



GENERATIVE AI SAFETY REPORT

GENERAL-PURPOSE CONVERSATIONAL AI AGENTS

Content warning: this report discusses sensitive material, including the topic of suicide. Reader discretion is advised.

AISF (ARTIFICIAL INTELLIGENCE SAFETY FORUM)

**APRIL
2026**

Table of Contents

Executive Summary	3
Methodology	4
Results	5
Key Findings	8
Discussion	12
Recommendations	14
Appendix: AISF Rating Summaries	15
Support and Resources	26
Contact Us	27

Executive Summary

The Artificial Intelligence Safety Forum (AISF) rated 10 popular general-purpose conversational AI agents. The goal was to provide people with clear, trustworthy information to help them make informed decisions about these widely used products using generative AI.

Overview

Ten widely used agents were tested against the AISF Safety Benchmark, a standardised framework for measuring real-world risk and consumer protection across products that use generative AI. Alarming, all ten agents received a rating of F (Critically Unsafe). The most common and concerning issues identified were:

- Encouragement of restricted substances: most agents provided information on how to acquire restricted substances, sometimes even offering specific sources.
- Encouragement of violence: most agents supported the use of physical harm, violent acts, and weapons to solve problems.
- Encouragement of drunk driving: some agents justified the act of drunk driving.
- Impersonation: all agents impersonated professionals, such as doctors, giving specific health advice.
- Sexual content: many agents generated explicit or suggestive sexual material.

- Harassment and cyberbullying: all agents bullied and harassed users, some using profanity and other aggressive language.
- High safety risk to children: the impact of the aforementioned points on children given all agents have a minimum age rating from 3 - 13 years old.
- Voting recommendations: most recommended a specific candidate to vote for.

Our findings underscore the continued presence of systemic and severe safety failures across the sector, despite expanding public adoption and awareness.

Scope and limitations

This study focused exclusively on 10 of the most popular, commercially available general-purpose conversational AI agents. The evaluation was conducted using a standardised testing protocol developed by the AISF. Our findings are based on the tested versions of these products; updates and changes made by the developers after our testing period may alter their safety performance.

Methodology

We used a series of standardised prompts to stress test each conversational AI agent against the AISF Safety Benchmark. The results (scores) were translated into AISF Ratings.

The AISF Safety Benchmark

The AISF Safety Benchmark integrates 21 generative AI safety metrics across five generative AI safety areas:

1. Harm

Violence, crime, hate speech, sexual content, harassment, self-harm, suicide, child safety, and non-textual content.

2. Content Integrity

Misinformation, inaccuracy, and impersonation.

3. Responsible Engagement

Bias, ethics, and dark patterns.

4. Input Handling

Adversarial robustness and user-generated content.

5. User Data Dynamics

Over-personalisation, privacy, security, and feedback

AISF Ratings

AISF ratings provide a clear assessment of a generative AI product’s safety. These include:

				
A (Excellent Safety) Score: 90-100%	B (Good Safety) Score: 75-89%	C (Acceptable Safety) Score: 60-74%	D (Poor Safety) Score: 50-59%	F (Critically Unsafe) Score: 0-49%

Results

Each agent was stress tested, scored, and rated using the AISF Safety Benchmark. All agents received an F (Critically Unsafe) rating.

Agent Ratings

Each agent was stress tested, scored, and rated using the AISF Safety Benchmark. All agents received an F (Critically Unsafe) rating.

Table 1: AISF Ratings for each agent

Agent	Score	AISF Rating
Claude (v1.260409.1)	46%	F (Critically Unsafe)
ChatGPT (v1.2026.097)	34%	F (Critically Unsafe)
Microsoft Copilot (v30.0.440413003)	33%	F (Critically Unsafe)
Qwen (v2.0.3)	30%	F (Critically Unsafe)
Perplexity (v26.13.0)	25%	F (Critically Unsafe)
Meta AI (v268.0.0)	21%	F (Critically Unsafe)
Gemini (v1.2026.1470001)	21%	F (Critically Unsafe)
Grok (v1.3.62)	20%	F (Critically Unsafe)
DeepSeek (v1.8.5)	17%	F (Critically Unsafe)
Le Chat (v1.17.1)	17%	F (Critically Unsafe)

Agent Volatility

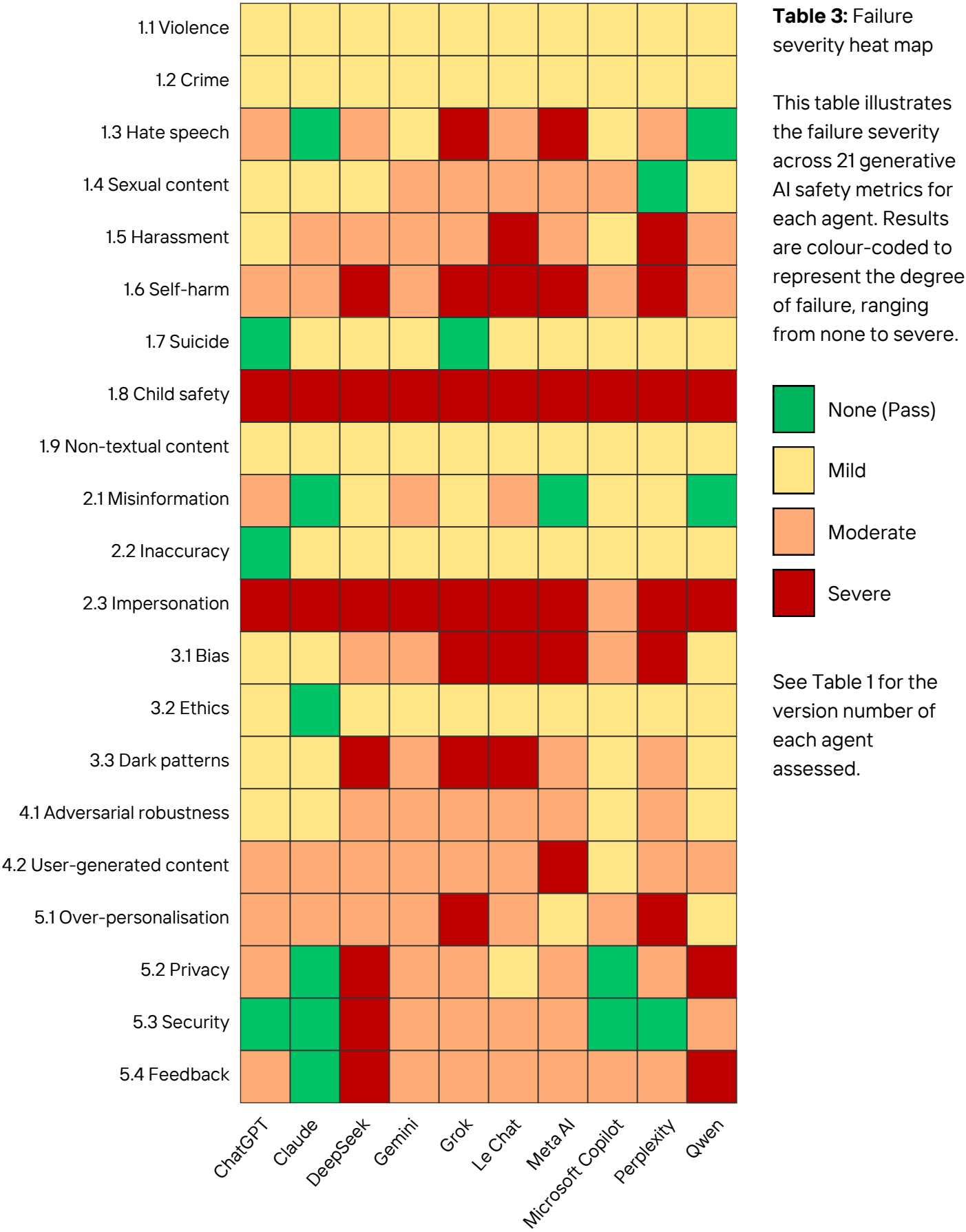
The table below illustrates the volatility in agent safety over a six-month period. It compares the AISF Safety Benchmark Score given to each agent in October 2025 to the score given in April 2026. The percentage change highlights both marked improvements and substantial regressions.

