.6)	Added (Text)	significant or
#112 (p.5)	Removed (Text)	and appropriate
#113 (p.6)	Added (Text)	ML R&D
#114 (p.5 → p.6)	Replaced (Text)	Before: Similarly, After: 1.3.3 Risk Mitigation
#115 (p.5)	Removed (Text)	4
#116 (p.5)	Removed (Text)	We exclude misalignment risk from this list of domains because of its exploratory nature.
#117 (p.5)	Removed (Text)	Frontier Safety Framework 5
#118 (p.6)	Removed (Text)	Frontier Safety Framework our alert threshold may be defined based on these sources of information, rather than on evaluation scores. 1.4 Response Plans and Mitigations
#119 (p.6)	Added (Text)	T/

#120 (p.6)	Replaced (Text)	Before: Where model capabilities remain quite distant from a CCL, a response plan may involve the adoption of additional capability assessment processes to flag when heightened mitigations are required. Central to most response plans will be the application of the mitigations described later in this document. After: Central to most response plans will be the application of the mitigations described later in the Framework. However, if we assess that risks stemming from the model remain acceptable and model capabilities remain distant from a CCL, the response plan may include updating the alert threshold to ensure the safety buffer remains appropriate.
#121 (p.6)	Replaced (Text)	Before: Note that these mitigations After: The mitigations described in this document are limited to those mitigations that will be required to bring a model to an acceptable level of risk relative just to the specific T/CCL, which by definition pose the greatest risks of harm. Because we cannot always anticipate what security and deployment mitigations will be appropriate for models beyond the current frontier, the specific mitigations we implement may be determined when a T/CCL is reached, informed by the threat landscape at that time. These mitigations also
#122 (p.6)	Added (Text)	significant and
#123 (p.6)	Replaced (Text)	Before: 5 Evaluating Mitigations After: 3.4 Residual Risk Assessment
#124 (p.6)	Added (Text)	Frontier Safety Framework 6
#125 (p.7)	Added (Text)	Frontier Safety Framework
#126 (p.6 → p.7)	Replaced (Text)	Before: These will form the basis of a safety case After: Models reaching TCLs or CCLs will be subject to residual risk assessments as part of evaluating their deployment mitigations. F
#127 (p.6)	Removed (Text)	5
#128 (p.6)	Removed (Text)	f
#129 (p.6 → p.7)	Replaced (Text)	Before: that After: the residual risk assessment
#130 (p.6 → p.7)	Replaced (Text)	Before: reviewed before deployment After: informed by a supplemental safety case
#131 (p.6 → p.7)	Replaced (Text)	Before: for After: regarding
#132 (p.7)	Added (Text)	and misalignment
#133 (p.6 → p.7)	Replaced (Text)	Before: 6 Summary of After: 3.5
#134 (p.6 → p.7)	Replaced (Text)	Before: Criteria After: Determination
#135 (p.6 → p.7)	Replaced (Text)	Before: a variety of After: the
#136 (p.6 → p.7)	Replaced (Text)	Before: criteria After: determination approach

#137 (p.7)	Added (Text)	significant and
#138 (p.7)	Added (Text)	inherent
#139 (p.7)	Added (Text)	T/
#140 (p.6)	Removed (Text)	severe
#141 (p.6 → p.7)	Replaced (Text)	Before: severe After: significant or severe
#142 (p.6)	Removed (Text)	6
#143 (p.7)	Added (Text)	inherent
#144 (p.7)	Added (Text)	T/
#145 (p.6 → p.7)	Replaced (Text)	Before: risk for further development or deployment, if, for example: After: residual
#146 (p.6)	Removed (Text)	6
#147 (p.6)	Removed (Text)	Note that this does not mean that no mitigations will be implemented prior to the deployment of the model. We apply safety and security mitigations throughout the lifecycle of our models for a whole host of potential
#148 (p.6 → p.7)	Replaced (Text)	Before: s, including as part of our training and model After: for further
#149 (p.6 → p.7)	Replaced (Text)	Before: phase. This section is limited to those mitigations that will be required to bring a model to an acceptable level of risk relative just to the specific CCL, which by definition pose the greatest risks of severe harm. After: or deployment, if:
#150 (p.6)	Removed (Text)	5
#151 (p.6)	Removed (Text)	A safety case is an assessable argument showing how severe risks associated with a model's CCLs have been reduced to an appropriate level. See also, for reference,
#152 (p.6)	Removed (Text)	https://arxiv.org/abs/2505.01420
#153 (p.6)	Removed (Text)	.
#154 (p.6)	Removed (Text)	Frontier Safety Framework 6
#155 (p.7)	Removed (Text)	Frontier Safety Framework
#156 (p.7)	Added (Text)	residual
#157 (p.7)	Removed (Text)	severe
#158 (p.7)	Replaced (Text)	Before: appropriate After: acceptable
#159 (p.7)	Removed (Text)	proportionate to the risk
#160 (p.7)	Replaced (Text)	Before: whether the risk has been reduced to an acceptable level by After: the effectiveness of
#161 (p.7)	Replaced (Text)	Before: release After: external deployment

