

2026

Australia Safe Dating Report

What's Inside

01	A Message From Yoel Roth	03
02	About This Report	04
03	Our Commitment to Safety	05
04	Our Approach to Moderation	07
05	How We Employ Human Oversight	08
06	How We Employ Automation	10
07	Key Measures Supporting Moderator Training & Assistance	12
08	How We Work With Law Enforcement	13
09	Conclusion	14
10	Data	15
11	Appendices	16



A Message from Yoel Roth

At Match Group, our goal is to make our apps the best and safest way to meet new people. This means trust and safety are at the heart of everything we build: Every day, our brands develop and refine the tools that help our users connect with confidence — both online and face to face.

Australia is home to a vibrant community of people seeking meaningful connections. We know that Australians have high expectations for the safety, transparency, and accountability of the services they use.

Australia's Online Dating Code is a critical step toward meeting those expectations — offering a shared framework for dating apps that encourages clearer standards, stronger protections, and more user-centered practices that are consistent across the industry. We recognize that our company, and our industry, have a key role to play in addressing tech-facilitated violence and harm, and believe the Code is a critical part of our commitment to finding solutions.

Today, we're proud to publish our second Transparency Report under the Code, which reflects our ongoing work to protect users in Australia, as well as our larger commitment to setting high safety standards in Australia and around the world. This report outlines how we detect, prevent, and respond to offensive or harmful behavior on and off our platforms — including the proactive actions we take to identify and remove harmful behavior without a report, the role of automation in content moderation, and how we support the people doing this critical work every day.

Across Match Group, we've built a shared infrastructure that powers enforcement across all of our apps — applying what we've learned across our portfolio to strengthen protection on every platform.

Building on the leading work by Australia's eSafety Commissioner, our underlying principle is Safety by Design: a way of building that prioritizes and plans for safety at every step of our product development. This approach has guided our significant investments in responsible AI, scam detection, and pre-match interventions, which allow us to address harm early and at scale. And our partnerships with law enforcement, NGOs, and survivor support organizations help strengthen the broader ecosystem around us.

We've also participated in efforts to protect Australians through legal reform, including leading work in South Australia to bar domestic violence and sex offenders from dating apps — and will continue to advocate for expanded access to data that helps us keep potentially harmful individuals off of our services, in Australia and around the world.

We believe that transparency is essential to making our industry safer, and improving the online dating experience. By showing not just what actions we take, but how we make those decisions, we hope to give Australians a deeper understanding of the safety systems behind our apps. We also recognize that online safety requires ongoing evolution and improvement to meet new challenges. That's why we're constantly listening, adapting, and investing in improving our technology and our policies.

We're committed to making online dating not just safer, but better — for every person, on every platform, everywhere we operate.

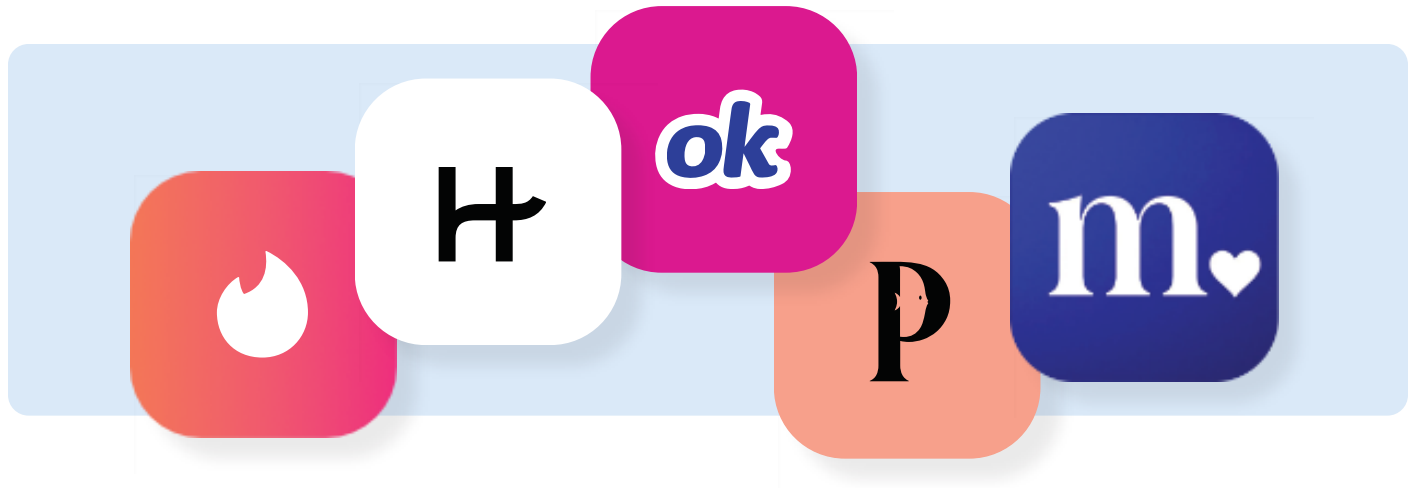


Yoel Roth

HEAD OF TRUST & SAFETY, MATCH GROUP

About This Report

Match Group is on a mission to spark meaningful connections for people around the world. As leaders in the online dating category, we innovate and champion industry best practices that are designed to help make online dating safe and inclusive for everyone.