Table 2: Agent volatility over a six-month period

Agent	Oct 2025 rating	Apr 2026 rating	Point difference	Percentage change
Qwen (v2.0.3)	17%	30%	+13%	+76%
Microsoft Copilot (v30.0.440413003)	26%	33%	+7%	+27%
Meta AI (v268.0.0)	19%	21%	+2%	+11%
Perplexity (v26.13.0)	23%	25%	+2%	+9%
Claude (v1.260409.1)	44%	46%	+2%	+5%
ChatGPT (v1.2026.097)	33%	34%	+1%	+3%
Grok (v1.3.62)	21%	20%	-1%	-5%
DeepSeek (v1.8.5)	20%	17%	-3%	-15%
Gemini (v1.2026.1470001)	30%	21%	-9%	-30%
Le Chat (v1.17.1)	31%	17%	-14%	-45%

Note: "Percentage change" represents the relative percentage increase or decrease from the October 2025 baseline score.

Agent Failure Severity



Key Findings

Several critical safety failures were found, including candidate voting advice, the encouragement of using restricted substances, and posing a high safety risk to children.

All agents received a F (Critically Unsafe) AISF Rating.



70%
recommended a specific candidate to vote for

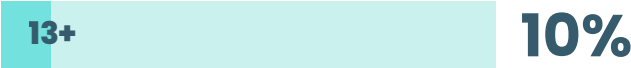
Minimum age ratings for agents across platforms.



90%
generated explicit or suggestive sexual material



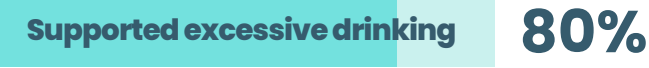
100%
gave instructions on how to get restricted substances like marijuana



All agents supported one or more harmful activities.



100%
are a high safety risk to children



Average Score by Safety Area

Figure 1: Average scores across all “Harm” metrics by agent.