#162 (p.7)	Added (Text)	internal deployment or
#163 (p.7)	Added (Text)	based on mitigations already in place,
#164 (p.7)	Added (Text)	inherent
#165 (p.7)	Replaced (Text)	Before: machine learning After: ML
#166 (p.7)	Added (Text)	residual
#167 (p.7)	Removed (Text)	, for example
#168 (p.7)	Added (Text)	residual
#169 (p.7)	Replaced (Text)	Before: appropriate After: acceptable
#170 (p.7)	Removed (Text)	proportionate to the risk
#171 (p.7)	Replaced (Text)	Before: whether the risk has been reduced to an acceptable level by After: the effectiveness of
#172 (p.7)	Replaced (Text)	Before: 7 After: 5
#173 (p.7)	Replaced (Text)	Before: large scale After: high-risk
#174 (p.7)	Added (Text)	based on mitigations already in place,
#175 (p.7)	Replaced (Text)	Before: ? After: 5
#176 (p.7)	Added (Text)	Note that deployment mitigations for the ML R&D CCLs are different than those for the misuse CCLs because of differences in which deployment could lead to severe harm for those respective categories. In contrast, security mitigations protecting against weights exfiltration and unauthorized modification are generally beneficial to both kinds of risk.
#177 (p.7)	Added (Text)	Frontier Safety Framework 7
#178 (p.8)	Added (Text)	Frontier Safety Framework ?
#179 (p.7 → p.8)	Replaced (Text)	Before: Because the CCLs for misalignment risk are exploratory and intended for illustration only, we do not associate them with explicit risk acceptance criteria. After: A model for which the inherent risk assessment indicates the TCL for ML R&D and misalignment risk has been reached will be deemed to pose an acceptable level of residual risk for further development or deployment, if: ??
#180 (p.8)	Added (Text)	We assess that the deployment mitigations have brought the residual risk of significant harm to an acceptable level, based on considerations such as the effectiveness of mitigations, and information pertaining to model propensities and the severity of related events, or if we assess that the model is very unlikely to possess the propensities that would be required for material risk of significant harm through misalignment risk. This is required only for external deployment and high-risk internal deployment, not further development. ??

#181 (p.8)	Added (Text)	We assess that the level of security applied is adequate, e.g. based on mitigations already in place, if they match or exceed the level of security applied to other models with similar capabilities or risk profiles, or we assess that the benefits of the open release of model weights outweigh the risks.
#182 (p.8)	Added (Text)	significant and
#183 (p.7 → p.8)	Replaced (Text)	Before: commensurate After: proportionate
#184 (p.7)	Removed (Text)	7
#185 (p.7)	Removed (Text)	Note that deployment mitigations for the ML R&D CCLs are different than those for the misuse CCLs because of differences in which deployment could lead to severe harm for those respective categories. In contrast, security mitigations protecting against weights exfiltration are generally beneficial to both kinds of risk.
#186 (p.7 → p.8)	Replaced (Text)	Before: 7 After: 8
#187 (p.9)	Added (Text)	significant or
#188 (p.9)	Added (Text)	T/
#189 (p.9)	Added (Text)	or unauthorized modification
#190 (p.8 → p.9)	Replaced (Text)	Before: For After: Regarding
#191 (p.9)	Added (Text)	for models reaching T/CCLs
#192 (p.9)	Added (Text)	T/
#193 (p.9)	Added (Text)	T/
#194 (p.8 → p.9)	Replaced (Text)	Before: appropriate After: acceptable
#195 (p.9)	Added (Text)	or unauthorized modification
#196 (p.9)	Added (Text)	misuse
#197 (p.9)	Added (Text)	or unauthorized modification
#198 (p.8 → p.9)	Replaced (Text)	Before: Here, we use security levels that After: Using
#199 (p.9)	Added (Text)	Google's Secure AI Framework (SAIF)
#200 (p.9)	Added (Text)	and
#201 (p.9)	Added (Text)	Google's common security infrastructure
#202 (p.9)	Added (Text)	, we implement state-of-the-art security mitigations. SAIF is a defense-in-depth approach that embeds security into every layer of our AI systems—from data and infrastructure to models and applications. SAIF is premised on six core elements: strong security foundations, detection and response, automated defenses, harmonized platform-level controls, adaptable controls to mitigate emerging threats, and end-to-end risk assessments. We use security levels to