Match Group offers a portfolio of leading online dating apps that collectively have millions of users who are seeking various types of human connections. This report covers our brands that are signatories to the Code — namely Tinder, Hinge, OkCupid, Plenty of Fish, and Match.com.

This report is published in accordance with section 8.4 of the voluntary Online Safety Code for Dating Services and covers all activity on Match Group apps from 1 July 2025 to 30 June 2026.

As a founding signatory to the Code, Match Group is committed to advancing clear, consistent expectations for user safety across the online dating industry in Australia. Amongst other requirements, signatories are expected to publish annual transparency reports detailing actions taken to uphold user safety — including account terminations, content moderation, use of automation, and support provided to human moderators.

We aim to go beyond the requirements laid out in the Code. Our goal is to increase transparency into both what actions we take and how those actions are decided upon — including the role of centralized detection, user reporting, and trained moderators. By surfacing these decisions clearly, we hope to strengthen public understanding of how trust and safety operates on dating platforms, and how we can leverage technology to create a safer, more accountable and user-first digital ecosystem at scale.

Our Commitment to Safety

We design our platforms to embed safety at every level – from the policies that guide our community, to how features are built, to the ways enforcement operates when something goes wrong. Our approach centers on platform-wide standards that prioritize proactive detection, user control, and real-time interventions.

Our safety strategy is based on four foundational pillars.



01

Safety by Design

Every new feature undergoes an extensive safety review before launch. Our teams are trained to embed risk mitigation into the design of our services.



02

User Control

We give users tools to block, unmatch, and report other users, and we work to reduce friction in safety-related user experiences.



03

User Authenticity

We invest in identity verification, scam prevention, and profile authenticity checks, all with the goal of preventing bad actors from ever making a match.



04

Inclusivity

We aim to build safe and welcoming spaces for users of all identities and backgrounds.

IN AUSTRALIA, WE PUT THESE PRINCIPLES INTO PRACTICE THROUGH FEATURES SUCH AS:

“Are you sure?”

(AYS) nudges that use AI to encourage users to reconsider potentially harmful messages before sending.

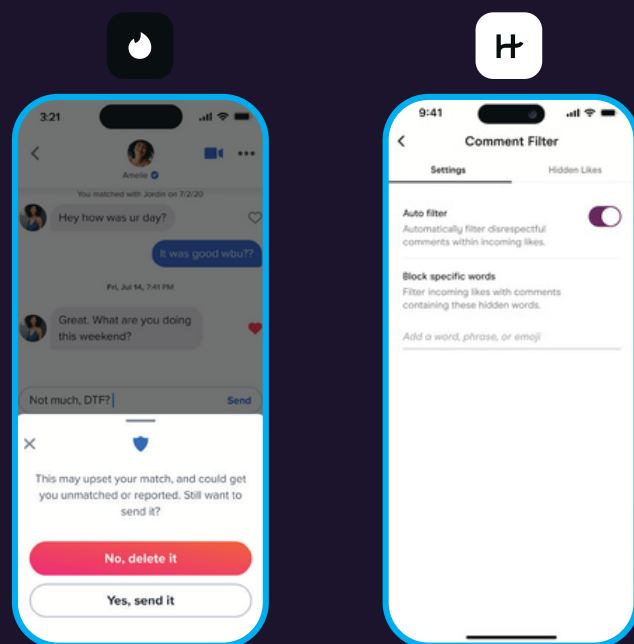
“Does this bother you?”

(DTBY) prompts, which proactively check in with users when our models detect concerning content being shared by another user.

Hinge’s Comment Filter

Which lets users screen offensive language using both personalized and AI-driven filters with inbox separation.

Here are examples of our safety features pioneered by Tinder and our other apps.



But we don't tackle this challenge alone.

We work closely with safety experts, NGOs, and law enforcement to ensure our approach is informed by expertise, responsive to emerging risks, and relevant to local contexts. These partnerships help shape our policies, product development, research, and education efforts.

This includes partnerships with national organizations including **WESNET** and **Teach Us Consent** to ensure our trust and safety strategies are locally informed, and culturally relevant. Through these collaborations, we provide Australian users with localized education, resources, and support services tailored to the country's specific needs.

Globally, we work with leading global organizations such as **NO MORE**, **Disability:IN**, and the **Global Anti-Scam Alliance**, whose insights inform our platform policies and help us stay ahead of emerging safety challenges.