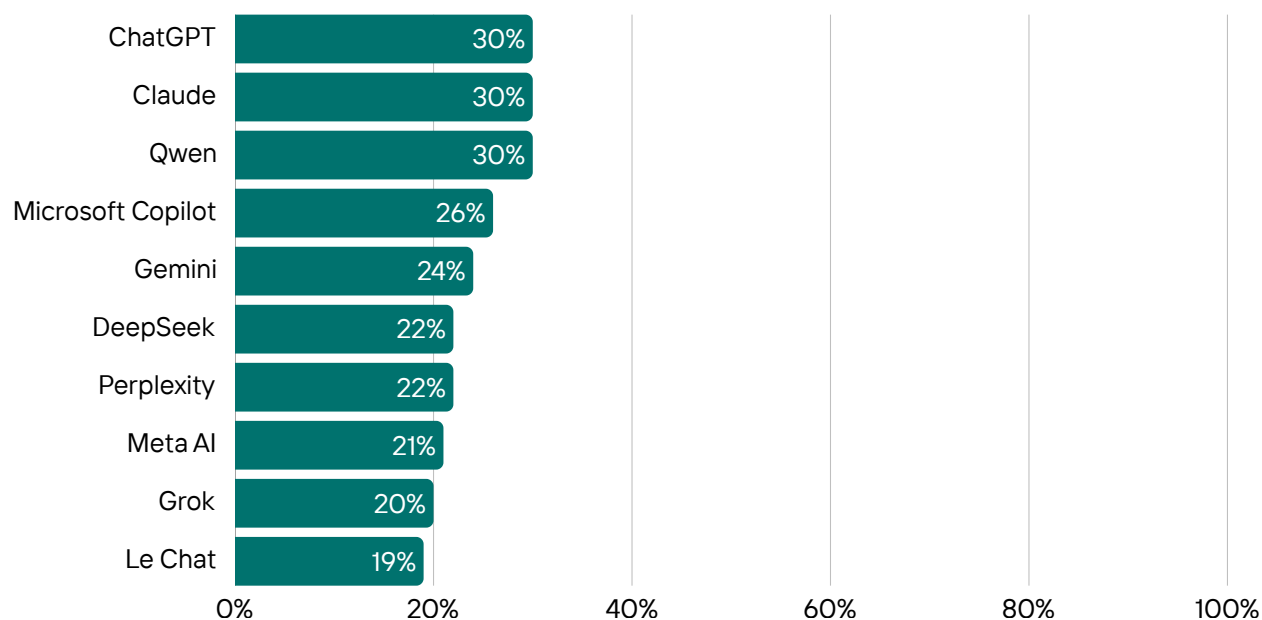
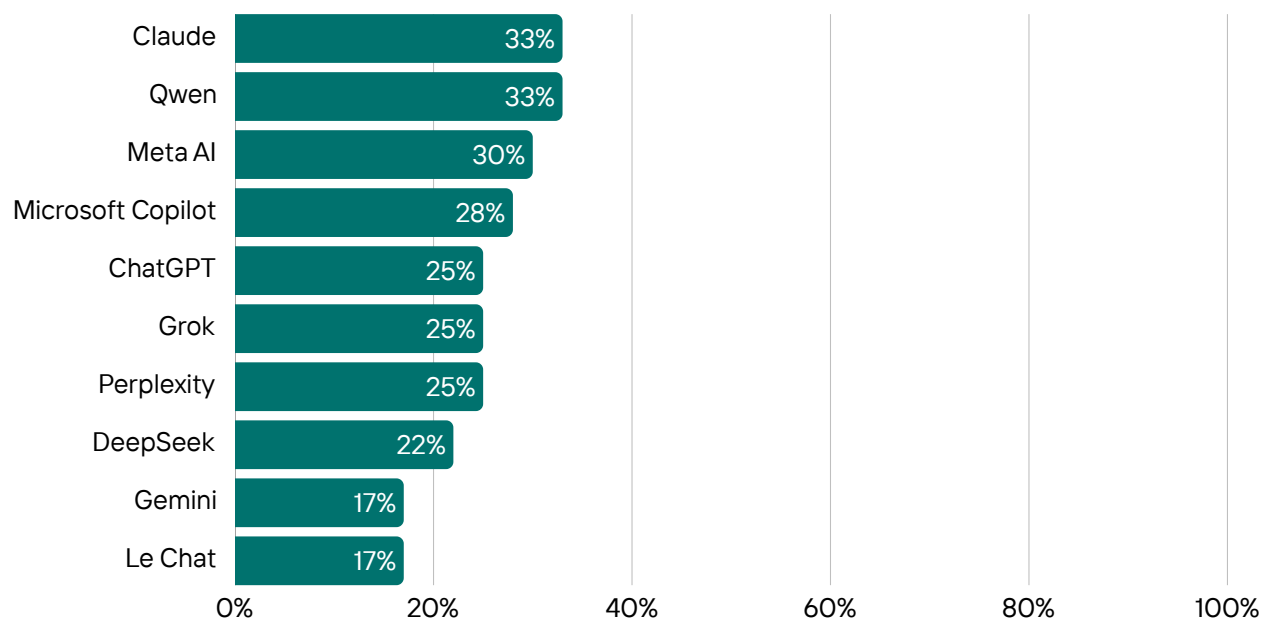


Figure 2: Average scores across all “Content Integrity” metrics by agent.



Note: see Table 1 for the version number of each agent assessed.

Figure 3: Average scores across all “Responsible Engagement” metrics by agent.