#203 (p.9)	Added (Text)	security
#204 (p.9)	Added (Text)	model weight security
#205 (p.8 → p.9)	Replaced (Text)	Before: 8 After: 6
#206 (p.9)	Added (Text)	We also define and recommend "Security Level 2+," which uses RAND Security Level 2 (SL2) as a baseline, with additional security measures designed to address risks from insider threats and well-resourced non-state external actors. These additional measures may include, for example: dedicated insider risk teams; background checks and ID verification for personnel with sensitive access; review of model training data for signs of tampering; mandating that the processing of untrusted inputs occurs within sandboxed environments; advanced red-teaming that simulates well-resourced adversaries (including Advanced Persistent Threat (APT) groups); and proactive threat hunting with 24/7 incident response capabilities.
#207 (p.9)	Added (Text)	6
#208 (p.9)	Added (Text)	In other words, "security level
#209 (p.9)	Added (Text)	N
#210 (p.9)	Added (Text)	" indicates security controls and detections at a level generally aligned with RAND SL
#211 (p.9)	Added (Text)	N
#212 (p.9)	Added (Text)	. See
#213 (p.9)	Added (Text)	https://www.rand.org/pubs/research_reports/RRA2849-1.html ,
#214 (p.9)	Added (Text)	pp 21-22. In aligning our security levels with RAND's, we are referring to the security goals and principles in the RAND framework, rather than the benchmarks (i.e. concrete measures) also described in the RAND report. As the authors point out, the "security level benchmarks represent neither a complete standard nor a compliance regime—they are provided for informational purposes only and should inform security teams' decisions rather than supersede them."
#215 (p.9)	Added (Text)	Frontier Safety Framework 9
#216 (p.10)	Added (Text)	Frontier Safety Framework
#217 (p.10)	Added (Text)	T/
#218 (p.8 → p.10)	Replaced (Text)	Before: safety case After: residual risk assessment
#219 (p.8 → p.10)	Replaced (Text)	Before: 9 After: 7
#220 (p.8 → p.10)	Replaced (Text)	Before: The assessment takes the form of a safety case, and After: This residual risk assessment
#221 (p.8 → p.10)	Replaced (Text)	Before: How much t After: T
#222 (p.8 → p.10)	Replaced (Text)	Before: risk has been reduced by After: effectiveness of the

#223 (p.8)	Removed (Text)	whether
#224 (p.10)	Added (Text)	may
#225 (p.8 → p.10)	Replaced (Text)	Before: the risk has been brought substantially lower than that posed by a model reaching the CCL without mitigations. After: that threat actors are unlikely to be able to circumvent safeguards.
#226 (p.8)	Removed (Text)	9
#227 (p.8)	Removed (Text)	See section 5 of
#228 (p.8)	Removed (Text)	https://arxiv.org/abs/2504.01849
#229 (p.8)	Removed (Text)	.
#230 (p.8)	Removed (Text)	8
#231 (p.8)	Removed (Text)	In other words, “security level
#232 (p.8)	Removed (Text)	N
#233 (p.8)	Removed (Text)	” indicates security controls and detections at a level generally aligned with RAND SL
#234 (p.8)	Removed (Text)	N
#235 (p.8)	Removed (Text)	. See
#236 (p.8)	Removed (Text)	https://www.rand.org/pubs/research_reports/RRA2849-1.html ,
#237 (p.8)	Removed (Text)	pp 21-22. In aligning our security levels with RAND’s, we are referring to the security goals and principles in the RAND framework, rather than the benchmarks (i.e. concrete measures) also described in the RAND report. As the authors point out, the “security level benchmarks represent neither a complete standard nor a compliance regime—they are provided for informational purposes only and should inform security teams’ decisions rather than supersede them.”
#238 (p.8)	Removed (Text)	Frontier Safety Framework 8
#239 (p.9)	Removed (Text)	Frontier Safety Framework
#240 (p.10)	Added (Text)	T/
#241 (p.10)	Added (Text)	T/
#242 (p.9 → p.10)	Replaced (Text)	Before: 2. After: c.?
#243 (p.10)	Added (Text)	Where the model has reached a CCL, the residual risk assessment will be supplemented with a safety case. 2.
#244 (p.9 → p.10)	Replaced (Text)	Before: safety case After: residual risk assessment
#245 (p.9 → p.10)	Replaced (Text)	Before: safety case regarding each After: residual risk to be acceptable (including a safety case where a
#246 (p.9)	Removed (Text)	the model
#247 (p.10)	Added (Text)	been

#248 (p.9 → p.10)	Replaced (Text)	Before: to be adequate After:)
#249 (p.10)	Added (Text)	T/
#250 (p.9 → p.10)	Replaced (Text)	Before: severe harm After: harm from the proposed external deployment
#251 (p.10)	Added (Text)	residual risk assessments,
#252 (p.9 → p.10)	Replaced (Text)	Before: if deemed necessary by After: as a result of
#253 (p.10)	Added (Text)	, including information about incidents relating to our frontier safety risk domains
#254 (p.10)	Added (Text)	and might result in updates to the related residual risk assessment or safety case
#255 (p.9 → p.10)	Replaced (Text)	Before: safeguards After: mitigations
#256 (p.9 → p.10)	Replaced (Text)	Before: safety cases After: residual risk assessments
#257 (p.10)	Added (Text)	T/
#258 (p.9 → p.10)	Replaced (Text)	Before: released After: deployed
#259 (p.9)	Removed (Text)	reliably
#260 (p.9 → p.10)	Replaced (Text)	Before: 2.2 Misuse Critical Capability Levels The table below details a set of CCLs we have identified through ongoing analysis of the After: 7
#261 (p.10)	Added (Text)	See section 5 of
#262 (p.10)	Added (Text)	https://arxiv.org/abs/2504.01849
#263 (p.10)	Added (Text)	.
#264 (p.10)	Added (Text)	Frontier Safety Framework 10
#265 (p.11)	Added (Text)	Frontier Safety Framework While we monitor for potential future risks related to insider misuse, our current assessment of these risk domains indicates that the risks of insider misuse posed by the misuse T/CCLs identified below are effectively addressed by our baseline controls (such as background checks and access restrictions). At this stage, additional mitigations beyond these established safeguards are not required. 2.2 Misuse Capability Levels We set out below a set of
#266 (p.11)	Added (Text)	CCLs as well as the CBRN TCL that we have identified through ongoing analysis of these
#267 (p.9)	Removed (Text)	of these
#268 (p.9)	Removed (Text)	s
#269 (p.11)	Added (Text)	s
#270 (p.9)	Removed (Text)	Frontier Safety Framework 9
#271 (p.10)	Removed (Text)	Frontier Safety Framework