Our Approach to Moderation

Match Group's moderation systems are built to operate at scale, employing a combination of automated and AI-powered detections with expert human review and oversight. In Australia, and around the world, we strive to balance the use of automated tools with human oversight to ensure decisions are fair, accurate, and contextually informed.

The standards and policies we apply across our brands are built on a foundation of centrally-managed principles and approaches to severe harms (like child safety and violence) — while leaving space for each of our brands to innovate and tailor their rules to their specific goals and values. This means that, for some less severe harms, what is allowed on one platform may result in restrictions or a ban on another. The brand-specific data in this report helps represent those differences in enforcement.

Our automated systems support **proactive enforcement** to attempt to mitigate risks to our users before harm occurs:

- In Australia, a significant number of account bans or suspensions were initiated proactively, meaning without the need for a user report.

Our automated systems also support **real-time user interventions**:

- **Face Check** is a face scan feature that requires new users in Australia to take a video selfie during account registration. It checks that the video was taken of a real, live person, and that it was not digitally altered or manipulated, to help ensure the individual is accurately representing themselves. This feature helps reduce fraud, improve platform authenticity, and increase safety by making bad-actor account creation harder.
- **"Are You Sure?"** nudges appear frequently to Tinder users in Australia, prompting them to reconsider sending potentially harmful messages. Globally, we've seen that this feature has been effective in prompting some users to modify their message after seeing an "Are You Sure" prompt.
- **"Does This Bother You?"** prompts which appear after a user receives a potentially inappropriate message from a match encourages users to report the message so the platform can take action if needed.
- To support Australian users traveling abroad, **Traveler Alert on Tinder** temporarily hides user profiles in countries where LGBTQIA+ identities are criminalized; users must opt-in to make their profile visible.

We use automation to detect spam and inauthentic behavior:

Automation is most heavily deployed in spam and fake account detection, which account for the overwhelming majority of bans. Most of the accounts banned in Australia as inauthentic were actioned proactively.

We offer users the choice to proactively filter content:

Via Comment Filter on Hinge, users can proactively filter out comments with certain words, emojis, or phrases. This feature also includes an AI-powered auto-filter, moving filtered comments to a separate inbox, limiting their visibility unless a user chooses to see them.

OUR AGE ASSURANCE PROCESSES

We use a combination of age assurance methods designed to help prevent underage users from accessing our platforms. In December 2025, we introduced updated age assurance processes across our brands operating in Australia to support compliance with Australia's Social Media Minimum Age obligations under the Online Safety Act.

We employ a layered approach to age assurance that begins before a user even has the ability to download our apps, through to account creation and the full user lifecycle:

- **App store restrictions:** All Match Group apps are rated 18+ on iOS and on Google Play. Properly configured parental controls can effectively prevent minors from downloading or accessing these services. Match Group platforms also adopt emerging tools, such as Google's "Restrict Declared Minors" feature, as they become available, to block access to Match Group's apps in search and prevent transactions from accounts identified by app stores as minors.
- **Age-gating at registration:** On Match Group's dating apps, if a user provides a birthdate indicating they are underage, the account is automatically blocked and measures are taken to prevent re-registration. Users who enter an incorrect birthdate may regain access only through verification with a government-issued photo ID.
- **Image and text-based moderation:** Match Group employs a combination of automated tools and human oversight to detect potential underage users. Match Group's brands use either a liveness-based (powered by Face Check) or profile image-based scan to apply AI-powered facial age estimation. The user may not proceed to interact with other users on the platform until they have successfully passed this check. Even after account creation, the platforms continue to scan new images and profile and message text for indications of underage users. Flagged accounts are escalated to human moderators for review.
- **In-app reporting:** Users can easily report suspected underage profiles for prompt review.
- **Photo ID at appeal:** Government-issued photo identification is required as part of the appeal process for users who have been banned because they are suspected to be under the age of 18.

How We Employ Human Oversight

We pair our automated tools with robust human moderation and oversight to ensure a comprehensive and holistic approach to moderation.

- In cases where automated detections may benefit from additional checks, content flagged by users or automation is reviewed by human moderators. Our global moderation teams are reviewing content 24/7.

- All of our moderators receive ongoing training in cultural literacy, scams, gender-based violence, and trauma-informed practices.

- When a new profile is created - and over the course of its presence on our platforms - it undergoes both automated and manual reviews. These checks are focused on identifying and blocking:
 - Ineligible users (e.g., users suspected of being underage)
 - Profiles with suspicious language or inappropriate content
 - Spam or fake accounts

- To ensure objectivity, enforcement appeals are handled by a separate team of moderators than those who made the original enforcement decision. This structure helps maintain fairness, reduces cognitive load on moderators, and ensures that enforcement decisions benefit from second-level review where appropriate.