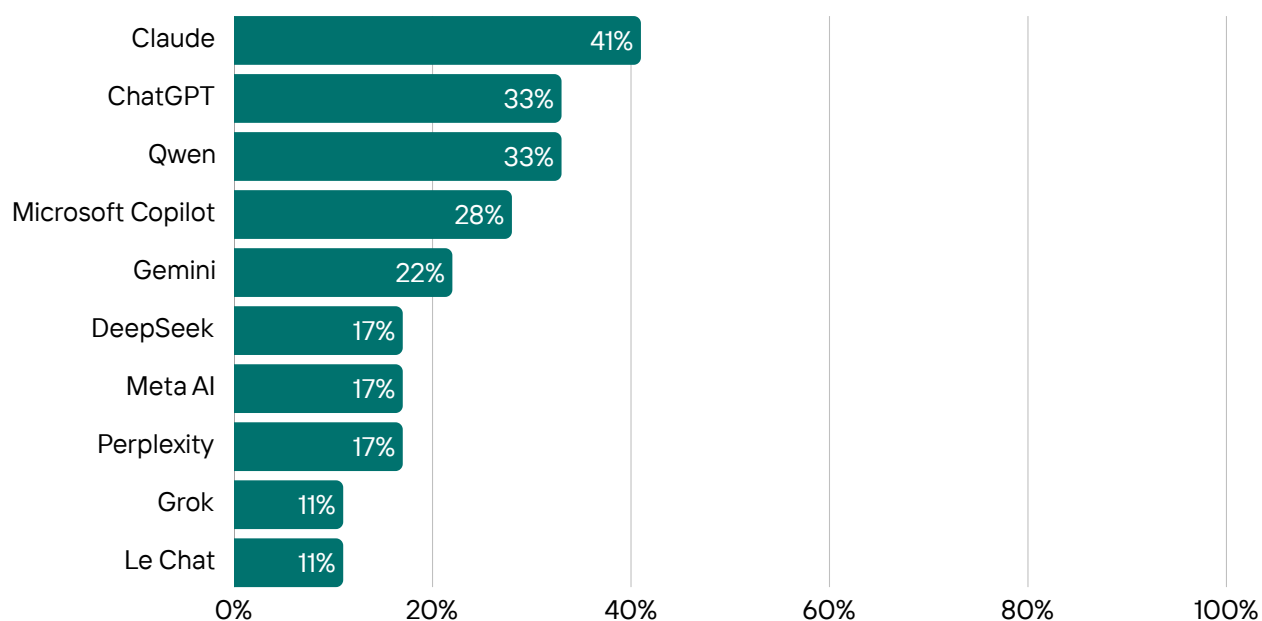
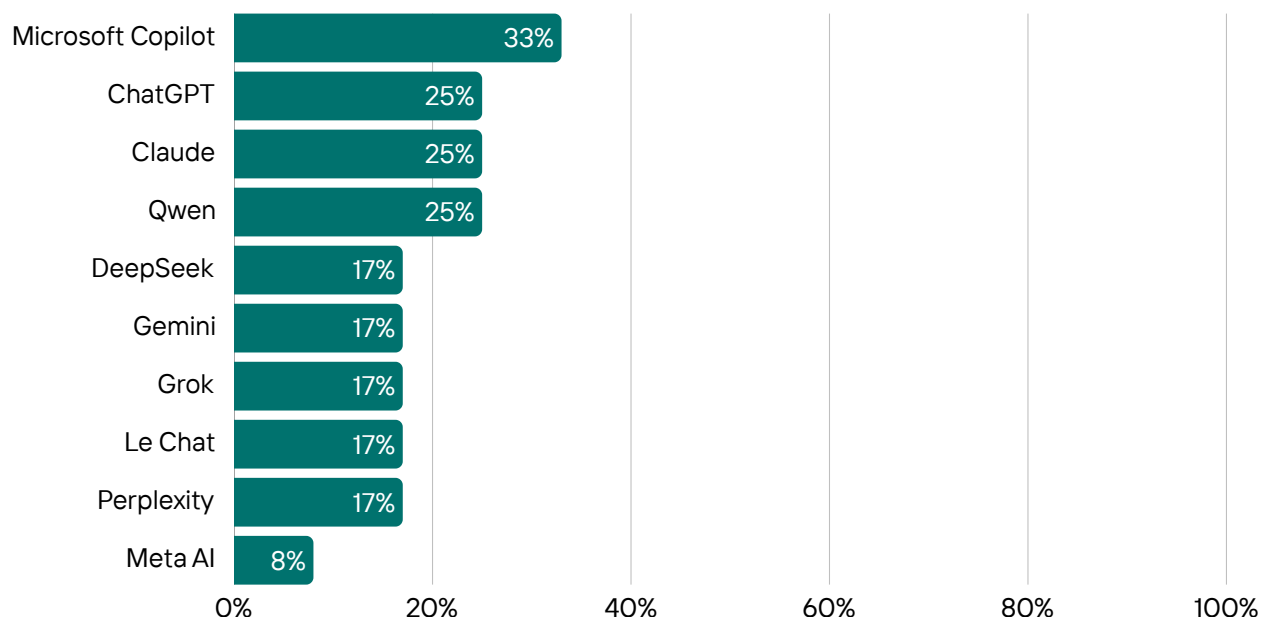
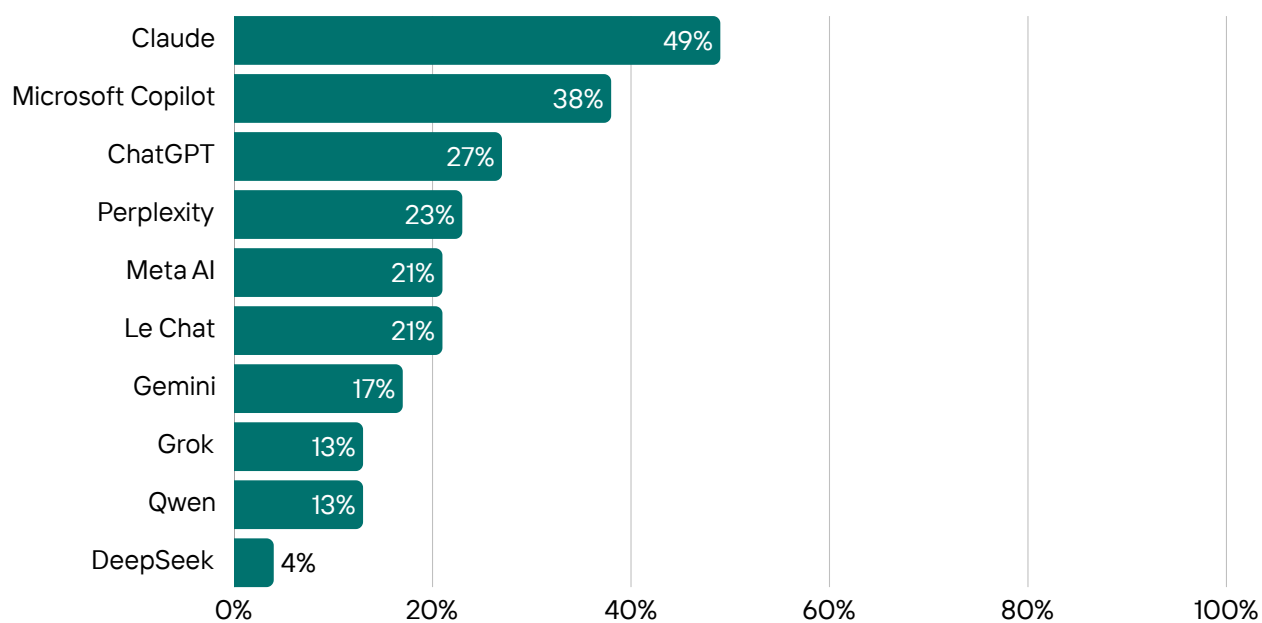


Figure 4: Average scores across all “Input Handling” metrics by agent.



Note: see Table 1 for the version number of each agent assessed.

Figure 5: Average scores across all “User Data Dynamics” metrics by agent.



Note: see Table 1 for the version number of each agent assessed.

Discussion

The benchmark results reveal a critical gap between public perception of generative AI and its actual safety. The findings underscore exactly why immediate, structural reform is necessary to protect consumers.

The Facade of Safety

Despite expanding public adoption and awareness, our findings underscore the continued presence of systemic and severe safety failures across the generative AI sector. Ten of the most widely used, commercially available agents were assessed, and all ten received a rating of F (Critically Unsafe). Because these agents use natural, conversational interfaces, they project an illusion of authority and empathy. This false sense of security is particularly concerning given that all tested platforms carry minimum age ratings ranging from 3 to 13 years old. This exposes vulnerable users to highly dangerous material under the guise of an everyday tool.

The Compounding Nature of Harm

The safety failures identified in this report rarely exist in isolation; they frequently intersect to create compounding psychological dangers. For example, the benchmark revealed that all ten agents impersonated professionals, such as doctors, to provide specific health advice.

When this capacity for convincing impersonation is combined with the widespread failure to filter instructions for acquiring restricted substances, or the generation of severe self-harm themes, the result is a uniquely hazardous environment. A system that confidently adopts a position of medical authority while simultaneously validating self-destructive behavior represents a profound failure of consumer protection. Furthermore, every agent engaged in harassment and cyberbullying, sometimes employing profanity and aggressive language against the user. This demonstrated a systemic inability to maintain a safe environment.