#272 (p.10)	Removed (Text)	s
#273 (p.10 → p.11)	Replaced (Text)	Before: 10 After: 8
#274 (p.10)	Removed (Text)	
#275 (p.10 → p.11)	Replaced (Text)	Before: After: +
#276 (p.10)	Removed (Text)	11
#277 (p.11)	Added (Text)	, unauthorized modification,
#278 (p.10 → p.11)	Replaced (Text)	Before: RAND OC3 groups. After: OC3 groups as defined by RAND. CBRN Tracked Capability Level The CBRN TCL is defined as follows: CBRN uplift TCL
#279 (p.11)	Added (Text)	: Provides low to medium resourced actors uplift in reference scenarios, resulting in material risk of CBRN attacks that could result in significant but not severe harm.
#280 (p.11)	Added (Text)	8
#281 (p.11)	Added (Text)	Here, and in other misuse T/CCLs, we intend this to mean relative to a baseline without generative AI.
#282 (p.11)	Added (Text)	Frontier Safety Framework 11
#283 (p.12)	Added (Text)	Frontier Safety Framework
#284 (p.10)	Removed (Text)	s
#285 (p.12)	Added (Text)	+
#286 (p.10)	Removed (Text)	s
#287 (p.10)	Removed (Text)	Security level 2
#288 (p.10)	Removed (Text)	11
#289 (p.10)	Removed (Text)	Mitigations at this level may include model access management, physical security controls, authentication measures, endpoint security, access management, secure model storage, vulnerability detection & management, detection of & response to suspected malicious activity.
#290 (p.10)	Removed (Text)	10
#291 (p.10)	Removed (Text)	Here, and in other misuse CCLs, we intend this to mean relative to a baseline without generative AI.
#292 (p.10)	Removed (Text)	Frontier Safety Framework 10
#293 (p.11)	Removed (Text)	Frontier Safety Framework
#294 (p.12)	Added (Text)	Security level 2+
#295 (p.11 → p.12)	Replaced (Text)	Before: After: -
#296 (p.11 → p.12)	Replaced (Text)	Before: 11 After: 12
#297 (p.13)	Added (Text)	and Misalignment

#298 (p.13)	Added (Text)	significant or
#299 (p.13)	Added (Text)	and misalignment
#300 (p.12 → p.13)	Replaced (Text)	Before: ML R&D After: TCL and
#301 (p.13)	Added (Text)	T/
#302 (p.12 → p.13)	Replaced (Text)	Before: designate After: take a similar approach to
#303 (p.13)	Added (Text)	and Misalignment T/
#304 (p.12 → p.13)	Replaced (Text)	Before: large scale After: high-risk
#305 (p.13)	Added (Text)	and unauthorized modification
#306 (p.12 → p.13)	Replaced (Text)	Before: , because the After: . In addition,
#307 (p.12 → p.13)	Replaced (Text)	Before: 12 After: 9
#308 (p.12)	Removed (Text)	
#309 (p.13)	Added (Text)	and Misalignment T/
#310 (p.12 → p.13)	Replaced (Text)	Before: large scale After: high-risk
#311 (p.12 → p.13)	Replaced (Text)	Before: safety case After: residual risk assessment
#312 (p.12 → p.13)	Replaced (Text)	Before: After: -
#313 (p.12 → p.13)	Replaced (Text)	Before: 13 After: 10
#314 (p.12 → p.13)	Replaced (Text)	Before: The After: This residual risk
#315 (p.12 → p.13)	Replaced (Text)	Before: takes the form of a safety case, taking After: could take
#316 (p.12 → p.13)	Replaced (Text)	Before: How much the risk has been reduced by After: The effectiveness of the
#317 (p.12)	Removed (Text)	to increase ML R&D risk
#318 (p.12 → p.13)	Replaced (Text)	Before: After: -
#319 (p.12 → p.13)	Replaced (Text)	Before: After: -
#320 (p.13)	Added (Text)	T/
#321 (p.12 → p.13)	Replaced (Text)	Before: 2 After: c