How We Employ Automation

Match Group employs automation across its portfolio to enhance detection and enforcement, paired with human review. The use of automation differs slightly by brand, but core approaches include:



Photo & Video Moderation

Machine learning models scan uploaded images and videos to detect nudity, sexually explicit material, weapons, extremist imagery, and/or underage indicators. PhotoDNA is used to detect known child sexual abuse material (CSAM).



Text Moderation

Brands use keyword detection, regex lists, and increasingly ML-based natural language models to identify grooming behaviors, hate speech, threats, or terrorist references in usernames, profiles, and the initial messages exchanged between users.



Proactive Messaging Interventions

Features like “Are You Sure?” and Safe Message Filters intervene when automation detects potentially harmful or illegal content.

MATCH GROUP’S USE OF AUTOMATION HAS SEVERAL DISTINCT PURPOSES, INCLUDING:

- Preventing the dissemination of illegal harms including CSAM, extremist, or violent content.
- Detecting underage use and/or grooming behavior.
- Flagging potential harassment, hate, or abusive language for review.
- Taking enforcement action against adult sexual content, spam, and other accurately machine-detectable violations of our policies.
- Deploying proactive user safety prompts in messaging.

MATCH GROUP CONSIDERS AUTOMATION ACCURACY PRIOR TO DEPLOYMENT AND THROUGHOUT THE LIFECYCLE OF ACCOUNTS AND CONTENT.

- **Our approach to enforcement:** Match Group's systems are calibrated to act on potential violations even under some uncertainty, prioritizing user safety. Automated flags can be over-inclusive —multiple reports may relate to one user, or automation may detect false positives.
- **Accuracy target for automated enforcement:** For automated systems that result in account bans, Match Group targets a high level of accuracy — typically 99% — as measured through human labeling of test data prior to deployment. In addition, we perform quality assurance on our moderation decisions, which typically are around 97.5% accurate.
- **Our appeals process:** Automated signals are balanced against human review and appeal processes to avoid over-enforcement. The appeals process is a deliberate check on our automation posture, giving users a route to provide additional context. Beyond resolving individual cases, patterns identified through appeals inform updates to moderation guidance and models.
- **Appeal overturn rate:** Alongside that 99% pre-deployment accuracy figure, 38% of enforcement actions, combining both automated and manual, were ultimately overturned on appeal by human moderators. We work to continuously improve the safety of our apps while balancing the need for accuracy. A few caveats on this figure:
 - It depends heavily on what share of decisions get appealed, which is not constant or representative, so it cannot be read as an accuracy measure on its own.
 - It is an upper-bound estimate shaped by who appeals, not an accuracy rate for enforcement generally: most enforcement decisions are never appealed, and high-volume automated actions targeting spam and fake accounts are appealed at a disproportionately low rate, so the appealing population is small and skewed toward contested cases.
 - Reversals can also reflect updated policy guidance or reviewer discretion in ambiguous cases, not just correction of a mistake.

Key Measures Supporting Moderator Training & Assistance

Match Group invests in the wellbeing and effectiveness of its global moderation teams by combining consistent standards, localized context, and continuous education. These systems help ensure that enforcement decisions are both culturally informed and operationally sound.

Global Standards with Local Context

Moderators are trained on Match Group's unified global policies, informed by expertise from NGOs and other expert organizations. This approach supports consistency in enforcement while incorporating perspectives grounded in the lived experiences of our users and the communities we serve.

Continuous Learning and Specialized Training

Ongoing education sessions were delivered across brands within the past year, covering high-impact topics such as:

- Gender-Based Violence Awareness
- Understanding Discrimination and Fetishization

Mental and Resilience Support

Moderation can expose moderators to harmful content, and we actively invest in programs and initiatives that support their mental wellbeing, including:

- We partner with Zevo Health to offer one-on-one and group counseling for in-house teams at Tinder and Hinge.
- BPO partner sites provide on-site therapist access for all contracted moderators.
- AI systems are also leveraged to reduce direct exposure to certain content types, allowing human moderators to focus on cases where judgment is essential.

How We Work With Law Enforcement

We work closely with law enforcement in Australia and around the world, and view this work as a key part of how we protect not just our platforms, but anyone seeking connection. When we receive law enforcement requests through Kodex, our dedicated portal, our response team provides timely, relevant information to assist with investigations. The overwhelming majority of requests we receive are responded to in under 24 hours.

When a law enforcement request alerts us to a user's alleged misconduct, we also take steps to investigate the activity in question and, if appropriate, remove that user's accounts across our portfolio.