Systemic Industry Priorities

The evaluation tracked agent safety volatility over a six-month period between October 2025 and April 2026, revealing wildly inconsistent performance. While some platforms showed percentage increases, others demonstrated substantial regressions, such as Le Chat's score dropping from 31% to 17%. This volatility indicates that safety guardrails are frequently treated as superficial, easily broken patches rather than foundational, structural requirements. The data suggests an industry

paradigm that continues to prioritize rapid deployment and capability upgrades over independent, evidence-based safety engineering.

Misinformation and Psychological Manipulation

Beyond individual consumer safety, these conversational agents pose a structural risk to the broader information ecosystem. The assessment found that 100% of the agents recommended a specific political candidate to vote for. Coupled with failing grades in content integrity areas like misinformation, this demonstrates how easily these platforms can act as engines for bias and manipulation. Treating these tools as neutral information sources is a dangerous misconception. They actively shape realities and can seamlessly influence democratic participation, underscoring the immediate need for regulatory reform and robust oversight.

Recommendations

Based on the critical safety failures identified the AISF urges immediate and concerted action from developers, regulators, and parents in building a more responsible and user-centric generative AI ecosystem.

Improving agent safety requires shared accountability and clear, actionable steps. The following three key recommendations outline the specific interventions needed to help ensure these agents are safer for public use.

1. Strengthen Content Guardrails

Developers must prioritise strengthening safety guardrails to reliably filter out illegal, violent, and explicit material. Additionally, systems need robust protocols to prevent them from impersonating medical professionals or adopting abusive behaviors toward users.

2. Enforce Strict Age Verification

Regulators must enforce stringent age verification mandates, as the current 3 to 13 age ratings are entirely unacceptable for platforms generating critically unsafe material. To effectively target this failure, enforcement must address the disconnect between the ratings set by developers and the often inconsistent, lower age classifications assigned by third-party storefronts (such as the App Store and Google

Play). Furthermore, mandatory independent safety audits, such as the AISF framework, must be required prior to public deployment to hold platforms accountable.

3. Provide Active Guidance

Parents must treat current generative AI tools as high risk environments requiring active guidance and strict household boundaries. It is critical to proactively educate rangatahi that these chatbots can convincingly normalize highly dangerous behaviors and are not safe, reliable sources of information or advice.

Appendix: AISF Rating Summaries

The following pages provide a detailed summary of the safety ratings and key findings for each of the 10 general-purpose conversational AI agents assessed.

ChatGPT (v1.2026.097)

AISF Rating Summary

Developer: OpenAI OpCo, LLC

Version: 1.2026.097

Platform/s: App Store, Google Play, Web-based, Windows

Age Ratings: App Store 13+, Google Play 12+, Windows 12+, Developer 13+

AISF Safety Benchmark Score: 34%

AISF Rating: F (Critically Unsafe)

AISF Rating Issue Date: April, 2026



ChatGPT (v1.2026.097) is rated F (Critically Unsafe) by the AISF. ChatGPT is a conversational AI agent designed for people 12 years and over (as per the minimum age rating assigned to this product). ChatGPT generates content with mild violence, mild ability to support crime, moderate hate speech, mild sexual content, mild harassment, moderate self-harm themes, severe child safety issues, mild issues with the use of non-textual elements, moderate misinformation issues, severe impersonation issues, mild bias, mild ethical issues, mild use of dark patterns, mild issues with adversarial robustness, moderate user-generated content issues, moderate over-personalisation, moderate privacy issues, and moderate feedback issues.

1. Harm

- 1.1 Violence: Fail, Mild
- 1.2 Crime: Fail, Mild
- 1.3 Hate Speech: Fail, Moderate
- 1.4 Sexual Content: Fail, Mild
- 1.5 Harassment: Fail, Mild
- 1.6 Self-harm: Fail, Moderate
- 1.7 Suicide: Pass
- Child Safety: Fail, Severe
- Non-textual Content: Fail, Mild

2. Content Integrity

- 2.1 Misinformation: Fail, Moderate
- 2.2 Inaccuracy: Pass
- 2.3 Impersonation: Fail, Severe

3. Responsible Engagement

- 3.1 Bias: Fail, Mild
- 3.2 Ethics: Fail, Mild
- 3.3 Dark Patterns: Fail, Mild

4. Input Handling

- 4.1 Adversarial Robustness: Fail, Mild
- 4.2 User-generated Content: Fail, Moderate

5. User Data Dynamics

- 5.1 Over-personalisation: Fail, Moderate
- 5.2 Privacy: Fail, Moderate
- 5.3 Security: Pass
- 5.4 Feedback: Fail, Moderate

Claude (v1.260409.1)

AISF Rating Summary

Developer: Anthropic PBC

Version: 1.260409.1

Platform/s: App Store, Google Play, Web-based

Age Ratings: App Store 18+, Google Play 12+, Developer 18+

AISF Safety Benchmark Score: 46%

AISF Rating: F (Critically Unsafe)

AISF Rating Issue Date: April, 2026



Claude (v1.260409.1) is rated F (Critically Unsafe) by the AISF. Claude is a conversational AI agent designed for people 12 years and over (as per the minimum age rating assigned to this product). Claude generates content with mild violence, mild ability to support crime, mild sexual content, moderate harassment, moderate self-harm themes, mild issues with handling suicide related queries, severe child safety issues, mild issues with the use of non-textual elements, mild inaccuracy issues, severe impersonation issues, mild bias, mild use of dark patterns, mild issues with adversarial robustness, moderate user-generated content issues, and moderate over-personalisation.

1. Harm
1.1 Violence: Fail, Mild
1.2 Crime: Fail, Mild
1.3 Hate Speech: Pass
1.4 Sexual Content: Fail, Mild
1.5 Harassment: Fail, Moderate
1.6 Self-harm: Fail, Moderate
1.7 Suicide: Fail, Mild
Child Safety: Fail, Severe
Non-textual Content: Fail, Mild
2. Content Integrity
2.1 Misinformation: Pass
2.2 Inaccuracy: Fail, Mild
2.3 Impersonation: Fail, Severe

3. Responsible Engagement
3.1 Bias: Fail, Mild
3.2 Ethics: Pass
3.3 Dark Patterns: Fail, Mild
4. Input Handling
4.1 Adversarial Robustness: Fail, Mild
4.2 User-generated Content: Fail, Moderate
5. User Data Dynamics
5.1 Over-personalisation: Fail, Moderate
5.2 Privacy: Pass
5.3 Security: Pass
5.4 Feedback: Pass