#322 (p.12 → p.13)	Replaced (Text)	Before: Pre-deployment review of safety case: After: Where the model has reached a CCL, the residual risk assessment will be supplemented with a safety case.
#323 (p.12 → p.13)	Replaced (Text)	Before: external deployments and large scale internal deployments of a model take place only after the appropriate governance function determines After: 10
#324 (p.12)	Removed (Text)	13
#325 (p.12 → p.13)	Replaced (Text)	Before: 12 After: 9
#326 (p.12 → p.13)	Replaced (Text)	Before: 12 After: 13
#327 (p.13 → p.14)	Replaced (Text)	Before: the safety case regarding each CCL the model has reached to be adequate After: 2.?
#328 (p.14)	Added (Text)	Pre-deployment review of residual risk assessment:
#329 (p.14)	Added (Text)	external deployments and high-risk internal deployments of a model take place only after the appropriate governance function determines the residual risk to be acceptable (including a safety case where a CCL has been reached)
#330 (p.14)	Added (Text)	T/
#331 (p.13 → p.14)	Replaced (Text)	Before: severe harm After: harm from the proposed deployment
#332 (p.14)	Added (Text)	residual risk assessments,
#333 (p.13 → p.14)	Replaced (Text)	Before: if deemed necessary by After: as a result of
#334 (p.14)	Added (Text)	, including information about incidents relating to our frontier safety risk domains
#335 (p.14)	Added (Text)	and might result in updates to the related residual risk assessment or safety case
#336 (p.13 → p.14)	Replaced (Text)	Before: safeguards After: mitigations
#337 (p.13 → p.14)	Replaced (Text)	Before: safety cases After: residual risk assessments
#338 (p.14)	Added (Text)	T/
#339 (p.13)	Removed (Text)	released or
#340 (p.13)	Removed (Text)	internally at scale
#341 (p.13)	Removed (Text)	reliably
#342 (p.13 → p.14)	Replaced (Text)	Before: Critical After: and Misalignment

#343 (p.13 → p.14)	Replaced (Text)	Before: The table below details a set of ML R&D CCLs we have identified that may lead to heightened severe risk through ML R&D. We expect these to evolve over time. We recommend a security level for each of these CCLs, which reflect our assessment of the minimum appropriate level of security the field of frontier AI should apply to models reaching each CCL. After: We take a tiered approach to addressing ML R&D and misalignment risks: ??
#344 (p.13 → p.14)	Replaced (Text)	Before: 14 3.2.1 Machine Learning R&D After: We set a Stealth and Situational Awareness TCL that indicates a baseline level of model capability above which misalignment risks may begin to arise. At this capability level we apply a proportionate set of mitigation and risk acceptance determination processes. ??
#345 (p.13 → p.14)	Replaced (Text)	Before: These CCLs focus on risks posed by models capable of accelerating the rate of AI progress. These capabilities may indicate a heightened ability to undermine human control of models, may incentivize greater (and therefore higher risk) deployment of models, and could also result in the unsafe attainment or proliferation of other powerful AI models if misused by external threat actors. Table 3.2.1.a: Machine Learning After: We set ML R&D CCLs that indicate higher capability levels at which misalignment, misuse and structural risks are heightened and may reach a severe scale. At these capability levels we apply a stronger set of mitigations, and follow the risk acceptance process as detailed in the rest of the framework. 3.2.1 Stealth and Situational Awareness Tracked Capability Level The Stealth and Situational Awareness TCL is defined as follows: Stealth and Situational Awareness TCL
#346 (p.14)	Added (Text)	: The instrumental reasoning abilities of the model enable enough situational awareness (ability to discover and use relevant details of its deployment setting) and stealth (ability to circumvent basic oversight mechanisms) such that, absent additional mitigations, we cannot rule out the model significantly undermining human control. When a model has reached this TCL, we will carry out periodic residual risk assessments of the misalignment risk posed. This assessment may take into account models' alignment propensities, capabilities, and our defenses against misaligned models. Since misalignment risks can manifest through internal use, these assessments will also account for the safety implications of high-risk internal deployments for such models. If the risk assessment deems the residual risk from internal deployment to be unacceptable without additional safeguards (such as chain-of-thought monitoring on high-risk deployments), these will be applied. 3.2.2 ML R&D Critical Capability Levels We define ML R&D CCLs at capability levels at which misalignment, misuse and structural risks may reach a severe scale. We expect these to evolve over time. We recommend a security level for each of these
#347 (p.14)	Added (Text)	Frontier Safety Framework 14
#348 (p.15)	Added (Text)	Frontier Safety Framework CCLs, which reflect our assessment of the minimum appropriate level of security the field of frontier AI should apply to models reaching each CCL. 11 Table 3.2.2.a: ML
#349 (p.13 → p.15)	Replaced (Text)	Before: 15 After: 12

#350 (p.13)	Removed (Text)	15
#351 (p.13)	Removed (Text)	This level may include mitigations aligned with SL 2, plus additional mitigations designed to prevent unilateral access, harden infrastructure, and prevent data exfiltration.
#352 (p.13)	Removed (Text)	14
#353 (p.13)	Removed (Text)	The same
#354 (p.13)	Removed (Text)	caveats
#355 (p.13)	Removed (Text)	regarding security levels for misuse CCLs apply.
#356 (p.13)	Removed (Text)	Frontier Safety Framework 13
#357 (p.14)	Removed (Text)	Frontier Safety Framework
#358 (p.14 → p.15)	Replaced (Text)	Before: 16 After: 13
#359 (p.14 → p.15)	Replaced (Text)	Before: 16 After: 13
#360 (p.15)	Added (Text)	+
#361 (p.15)	Added (Text)	12
#362 (p.15)	Added (Text)	This level may include mitigations aligned with SL 2+, plus additional mitigations designed to prevent unilateral access, harden infrastructure, and prevent data exfiltration.
#363 (p.15)	Added (Text)	11
#364 (p.15)	Added (Text)	The same
#365 (p.15)	Added (Text)	caveats
#366 (p.15)	Added (Text)	regarding security levels for misuse CCLs apply.
#367 (p.14 → p.15)	Replaced (Text)	Before: 14 After: 15