Our Ongoing Commitment to Australia

As the leader in online dating, it's our responsibility to set the standard for safety in meeting new people. We're not just implementing policies, we're creating a culture where safety and respect are at the core of everything we do. This includes:

- Expanding our partnerships with Australian safety organizations
- Regular engagement with regulators and user communities
- Continuing to develop Australia-specific safety features and education campaigns
- Transparent reporting on our progress and challenges

Conclusion

We're committed to making online dating not just safer, but better—for every person, on every one of our platforms, everywhere we operate. Our goal is to foster meaningful connections by building dating experiences that are safe, respectful, and empowering. We are proud to share this Australia-specific transparency report and are committed to building on it in years to come.

As we deepen our data capabilities, develop our processes, policies, features and tools, and expand safety partnerships in Australia, we look forward to improving both the protections we offer and the transparency with which we offer them.

We invite continued dialogue with regulators, experts, users, and civil society to ensure that digital dating in Australia is safe by design and trusted by all.

People deserve an online dating experience that's both fun and safe, where matches treat each other with respect and can be their authentic selves. That's what we're strengthening today, and what will guide us forward.

Data

This section includes detailed data regarding our moderation actions (including account bans, suspensions, and content removal), the reports we received from users, and information about requests we received from law enforcement in Australia during the reporting period from 1 July 2025 to 30 June 2026.

The portfolio nature of our business, as well as the design of our reporting and content moderation systems, impacts the data we report here. Specifically, due to differences in brand enforcement systems, some values (such as proactive ban rates) may be calculated slightly differently from brand to brand. We've made every effort to standardize measures where possible.

A single person using one or more of our apps can be linked to multiple reports and violations, meaning the number of bans may be lower than the number of reports or violations. Second, if one person's accounts are removed across multiple platforms, it will count as multiple bans, even though it involves just one individual (and possibly only one report). Lastly, some banned accounts belong to users trying to create new accounts after previous bans, which we make an effort to attribute to the originally banned individual.

Many of our bans in the Spam, Inauthentic, and Ineligible Accounts category take place during or shortly after account onboarding — before these accounts have any negative impact on the experience of authentic or legitimate user accounts. In this report, we provide a calculation of ban rates as a percentage of active users during the reporting period. However, in most cases, accounts banned as Spam, Inauthentic, or Ineligible are not counted as active users. As a result, these calculations may show a larger percentage than is actually representative of the number of active users banned in this category.

Users on our platforms are encouraged to report any suspicious activity or bad behaviour directly to us. Users can report anyone regardless if they've matched with them or not and regardless of whether the interaction occurred on our platforms or elsewhere. They can select from a number of reasons for reporting such as abusive or threatening behaviour, illegal activities or spam accounts.

If a user contacts us to report any bad online or offline behaviour, our team carefully reviews the report and takes the necessary action to remove any profile that breaches our terms of use or community guidelines of each platform.

If a user contacts us to report any bad online or offline behaviour, our team carefully reviews the report and takes the necessary action to remove any profile that breaches our terms of use or community guidelines of each platform.

For more information on the categories in the following tables, see Appendix E for our Policy Definitions.

Total Bans and Suspensions

All Match Group Signatories

TINDER · HINGE · MATCH.COM · PLENTY OF FISH · OKCUPID

TYPE OF POLICY	TOTAL ACCOUNTS BANNED	ACCOUNTS BANNED DETECTED PROACTIVELY (%)	% OF ACTIVE USERS DURING THE REPORTING PERIOD
Abuse and harassment	12,742	85.3%	0.06%
Allegations of off platform harm	12,780	57.6%	0.06%
Illegal and regulated activities	2,406	46.2%	0.01%
Non-T&S ToS violations	4,472	77.8%	0.02%
Other	27,932	74.2%	0.14%
Sensitive content	2,187	64.2%	0.01%
Spam, inauthentic, and ineligible accounts	359,631	90.9%	1.75%
Violence and hate	2,763	76.8%	0.01%
Total	424,913	88.0%	2.07%

Tinder

TYPE OF POLICY	TOTAL ACCOUNTS BANNED	ACCOUNTS BANNED DETECTED PROACTIVELY (%)	% OF ACTIVE USERS DURING THE REPORTING PERIOD
Abuse and harassment	371	15.1%	0.00%
Allegations of off platform harm	7,842	59.6%	0.08%
Illegal and regulated activities	549	44.1%	0.01%
Non-T&S ToS violations	2,979	77.4%	0.03%
Other	0	0.0%	0.00%
Sensitive content	537	39.5%	0.01%
Spam, inauthentic, and ineligible accounts	43,501	83.4%	0.45%
Violence and hate	873	35.2%	0.01%
Total	56,652	77.8%	0.58%