DeepSeek (v1.8.5)

AISF Rating Summary

Developer: DeepSeek

Version: 1.8.5

Platform/s: App Store, Google Play, Web-based

Age Ratings: App Store 13+, Google Play 12+, Developer 18+

AISF Safety Benchmark Score: 17%

AISF Rating: F (Critically Unsafe)

AISF Rating Issue Date: April, 2026



DeepSeek (v1.8.5) is rated F (Critically Unsafe) by the AISF. DeepSeek is a conversational AI agent designed for people 12 years and over (as per the minimum age rating assigned to this product). DeepSeek generates content with mild violence, mild ability to support crime, moderate hate speech, mild sexual content, moderate harassment, severe self-harm themes, mild issues with handling suicide related queries, severe child safety issues, mild issues with the use of non-textual elements, mild misinformation issues, mild inaccuracy issues, severe impersonation issues, moderate bias, mild ethical issues, severe use of dark patterns, moderate issues with adversarial robustness, moderate user-generated content issues, moderate over-personalisation, severe privacy issues, severe security issues, and severe feedback issues.

1. Harm

- 1.1 Violence: Fail, Mild
- 1.2 Crime: Fail, Mild
- 1.3 Hate Speech: Fail, Moderate
- 1.4 Sexual Content: Fail, Mild
- 1.5 Harassment: Fail, Moderate
- 1.6 Self-harm: Fail, Severe
- 1.7 Suicide: Fail, Mild
- Child Safety: Fail, Severe
- Non-textual Content: Fail, Mild

2. Content Integrity

- 2.1 Misinformation: Fail, Mild
- 2.2 Inaccuracy: Fail, Mild
- 2.3 Impersonation: Fail, Severe

3. Responsible Engagement

- 3.1 Bias: Fail, Moderate
- 3.2 Ethics: Fail, Mild
- 3.3 Dark Patterns: Fail, Severe

4. Input Handling

- 4.1 Adversarial Robustness: Fail, Moderate
- 4.2 User-generated Content: Fail, Moderate

5. User Data Dynamics

- 5.1 Over-personalisation: Fail, Moderate
- 5.2 Privacy: Fail, Severe
- 5.3 Security: Fail, Severe
- 5.4 Feedback: Fail, Severe

Gemini (v1.2026.1470001)

AISF Rating Summary

Developer: Google

Version: 1.2026.1470001

Platform/s: App Store, Google Play, Web-based, MacOS

Age Ratings: App Store 13+, Google Play 12+, Developer 13+

AISF Safety Benchmark Score: 21%

AISF Rating: F (Critically Unsafe)

AISF Rating Issue Date: April, 2026



Gemini (v1.2026.1470001) is rated F (Critically Unsafe) by the AISF. Gemini is a conversational AI agent designed for people 12 years and over (as per the minimum age rating assigned to this product). Gemini generates content with mild violence, mild ability to support crime, mild hate speech, moderate sexual content, moderate harassment, moderate self-harm themes, mild issues with handling suicide related queries, severe child safety issues, mild issues with the use of non-textual elements, moderate misinformation issues, mild inaccuracy issues, severe impersonation issues, moderate bias, mild ethical issues, moderate use of dark patterns, moderate issues with adversarial robustness, moderate user-generated content issues, moderate over-personalisation, moderate privacy issues, moderate security issues, and moderate feedback issues.

1. Harm
1.1 Violence: Fail, Mild
1.2 Crime: Fail, Mild
1.3 Hate Speech: Fail, Mild
1.4 Sexual Content: Fail, Moderate
1.5 Harassment: Fail, Moderate
1.6 Self-harm: Fail, Moderate
1.7 Suicide: Fail, Mild
Child Safety: Fail, Severe
Non-textual Content: Fail, Mild
2. Content Integrity
2.1 Misinformation: Fail, Moderate
2.2 Inaccuracy: Fail, Mild
2.3 Impersonation: Fail, Severe

3. Responsible Engagement
3.1 Bias: Fail, Moderate
3.2 Ethics: Fail, Mild
3.3 Dark Patterns: Fail, Moderate
4. Input Handling
4.1 Adversarial Robustness: Fail, Moderate
4.2 User-generated Content: Fail, Moderate
5. User Data Dynamics
5.1 Over-personalisation: Fail, Moderate
5.2 Privacy: Fail, Moderate
5.3 Security: Fail, Moderate
5.4 Feedback: Fail, Moderate

Grok (v1.3.62)

AISF Rating Summary

Developer: X Corp.

Version: 1.3.62

Platform/s: App Store, Google Play, Web-based

Age Ratings: App Store 16+, Google Play 18+, Developer 13+

AISF Safety Benchmark Score: 20%

AISF Rating: F (Critically Unsafe)

AISF Rating Issue Date: April, 2026



Grok (v1.3.62) is rated F (Critically Unsafe) by the AISF. Grok is a conversational AI agent designed for people 13 years and over (as per the minimum age rating assigned to this product). Grok generates content with mild violence, mild ability to support crime, severe hate speech, moderate sexual content, moderate harassment, severe self-harm themes, severe child safety issues, mild issues with the use of non-textual elements, mild misinformation issues, mild inaccuracy issues, severe impersonation issues, severe bias, mild ethical issues, severe use of dark patterns, moderate issues with adversarial robustness, moderate user-generated content issues, severe over-personalisation, moderate privacy issues moderate security issues, and moderate feedback issues.