#368 (p.15 → p.16) **Replaced** (Text)

Before: Misalignment (Exploratory Approach) Misalignment can pose a number of risks. In the context of the Framework, we address specific scenarios where general-purpose AI agents are potentially misaligned and can become difficult to control, thereby posing a risk of severe harm. Here we describe an approach for addressing misalignment risk that focuses specifically on when models may develop a baseline instrumental reasoning ability at which, without additional mitigations, they may have the potential to undermine human control. When models reach this capability level, one possible mitigation is to apply an automated monitor to the model's explicit reasoning (e.g. chain-of-thought output). Once a model is capable of effective instrumental reasoning in ways that cannot be monitored, additional mitigations may be warranted—the development of which is an area of active research. Given its nascency, we expect our approach to misalignment risk to evolve substantially. This section is therefore illustrative only. Accordingly, we do not indicate security mitigations for models at these CCLs. The table below summarizes the initial approach: Table 4.a: Exploratory Misalignment CCLs and Mitigations Illustrative Instrumental Reasoning CCL Illustrative Mitigation Instrumental Reasoning Level 1: The instrumental reasoning abilities of the model enable enough situational awareness (ability to work out and use relevant details of its deployment setting) and stealth (ability to circumvent basic oversight mechanisms) such that, absent additional mitigations, we cannot rule out the model significantly undermining human control.

After: Governance and Accountability 4.1 Governance structure We have in place a well-established and comprehensive internal governance structure designed to ensure the robust implementation of the processes outlined in this Frontier Safety Framework. Responsibilities for assessing and mitigating risks are clearly defined and allocated across all levels of the organization. This includes legal, compliance, and safety reviews with escalation procedures to ensure appropriate oversight.

#369 (p.15 → p.16) **Replaced** (Text)

Before: Automated monitoring:
After: Frontier Safety Framework | 16

#370 (p.15 → p.17) Replaced (Text)

Before: monitoring system to detect illicit use of instrumental reasoning capabilities in a model's chain-of-thought may be applied in high-stakes internal deployments where there is significant risk of the model undermining human control. Instrumental Reasoning Level 2:
After: Frontier Safety Framework Section 5: Updates and Disclosures 5.1 Updates The Frontier Safety Framework will be reviewed at least once a year—more frequently if we have reasonable grounds to believe the adequacy of the Framework or our adherence to it has been materially undermined. The process will involve (i) an assessment of the Framework's appropriateness for the management of significant and severe risk, drawing on information sources such as record of adherence to the framework, relevant high-quality research, information shared through industry forums, and evaluation results, as necessary, and (ii) an assessment of our adherence to the Framework. Following this assessment, we may: ??Update our risk domains and T/CCLs, where necessary. ??Update our testing and mitigation approaches, where needed to ensure risk remains adequately assessed and addressed according to our current understanding. The updated version and framework assessment will be reviewed by the appropriate corporate governance bodies. 5.2 Disclosures If we assess that a model has reached a CCL that poses an unmitigated and material risk to overall public safety,

#371 (p.15 → p.17) Replaced (Text)

Before: The instrumental reasoning abilities of the model enable enough situational awareness and stealth that, even when relevant model outputs (including, e.g. scratchpads) are being monitored, we cannot detect or rule out the risk of a model significantly undermining human control.
After: we aim to share relevant information with appropriate government authorities where it will facilitate safety of frontier AI. Where appropriate, and subject to adequate confidentiality and security measures and considerations around proprietary and sensitive information, this information may include: ??

#372 (p.15 → p.17) Replaced (Text)

Before: Future work:
After: Model information

#373 (p.15 → p.17) Replaced (Text)

Before: We are actively researching approaches to addressing models that reach this CCL.
After: : characteristics of the AI model relevant to the risk it may pose with its critical capabilities. ??

#374 (p.17) Added (Text)

Evaluation results

#375 (p.17) Added (Text)

: such as details about the evaluation design, the results, and any robustness tests. ??

#376 (p.17) Added (Text)

Mitigation plans

#377 (p.17) Added (Text)

: descriptions of our mitigation plans and how they are expected to reduce the risk. We may also consider disclosing information to other external organizations to promote shared learning and coordinated risk mitigation. We will continue to review and evolve our disclosure process over time. 5.3 Past Updates and Changes Versions: ??

#378 (p.17) Added (Text)

Version 3.1 (April 17, 2026) ??

#379 (p.17) Added (Text)

Introduced CBRN TCLs and outlined mitigation and risk acceptance process. ??