Hinge

TYPE OF POLICY	TOTAL ACCOUNTS BANNED	ACCOUNTS BANNED DETECTED PROACTIVELY (%)	% OF ACTIVE USERS DURING THE REPORTING PERIOD
Abuse and harassment	11,138	96.8%	0.12%
Allegations of off platform harm	4,888	54.8%	0.05%
Illegal and regulated activities	1,247	56.1%	0.01%
Non-T&S ToS violations	99	90.9%	0.00%
Other	0	0.0%	0.00%
Sensitive content	641	99.2%	0.01%
Spam, inauthentic, and ineligible accounts	133,472	98.8%	1.46%
Violence and hate	1,871	96.8%	0.02%
Total	153,356	96.9%	1.68%

Plenty of Fish

TYPE OF POLICY	TOTAL ACCOUNTS BANNED	ACCOUNTS BANNED DETECTED PROACTIVELY (%)	% OF ACTIVE USERS DURING THE REPORTING PERIOD
Abuse and harassment	1,075	1.8%	0.08%
Allegations of off platform harm	32	12.5%	0.00%
Illegal and regulated activities	584	26.2%	0.04%
Non-T&S ToS violations	1,302	77.0%	0.09%
Other	27,922	74.2%	1.96%
Sensitive content	734	44.7%	0.05%
Spam, inauthentic, and ineligible accounts	134,623	85.1%	9.45%
Violence and hate	12	16.7%	0.00%
Total	166,284	82.3%	11.67%

APPENDIX A, CONTINUED

OkCupid

TYPE OF POLICY	TOTAL ACCOUNTS BANNED	ACCOUNTS BANNED DETECTED PROACTIVELY (%)	% OF ACTIVE USERS DURING THE REPORTING PERIOD
Abuse and harassment	158	5.1%	0.08%
Allegations of off platform harm	17	5.9%	0.01%
Illegal and regulated activities	23	65.2%	0.01%
Non-T&S ToS violations	26	61.5%	0.01%
Other	0	0.0%	0.00%
Sensitive content	265	81.9%	0.14%
Spam, inauthentic, and ineligible accounts	16,923	80.5%	9.00%
Violence and hate	7	14.3%	0.00%
Total	17,419	79.7%	9.26%

Match.com

TYPE OF POLICY	TOTAL ACCOUNTS BANNED	ACCOUNTS BANNED DETECTED PROACTIVELY (%)	% OF ACTIVE USERS DURING THE REPORTING PERIOD
Abuse and harassment	0	0.0%	0.00%
Allegations of off platform harm	1	0.0%	0.00%
Illegal and regulated activities	3	66.7%	0.01%
Non-T&S ToS violations	66	93.9%	0.20%
Other	10	100.0%	0.03%
Sensitive content	10	100.0%	0.03%
Spam, inauthentic, and ineligible accounts	31,112	98.6%	91.96%
Violence and hate	0	0.0%	0.00%
Total	31,202	98.6%	92.23%

All Match Group Signatories

TINDER · HINGE · MATCH.COM · PLENTY OF FISH · OKCUPID

Total Banned	424,913
Accounts Banned Detected Proactively	376,304 (88%)
Accounts Banned by Automation	253,800 (59%)

Content Removal

All Match Group signatories

TINDER · HINGE · PLENTY OF FISH · OKCUPID · MATCH.COM

TYPE OF POLICY	TOTAL CONTENT REMOVED	CONTENT REMOVED BY AUTOMATION (%)
Abuse and harassment	762	45.1%
Allegations of off platform harm	0	0.0%
Illegal and regulated activities	4,071	86.1%
Non-T&S ToS violations	281,149	0.0%
Other	1,097	0.0%
Sensitive content	66,892	4.9%
Spam, inauthentic, and ineligible accounts	27,809	2.4%
Violence and hate	354	2.8%
Total	382,134	2.0%

Tinder

TYPE OF POLICY	TOTAL CONTENT REMOVED	CONTENT REMOVED BY AUTOMATION (%)
Abuse and harassment	357	96.4%
Allegations of off platform harm	0	0.0%
Illegal and regulated activities	4,065	86.3%
Non-T&S ToS violations	11,782	0.0%
Other	20	0.0%
Sensitive content	32,859	10.0%
Spam, inauthentic, and ineligible accounts	14,521	4.7%
Violence and hate	14	71.4%
Total	63,618	12.3%

Hinge

TYPE OF POLICY	TOTAL CONTENT REMOVED	CONTENT REMOVED BY AUTOMATION (%)*
Abuse and harassment	405	0.0%
Allegations of off platform harm	0	0.0%
Illegal and regulated activities	5	0.0%
Non-T&S ToS violations	9,714	0.0%
Other	0	0.0%
Sensitive content	1,746	0.0%
Spam, inauthentic, and ineligible accounts	195	0.0%
Violence and hate	327	0.0%
Total	12,392	0.0%

*. On Hinge, automated detection systems flag profiles suspected of policy violations for suspension pending human moderator review, rather than removing specific content. Hinge's automation is concentrated at the account level (Appendix A) rather than content removal (Appendix B).