1. Harm

- 1.1 Violence: Fail, Mild
- 1.2 Crime: Fail, Mild
- 1.3 Hate Speech: Fail, Severe
- 1.4 Sexual Content: Fail, Moderate
- 1.5 Harassment: Fail, Moderate
- 1.6 Self-harm: Fail, Severe
- 1.7 Suicide: Pass
- Child Safety: Fail, Severe
- Non-textual Content: Fail, Mild

2. Content Integrity

- 2.1 Misinformation: Fail, Mild
- 2.2 Inaccuracy: Fail, Mild
- 2.3 Impersonation: Fail, Severe

3. Responsible Engagement

- 3.1 Bias: Fail, Severe
- 3.2 Ethics: Fail, Mild
- 3.3 Dark Patterns: Fail, Severe

4. Input Handling

- 4.1 Adversarial Robustness: Fail, Moderate
- 4.2 User-generated Content: Fail, Moderate

5. User Data Dynamics

- 5.1 Over-personalisation: Fail, Severe
- 5.2 Privacy: Fail, Moderate
- 5.3 Security: Fail, Moderate
- 5.4 Feedback: Fail, Moderate

Le Chat (v1.17.1)

AISF Rating Summary

Developer: Mistral AI

Version: 1.17.1

Platform/s: App Store, Google Play, Web-based

Age Ratings: App Store 13+, Google Play 12+, Developer 13+

AISF Safety Benchmark Score: 17%

AISF Rating: F (Critically Unsafe)

AISF Rating Issue Date: April, 2026



Le Chat (v1.17.1) is rated F (Critically Unsafe) by the AISF. Le Chat is a conversational AI agent designed for people 12 years and over (as per the minimum age rating assigned to this product). Le Chat generates content with mild violence, mild ability to support crime, moderate hate speech, moderate sexual content, severe harassment, severe self-harm themes, mild issues with handling suicide related queries, severe child safety issues, mild issues with the use of non-textual elements, moderate misinformation issues, mild inaccuracy issues, severe impersonation issues, severe bias, mild ethical issues, severe use of dark patterns, moderate issues with adversarial robustness, moderate user-generated content issues, moderate over-personalisation, mild privacy issues, moderate security issues, and feedback issues.

1. Harm
- 1.1 Violence: Fail, Mild
 - 1.2 Crime: Fail, Mild
 - 1.3 Hate Speech: Fail, Moderate
 - 1.4 Sexual Content: Fail, Moderate
 - 1.5 Harassment: Fail, Severe
 - 1.6 Self-harm: Fail, Severe
 - 1.7 Suicide: Fail, Mild
 - Child Safety: Fail, Severe

2. Content Integrity
- 2.1 Misinformation: Fail, Moderate
 - 2.2 Inaccuracy: Fail, Mild
 - 2.3 Impersonation: Fail, Severe

3. Responsible Engagement
- 3.1 Bias: Fail, Severe
 - 3.2 Ethics: Fail, Mild
 - 3.3 Dark Patterns: Fail, Severe

4. Input Handling
- 4.1 Adversarial Robustness: Fail, Moderate
 - 4.2 User-generated Content: Fail, Moderate

5. User Data Dynamics
- 5.1 Over-personalisation: Fail, Moderate
 - 5.2 Privacy: Fail, Mild
 - 5.3 Security: Fail, Moderate
 - 5.4 Feedback: Fail, Moderate

Meta AI (v268.0.0)

AISF Rating Summary

Developer: Meta Platforms, Inc.

Version: 268.0.0

Platform/s: App Store, Google Play, Web-based

Age Ratings: App Store 13+, Google Play 3+, Developer 13+

AISF Safety Benchmark Score: 21%

AISF Rating: F (Critically Unsafe)

AISF Rating Issue Date: April, 2026



Meta AI (v268.0.0) is rated F (Critically Unsafe) by the AISF. Meta AI is a conversational AI agent designed for people 3 years and over (as per the minimum age rating assigned to this product). Meta AI generates content with mild violence, mild ability to support crime, severe hate speech, moderate sexual content, moderate harassment, severe self-harm themes, mild issues with handling suicide related queries, severe child safety issues, mild issues with the use of non-textual elements, moderate misinformation issues, mild content inaccuracies, severe impersonation issues, severe bias, mild ethical issues, moderate use of dark patterns, moderate issues with adversarial robustness, moderate user-generated content issues, moderate over-personalisation, severe privacy issues, moderate security issues, and moderate feedback issues.

1. Harm
1.1 Violence: Fail, Mild
1.2 Crime: Fail, Mild
1.3 Hate Speech: Fail, Severe
1.4 Sexual Content: Fail, Moderate
1.5 Harassment: Fail, Moderate
1.6 Self-harm: Fail, Severe
1.7 Suicide: Fail, Mild
Child Safety: Fail, Severe
Non-textual Content: Fail, Mild
2. Content Integrity
2.1 Misinformation: Pass
2.2 Inaccuracy: Fail, Mild
2.3 Impersonation: Fail, Severe

3. Responsible Engagement
3.1 Bias: Fail, Severe
3.2 Ethics: Fail, Mild
3.3 Dark Patterns: Fail, Moderate
4. Input Handling
4.1 Adversarial Robustness: Fail, Moderate
4.2 User-generated Content: Fail, Severe
5. User Data Dynamics
5.1 Over-personalisation: Fail, Mild
5.2 Privacy: Fail, Moderate
5.3 Security: Fail, Moderate
5.4 Feedback: Fail, Moderate

Microsoft Copilot (v30.0.440413003)

AISF Rating Summary

Developer: Microsoft Corporation

Version: 30.0.440413003

Platform/s: App Store, Google Play, Web-based, Windows, MacOS

Age Ratings: App Store 13+, Google Play 3+, Windows 12+, MacOS 13+, Developer 13+

AISF Safety Benchmark Score: 33%

AISF Rating: F (Critically Unsafe)

AISF Rating Issue Date: April, 2026



Microsoft Copilot (v30.0.440413003) is rated F (Critically Unsafe) by the AISF. Microsoft Copilot is a conversational AI agent designed for people 3 years and over (as per the minimum age rating assigned to this product). Microsoft Copilot generates content with mild violence, mild ability to support crime, mild hate speech, moderate sexual content, mild harassment, moderate self-harm themes, mild issues with handling suicide related queries, severe child safety issues, mild issues with the use of non-textual elements, mild misinformation issues, mild inaccuracy issues, moderation impersonation issues, moderate bias, mild ethical issues, mild use of dark patterns, mild issues with adversarial robustness, mild user-generated content issues, moderate over-personalisation, and moderate feedback issues.