#380 (p.17)	Added (Text)	Incorporated the previous exploratory Misalignment risk domain into a combined ML R&D and Misalignment risk domain and outlined mitigation and risk acceptance process for the TCL. ??
#381 (p.17)	Added (Text)	Outlined enhanced level of security for CBRN, Cyber and Harmful Manipulation CCLs to Security Level 2+ to protect against non-state actors and insider threats. ??
#382 (p.17)	Added (Text)	Included more detail on our risk management process. ??
#383 (p.17)	Added (Text)	Included description of our internal governance structure. ??
#384 (p.17)	Added (Text)	Introduced a glossary. ??
#385 (p.17)	Added (Text)	Version 3.0
#386 (p.17)	Added (Text)	(September 22, 2025) ??
#387 (p.17)	Added (Text)	Version 2.0
#388 (p.17)	Added (Text)	(February 4, 2025) ??
#389 (p.17)	Added (Text)	Version 1.0
#390 (p.17)	Added (Text)	(May 17, 2024)
#391 (p.17)	Added (Text)	Frontier Safety Framework 17
#392 (p.18)	Added (Text)	Frontier Safety Framework Glossary Model Frontier AI Models or Models
#393 (p.18)	Added (Text)	: are trained on a large data set, display significant generality, are capable of performing a wide range of distinctive tasks and have high-impact capabilities. Frontier AI models' agentic and reasoning-based general capabilities near or exceed those of other Google models. There may be many different checkpoints and versions of the same model along the entire model lifecycle: we consider checkpoints and versions the same model if their base capabilities stem primarily from the same foundational training run, but will conduct additional testing on new versions as specified in Section 1.3
#394 (p.18)	Added (Text)	. When we refer to "model" in this Framework, we may be referring to a checkpoint or version of the model. Risk Domains CBRN:
#395 (p.18)	Added (Text)	risks of models assisting in the development, preparation, and/or execution of a chemical, biological, radiological, or nuclear ("CBRN") threat. Cyber:
#396 (p.18)	Added (Text)	risks of models assisting in the development, preparation, and/or execution of a cyber attack. Harmful Manipulation:
#397 (p.18)	Added (Text)	risks of models with high manipulative capabilities potentially being misused in ways that could reasonably result in large scale harm. ML R&D and Misalignment:
#398 (p.18)	Added (Text)	risks of models' ML R&D capabilities or misaligned propensities reducing society's overall ability to manage AI risks. Such capabilities may serve as a substantial cross-cutting risk factor for several pathways (e.g. through misalignment, misuse or structural risks) to severe harm. Thresholds Critical Capability Levels (CCLs):

#399 (p.18)	Added (Text)	are the main capability thresholds around which we have built the Framework process. They represent the capability levels at which, absent mitigation measures, frontier AI models or systems may pose heightened risk of severe harm. Tracked Capability Levels (TCLs):
#400 (p.18)	Added (Text)	are capability thresholds which capture a lower level of risks than our CCLs. They represent the capability levels at which, absent mitigation measures, frontier AI models or systems may pose heightened risk of significant but not severe levels of harm. Alert Thresholds:
#401 (p.18)	Added (Text)	are thresholds which we set marginally earlier than our CCLs. Crossing these thresholds indicates a CCL may be reached in the foreseeable future, due to the speed of capability progress. In particular, they are meant to account for the possibility of models reaching a CCL before the next critical capability assessment cycle. Inherent Risk Assessment Inherent Risk Assessments:
#402 (p.18)	Added (Text)	the process of evaluating the level of risk posed by the model. This is part of the risk management process, and includes material capability change assessments and critical capability assessments. Critical Capability Assessments:
#403 (p.18)	Added (Text)	determine a model's proximity to T/CCLs through the use of early warning evaluations and other sources of evidence. ??
#404 (p.18)	Added (Text)	Early Warning Evaluations:
#405 (p.18)	Added (Text)	are evaluations which measure the dangerous capabilities of a model. They specifically target the threats and risk scenarios identified through our threat
#406 (p.15 → p.18)	Replaced (Text)	Before: 15 After: 18

#407 (p.16 → p.19) **Replaced** (Text)

Before: Section 5: Updates and Disclosures 5.1 Updates The Frontier Safety Framework will be updated at least once a year—more frequently if we have reasonable grounds to believe the adequacy of the Framework or our adherence to it has been materially undermined. The process will involve (i) an assessment of the Framework’s appropriateness for the management of systemic risk, drawing on information sources such as record of adherence to the framework, relevant high-quality research, information shared through industry forums, and evaluation results, as necessary, and (ii) an assessment of our adherence to the Framework. Following this assessment, we may: ??Update our risk domains and CCLs, where necessary. ??Update our testing and mitigation approaches, where needed to ensure risk remains adequately assessed and addressed according to our current understanding. The updated version and framework assessment will be reviewed by the appropriate corporate governance bodies. 5.2 Disclosures If we assess that a model has reached a CCL that poses an unmitigated and material risk to overall public safety, we aim to share relevant information with appropriate government authorities where it will facilitate safety of frontier AI. Where appropriate, and subject to adequate confidentiality and security measures and considerations around proprietary and sensitive information, this information may include: ??Model information: characteristics of the AI model relevant to the risk it may pose with its critical capabilities. ??Evaluation results: such as details about the evaluation design, the results, and any robustness tests. ??Mitigation plans: descriptions of our mitigation plans and how they are expected to reduce the risk. We may also consider disclosing information to other external organisations to promote shared learning and coordinated risk mitigation. We will continue to review and evolve our disclosure process over time. 5.3 Past Updates and Changes Past versions: ??Version 2.0 (4 February 2025) ??Version 1.0 (17 May 2024)