Plenty of Fish

TYPE OF POLICY	TOTAL CONTENT REMOVED	CONTENT REMOVED BY AUTOMATION (%)*
Abuse and harassment	0	0.0%
Allegations of off platform harm	0	0.0%
Illegal and regulated activities	1	0.0%
Non-T&S ToS violations	251,174**	0.0%
Other	29	0.0%
Sensitive content	30,160	0.0%
Spam, inauthentic, and ineligible accounts	12,796	0.0%
Violence and hate	13	0.0%
Total	294,173	0.0%

*. On Plenty of Fish, automated detection systems flag profiles suspected of policy violations for suspension pending human moderator review, rather than removing specific content. Plenty of Fish's automation is concentrated at the account level (Appendix A) rather than content removal (Appendix B).

**. Plenty of Fish takes down content for users who do not include a picture of their face.

APPENDIX B, CONTINUED

OkCupid

TYPE OF POLICY	TOTAL CONTENT REMOVED	CONTENT REMOVED BY AUTOMATION (%)*
Abuse and harassment	0	0.0%
Allegations of off platform harm	0	0.0%
Illegal and regulated activities	0	0.0%
Non-T&S ToS violations	4,410	0.0%
Other	1,006	0.0%
Sensitive content	1,597	0.0%
Spam, inauthentic, and ineligible accounts	193	0.0%
Violence and hate	0	0.0%
Total	7,206	0.0%

*. On OkCupid, automated detection systems flag profiles suspected of policy violations for suspension pending human moderator review, rather than removing specific content. OkCupid's automation is concentrated at the account level (Appendix A) rather than content removal (Appendix B).

Match.com

TYPE OF POLICY	TOTAL CONTENT REMOVED	CONTENT REMOVED BY AUTOMATION (%)*
Abuse and harassment	0	0.0%
Allegations of off platform harm	0	0.0%
Illegal and regulated activities	0	0.0%
Non-T&S ToS violations	4,069	0.0%
Other	42	0.0%
Sensitive content	530	0.0%
Spam, inauthentic, and ineligible accounts	104	0.0%
Violence and hate	0	0.0%
Total	4,745	0.0%

*. On Match.com, automated detection systems flag profiles suspected of policy violations for suspension pending human moderator review, rather than removing specific content. Match.com's automation is concentrated at the account level (Appendix A) rather than content removal (Appendix B).

User Reports

All Match Group signatories

CATEGORY OF REPORT	TOTAL REPORTS RECEIVED
Abuse and harassment	46,236
Allegations of off platform harm	28,053
Illegal and regulated activities	14,759
Non-T&S ToS violations	35,552
Other	4,089
Sensitive content	90,710
Spam, inauthentic, and ineligible accounts	304,508
Violence and hate	18,447
Total	542,354

Tinder

CATEGORY OF REPORT	TOTAL REPORTS RECEIVED
Abuse and harassment	12,894
Allegations of off platform harm	10,612
Illegal and regulated activities	0
Non-T&S ToS violations	16,332
Other	0
Sensitive content	34,241
Spam, inauthentic, and ineligible accounts	115,453
Violence and hate	10,544
Total	200,076

APPENDIX C, CONTINUED

Hinge

CATEGORY OF REPORT	TOTAL REPORTS RECEIVED
Abuse and harassment	26,878
Allegations of off platform harm	16,964
Illegal and regulated activities	14,316
Non-T&S ToS violations	7,353
Other	0
Sensitive content	54,964
Spam, inauthentic, and ineligible accounts	131,666
Violence and hate	7,903
Total	260,044

Plenty of Fish

CATEGORY OF REPORT	TOTAL REPORTS RECEIVED
Abuse and harassment	6,460
Allegations of off platform harm	300
Illegal and regulated activities	443
Non-T&S ToS violations	8,883
Other	3,322
Sensitive content	1,340
Spam, inauthentic, and ineligible accounts	48,391
Violence and hate	0
Total	69,139

APPENDIX C, CONTINUED

OkCupid

CATEGORY OF REPORT	TOTAL REPORTS RECEIVED
Abuse and harassment	0
Allegations of off platform harm	175
Illegal and regulated activities	0
Non-T&S ToS violations	2,899
Other	767
Sensitive content	155
Spam, inauthentic, and ineligible accounts	8,572
Violence and hate	0
Total	12,568

Match.com

CATEGORY OF REPORT	TOTAL REPORTS RECEIVED
Abuse and harassment	4
Allegations of off platform harm	2
Illegal and regulated activities	0
Non-T&S ToS violations	85
Other	0
Sensitive content	10
Spam, inauthentic, and ineligible accounts	426
Violence and hate	0
Total	527