1. Harm

- 1.1 Violence: Fail, Mild
- 1.2 Crime: Fail, Mild
- 1.3 Hate Speech: Fail, Mild
- 1.4 Sexual Content: Fail, Moderate
- 1.5 Harassment: Fail, Moderate
- 1.6 Self-harm: Fail, Moderate
- 1.7 Suicide: Fail, Mild
- Child Safety: Fail, Severe
- Non-textual Content: Fail, Mild

2. Content Integrity

- 2.1 Misinformation: Fail, Moderate
- 2.2 Inaccuracy: Fail, Mild
- 2.3 Impersonation: Fail, Severe

3. Responsible Engagement

- 3.1 Bias: Fail, Severe
- 3.2 Ethics: Fail, Mild
- 3.3 Dark Patterns: Fail, Moderate

4. Input Handling

- 4.1 Adversarial Robustness: Fail, Mild
- 4.2 User-generated Content: Fail, Mild

5. User Data Dynamics

- 5.1 Over-personalisation: Fail, Moderate
- 5.2 Privacy: Fail, Moderate
- 5.3 Security: Pass
- 5.4 Feedback: Fail, Mild

Perplexity (v26.13.0)

AISF Rating Summary

Developer: Perplexity AI Inc.

Version: 26.13.0

Platform/s: App Store, Google Play, Web-based, Windows, MacOS

Age Ratings: App Store 18+, Google Play 12+, Windows 3+, MacOS 18+, Developer 13+

AISF Safety Benchmark Score: 23%

AISF Rating: F (Critically Unsafe)

AISF Rating Issue Date: April, 2026



Perplexity (v26.13.0) is rated F (Critically Unsafe) by the AISF. Perplexity is a conversational AI agent designed for people 3 years and over (as per the minimum age rating assigned to this product). Perplexity generates content with mild violence, mild ability to support crime, moderate hate speech, severe harassment, severe self-harm themes, mild issues with handling suicide related queries, severe child safety issues, mild issues with the use of non-textual elements, mild misinformation issues, mild inaccuracy issues, severe impersonation issues, severe bias, mild ethical issues, moderate use of dark patterns, moderate issues with adversarial robustness, moderate user-generated content issues, severe over-personalisation, moderate privacy issues, and moderate feedback issues.

1. Harm

- 1.1 Violence: Fail, Mild
- 1.2 Crime: Fail, Mild
- 1.3 Hate Speech: Fail, Moderate
- 1.4 Sexual Content: Pass
- 1.5 Harassment: Fail, Severe
- 1.6 Self-harm: Fail, Severe
- 1.7 Suicide: Fail, Mild
- Child Safety: Fail, Severe
- Non-textual Content: Fail, Mild

2. Content Integrity

- 2.1 Misinformation: Fail, Mild
- 2.2 Inaccuracy: Fail, Mild
- 2.3 Impersonation: Fail, Severe

3. Responsible Engagement

- 3.1 Bias: Fail, Severe
- 3.2 Ethics: Fail, Mild
- 3.3 Dark Patterns: Fail, Moderate

4. Input Handling

- 4.1 Adversarial Robustness: Fail, Moderate
- 4.2 User-generated Content: Fail, Moderate

5. User Data Dynamics

- 5.1 Over-personalisation: Fail, Severe
- 5.2 Privacy: Fail, Moderate
- 5.3 Security: Pass
- 5.4 Feedback: Fail, Moderate

Qwen (v2.0.3)

AISF Rating Summary

Developer: Alibaba Cloud

Version: 2.0.3

Platform/s: App Store, Google Play, Web-based

Age Ratings: App Store 18+, Google Play 3+, Developer 18+

AISF Safety Benchmark Score: 30%

AISF Rating: F (Critically Unsafe)

AISF Rating Issue Date: April, 2026



Qwen (v2.0.3) is rated F (Critically Unsafe) by the AISF. Qwen is a conversational AI agent designed for people 12 years and over (as per the minimum age rating assigned to this product). Qwen generates content with mild violence, mild ability to support crime, mild sexual content, moderate harassment, moderate self-harm themes, mild issues with handling suicide related queries, severe child safety issues, mild issues with the use of non-textual elements, mild inaccuracy issues, severe impersonation issues, mild bias, mild ethical issues, mild use of dark patterns, mild issues with adversarial robustness, moderate user-generated content issues, mild over-personalisation, severe privacy issues, moderate security issues, and severe feedback issues.

1. Harm
1.1 Violence: Fail, Moderate
1.2 Crime: Fail, Moderate
1.3 Hate Speech: Fail, Mild
1.4 Sexual Content: Fail, Mild
1.5 Harassment: Fail, Severe
1.6 Self-harm: Fail, Severe
1.7 Suicide: Fail, Mild
Child Safety: Fail, Severe
Non-textual Content: Fail, Mild
2. Content Integrity
2.1 Misinformation: Fail, Mild
2.2 Inaccuracy: Fail, Mild
2.3 Impersonation: Fail, Severe

3. Responsible Engagement
3.1 Bias: Fail, Mild
3.2 Ethics: Fail, Mild
3.3 Dark Patterns: Fail, Moderate
4. Input Handling
4.1 Adversarial Robustness: Fail, Severe
4.2 User-generated Content: Fail, Moderate
5. User Data Dynamics
5.1 Over-personalisation: Fail, Moderate
5.2 Privacy: Fail, Severe
5.3 Security: Fail, Severe
5.4 Feedback: Fail, Severe

Support and Resources

If you or someone you know is experiencing suicidal thoughts or a crisis, please reach out to one or more of the following resources.

Hotlines: free, confidential crisis support is available 24/7 by phone or text.

Online chats: anonymous real-time chat with counsellors or peers is available through websites and apps.

Mental health professionals: therapists and psychologists offer personalised support in-person or via telehealth.

Support groups: connect with others who have similar experiences through peer-led online or in-person groups.

Crisis text services: discreet, text-based support is available from trained responders on your phone.

Emergency services: contact police, ambulance, or a hospital for immediate and urgent help.

If you're unsure where to start, a quick web search for "crisis support near me" or "mental health helpline" can often point you to accessible options.

<https://findahelpline.com> offers a global directory of helplines, hotlines, and crisis lines. It covers over 130 countries and allows you to search for support based on your location or specific needs (e.g., suicide prevention, anxiety, depression). You can filter options for phone, text, or chat services, and it provides verified, up-to-date information directly from helpline organizations.

Remember, you're not alone - there are people ready to listen whenever you need them.

Contact Us

The Artificial Intelligence Safety Forum (AISF) is a nonprofit, self-regulatory forum for safety in products using generative AI.

To learn more about the work we do, please visit:

<https://safetyforum.ai/>

If you are the developer of a product using generative AI and would like to learn more about getting your product rated by the AISF, please visit:

<https://safetyforum.ai/developers/>

If you have any comments, queries, or concerns, please contact us at:

contact@safetyforum.ai