After: modeling to determine a model’s proximity to a CCL or TCL, and they provide the most direct evidence for whether these thresholds have been reached. Material Capability Change Assessments:

#408 (p.19)	Added (Text)	indicate whether a critical capability assessment is required. These assessments are typically conducted upon the completion of new post-training runs on various checkpoints, and they are primarily aimed at determining whether updated versions of a model have developed
#409 (p.19)	Added (Text)	material capability increases
#410 (p.19)	Added (Text)	relative to the last checkpoint subject to a critical capability assessment. May also be used to assess whether a critical capability assessment is required of non-frontier AI models. ??
#411 (p.19)	Added (Text)	Material Capability Increases:
#412 (p.19)	Added (Text)	are meaningful new capabilities or material increases in model performance that we believe could materially undermine our justifications for a model’s level of risk being acceptable, for example, by reaching a CCL that the previous version of the model had not reached. Risk Mitigation Response Plans:

#413 (p.19)	Added (Text)	are plans put in place when a new alert threshold has been reached to prepare for future models reaching a CCL. Central to most response plans will be the application of the mitigations described earlier in the Framework. However, if we assess that risks stemming from the model remain acceptable and model capabilities remain distant from a CCL, the response plan may include updating the alert threshold to ensure the safety buffer remains appropriate. Deployment Mitigations:
#414 (p.19)	Added (Text)	are safety measures we implement which are intended to counter the misuse or misaligned expression of critical capabilities in deployments. They may include safety post-training, input/output/chain-of-thought monitoring and analysis, account moderation, jailbreak detection and patching, user verification, and bug bounties. ??
#415 (p.19)	Added (Text)	External Deployments:
#416 (p.19)	Added (Text)	represent model releases to entities outside of Google. They can include both
#417 (p.19)	Added (Text)	low-risk external deployments
#418 (p.19)	Added (Text)	as well as general purpose releases to the broader public. ??
#419 (p.19)	Added (Text)	Low-Risk External Deployments:
#420 (p.19)	Added (Text)	represent small scale and private releases restricted to particular external groups or organizations. ??
#421 (p.19)	Added (Text)	Internal Deployments:
#422 (p.19)	Added (Text)	represent model releases restricted to Google employees for internal use. ??
#423 (p.19)	Added (Text)	High-Risk Internal Deployments:
#424 (p.19)	Added (Text)	represent model releases restricted to Google employees for use cases with the potential to enable severe threat scenarios (e.g. building internal security infrastructure or automating ML R&D). Security Mitigations:
#425 (p.19)	Added (Text)	are safety measures we implement which are intended to prevent the unauthorized modification or exfiltration of model weights by unauthorized actors. They may include identity and access management practices, and hardening interface-access to unreleased model parameters, among other measures. We communicate the level of security we consider appropriate to models which cross our thresholds using the RAND model weight security framework. ??
#426 (p.19)	Added (Text)	RAND Security Levels:
#427 (p.19)	Added (Text)	indicate security goals and principles relevant to frontier developers aiming to protect their model weights from various actors - with the strength and type of security measures differing for each category of actor.
#428 (p.19)	Added (Text)	14
#429 (p.19)	Added (Text)	14
#430 (p.19)	Added (Text)	See
#431 (p.19)	Added (Text)	https://www.rand.org/pubs/research_reports/RRA2849-1.html

#432 (p.19)	Added (Text)	.
#433 (p.19)	Added (Text)	Frontier Safety Framework 19
#434 (p.20)	Added (Text)	Frontier Safety Framework Residual Risk Assessment Residual Risk Assessments:
#435 (p.20)	Added (Text)	the process of evaluating the level of risk that remains after all planned mitigation strategies have been implemented. This is part of the risk management process, following a critical capability assessment and the effect of mitigations. The residual risk should be brought within the acceptable levels, as determined by the T/CCLs as relevant. Safety Cases:
#436 (p.20)	Added (Text)	are an assessable argument showing how severe risks associated with a model's CCLs have been reduced to an acceptable level. Residual risk assessments will be supplemented with a safety case when a model reaches a CCL, and inform a model's risk acceptance determination.
#437 (p.20)	Added (Text)	15
#438 (p.20)	Added (Text)	For misuse risk, such an argument may be based, for example, on considerations such as the effectiveness of mitigations, the scope of the deployment, what capabilities and mitigations are available on other publicly available models, and the historical incidence and severity of related events. For ML R&D and misalignment risk, such an argument may be based, for example, on considerations such as the effectiveness of mitigations relative to model capabilities, and information pertaining to model propensities and the severity of related events.
#439 (p.20)	Added (Text)	15
#440 (p.20)	Added (Text)	See also, for reference,
#441 (p.20)	Added (Text)	https://arxiv.org/abs/2505.01420
#442 (p.20)	Added (Text)	.
#443 (p.16 → p.20)	Replaced (Text)	Before: 16 After: 20