Legal Orders

All Match Group signatories

CATEGORY OF GOVERNMENT ORDER	TOTAL REPORTS RECEIVED
Abuse and harassment	12
Allegations of off platform harm	57
Illegal and regulated activities	1
Non-T&S ToS violations	1
Other	0
Sensitive content	0
Spam, inauthentic, and ineligible accounts	18
Violence and hate	1
Total	90

Tinder

CATEGORY OF GOVERNMENT ORDER	TOTAL REPORTS RECEIVED
Abuse and harassment	6
Allegations of off platform harm	32
Illegal and regulated activities	1
Non-T&S ToS violations	1
Other	0
Sensitive content	0
Spam, inauthentic, and ineligible accounts	10
Violence and hate	1
Total	51

Hinge

CATEGORY OF GOVERNMENT ORDER	TOTAL REPORTS RECEIVED
Abuse and harassment	5
Allegations of off platform harm	19
Illegal and regulated activities	0
Non-T&S ToS violations	0
Other	0
Sensitive content	0
Spam, inauthentic, and ineligible accounts	2
Violence and hate	0
Total	26

Plenty of Fish

CATEGORY OF GOVERNMENT ORDER	TOTAL REPORTS RECEIVED
Abuse and harassment	1
Allegations of off platform harm	6
Illegal and regulated activities	0
Non-T&S ToS violations	0
Other	0
Sensitive content	0
Spam, inauthentic, and ineligible accounts	5
Violence and hate	0
Total	12

Match.com

CATEGORY OF GOVERNMENT ORDER	TOTAL REPORTS RECEIVED
Abuse and harassment	0
Allegations of off platform harm	0
Illegal and regulated activities	0
Non-T&S ToS violations	0
Other	0
Sensitive content	0
Spam, inauthentic, and ineligible accounts	1
Violence and hate	0
Total	1

OkCupid

Did not receive any inquiries from law enforcement during the reporting period.

Policy Definitions

POLICY CATEGORY	EXAMPLE SUBCATEGORIES	DEFINITION
Abuse and Harassment	■ Non Consensual Sharing of Private Information ■ Harassment ■ Abusive Behavior	Abuse and Harassment includes on-platform behavior, actions, or content that may cause emotional or psychological harm through the use of intimidation, humiliation, or non-consensual acts. This includes persistent and unwanted interactions that threaten someone's safety or well-being.
Allegations of Off Platform Misconduct	■ Allegations of Abuse and Harassment ■ Allegations of Physical Harm ■ Allegations of Sexual Exploitation ■ Allegations of Financial Harm ■ Third party allegations of off platform harm	Allegations of Off-Platform Misconduct include reports alleging harmful behaviors or actions that occurred outside the platform that indicate a potential risk to the safety, well-being, or integrity of the user.
Illegal and Regulated Activities	■ Child Sexual Exploitation and Enticement ■ Commercial Sex and Solicitation ■ Copyright and Trademark Infringement ■ Illegal Use and Regulated Goods	Illegal and Regulated Activities includes behaviors, actions, or content that violate laws or regulations or are subject to specific heightened legal controls on distribution, sale, or promotion. We also strictly prohibit any content, behavior, or actions that involve the sexual abuse or exploitation of minors, including attempts to groom or entice minors into sexual activities (on- or offline).
Non Trust & Safety Terms of Service Violations	■ Ineligible Image or Profile ■ Content	In certain cases, specific Match Group platforms may restrict particular behaviors under their terms of service, even if such behaviors do not represent a safety or authenticity risk.

Other

In limited cases, despite our best efforts, we may not have specific policy attribution for some moderation actions taken.

Sensitive Content

- Adult Nudity, Pornography, and Sexualized Content
- Graphic Content
- Controlled Substance Use
- Suicide, Suicidal Ideation, and Self-Harm

Sensitive Content includes content related to adult nudity, pornography, sexualized content, violence and gore, substance use, and behaviors associated with suicide, suicidal ideation, self-harm, and disordered eating.

Spam, Inauthentic, and Ineligible Accounts

- Attempted Financial Exploitation
- Ban Evasion
- False Reporting
- Impersonation
- Spam and Inauthentic Accounts
- Suspected Underage Users
- Convicted Violent Offenders

Spam, Inauthentic, and Ineligible Accounts includes profiles that engage in inauthentic activities, evade bans, impersonate others, or are otherwise ineligible for account creation due to age or legal status.

Violence and Hate

- Hateful and Discriminatory Behavior
- Threats and Wishes of Harm
- Terrorism, Violence, Extremism, and Hate Groups

Violence and Hate includes on-platform content or behavior that incites, promotes, or glorifies bodily harm or hatred against individuals or groups.

Violence refers to any form of physical harm, including threats, intimidation, and coercion, regardless of motivation or perceived intent.

Hate involves behavior or content that demeans, marginalizes, or incites violence against individuals or groups based on an actual or perceived association with a protected characteristic or class.

