

HATE HOSTED HERE.

AN **ANTISEMITISM**
ACCOUNTABILITY
REPORT ON
SOCIAL MEDIA



AN EXAMINATION OF
ENFORCEMENT **FAILURES**
AND **SUCCESSES**

EDITION I | 2025 - 2026



Gas all the 🍌🍌
2w 123 like Reply



The Painter was right.
1w 276 like Reply



The 🍷🍷 run the world.
3w 39 like Reply



6M isn't a loss #HH
1w 75 like Reply



⚡ Should have killed them all ⚡
4w 436 like Reply

CONTENTS

I.	INTRODUCTION	03
	<i>Why this report matters</i>	
II.	METHODOLOGY	07
	<i>Reporting process and evaluation criteria</i>	
III.	KEY FINDINGS	09
	<i>Summary of major findings</i>	
IV.	PLATFORM-BY-PLATFORM ANALYSIS	13
	<i>TikTok, Meta, and X</i>	
V.	REINSTATEMENT AND BAN EVASION	19
	<i>When removal fails to stop the spread of extremist content</i>	
VI.	RECOMMENDATIONS	22
	<i>Five measures for platforms and policymakers</i>	
VII.	CONCLUSION	24
	<i>Findings and the path forward</i>	
VIII.	STOPANTISEMITISM	25
	<i>About us & Contact</i>	

Why This Report Matters

*In an era of rampant online antisemitism, we have to ask:
Why aren't major social media platforms doing more to combat it?*

STOPANTISEMITISM, POLICY TEAM

Purpose of the Report

Jew hatred has existed for centuries, adapting itself to every major communication technology and social institution it encounters. From sermons and pamphlets to newspapers and television, anti-Jewish conspiracy theories have repeatedly found new ways to spread, evolve, and reach mass audiences. Today, social media has fundamentally transformed the speed, scale, and reach of that process. Antisemitic propaganda can now be distributed to millions of people within minutes, accelerating the spread of anti-Jewish hatred and increasing the risk of radicalization and real-world violence.

Major social media platforms maintain publicly stated guidelines governing harmful content. Yet, little public information exists regarding how consistently those guidelines are enforced when antisemitic content is reported. This is not merely an accountability issue; it is a public safety issue.

Unlike government agencies that publish detailed breakdowns of hate crime data, major social media platforms generally do not publicly disclose category-specific data that would allow researchers, policymakers, or the public to assess how antisemitic content is reported, reviewed, and addressed relative to other forms of hate.

As a result, the public is largely expected to trust that antisemitic content is being reviewed and addressed appropriately despite having little ability to independently evaluate those claims. At the same time, the sheer volume of antisemitic content that remains readily accessible across major social media platforms raises an important question: **Are these companies consistently enforcing the standards they publicly promote, or do those standards function more as public-facing commitments rather than operational practice?**

Understanding whether platforms are effectively enforcing their own policies is therefore critical to evaluating their role in combating online antisemitism and reducing the risk of further radicalization and violence.

Understanding Modern Online Antisemitism

All examples of antisemitic content referenced throughout this report were evaluated using the **International Holocaust Remembrance Alliance (IHRA)** Working Definition of Antisemitism. This internationally recognized framework provides the analytical basis for identifying the antisemitic content discussed throughout this report and helps distinguish legitimate political expression from rhetoric that targets or demonizes Jews as a people.

Not all antisemitic content is equally easy to identify or moderate. While some material clearly violates platform policies through explicit hate speech, praise for terrorist organizations, or direct calls for violence, a substantial portion of the content documented during this review relied on misinformation, historical distortion, conspiracy theories, and coded language that is significantly more challenging to detect and evaluate.

Unlike explicit violations, misinformation often presents itself as news reporting, historical analysis, political commentary, or humanitarian advocacy. False or misleading claims may be communicated through selectively edited videos, fabricated statistics, manipulated images, AI-generated media, or emotionally charged narratives that appear credible to audiences unfamiliar with the underlying facts. In many cases, the antisemitic message is implied rather than stated outright, making it less likely to be identified by automated moderation systems or reviewers without relevant historical or geopolitical context.

This challenge is further complicated by the growing use of “algospeak,” intentional misspellings, coded language, emojis, symbols, and euphemisms designed to evade automated detection while conveying substantially the same message.

Common Examples Include References Such As:

- “HH” – “Heil Hitler.”
- “Austrian painter” – a euphemism for Adolf Hitler.
- Lightning bolt emojis ⚡ – used to reference the SS.
- Juice box emojis 🍹 – used as coded references to “Jew.”
- Exaggerated nose emojis 🍑 – used as antisemitic imagery.
- Deliberately altered spellings intended to evade automated moderation systems.

As moderation systems become more sophisticated, users continually adapt their language and imagery to avoid enforcement.

While not every instance of misinformation violates a platform’s policies, repeated exposure to false and dehumanizing narratives can contribute to an environment in which hatred becomes normalized and violence becomes more easily justified. For this reason, effective and consistent moderation requires more than detecting explicit slurs or threats; it also demands the ability to recognize evolving narratives, historical tropes, and coded forms of antisemitic expression that increasingly characterize online discourse.

Why Measuring Enforcement Matters

We have repeatedly seen how antisemitic rhetoric and conspiracy theories amplified online can radicalize individuals, normalize hatred of Jews, and ultimately fuel deadly real-world violence – making the fight against online antisemitism a matter of public safety, not just speech.

LIORA REZ, STOPANTISEMITISM FOUNDER

While not every individual exposed to online hate becomes radicalized, almost every major antisemitic attack has involved perpetrators who consumed, shared, or actively participated in online ecosystems promoting antisemitic conspiracy theories, extremist ideologies, and violence. Researchers and law enforcement agencies have repeatedly documented the role online platforms can play in accelerating radicalization, spreading conspiracy theories, and amplifying extremist ideologies.

In 2018, before murdering eleven worshippers at the Tree of Life synagogue in Pittsburgh, the attacker regularly posted antisemitic conspiracy theories on Gab. In 2019, the San Diego Poway synagogue shooter published an antisemitic manifesto on 8chan moments before carrying out his attack. The 2019 perpetrators of the Jersey City kosher supermarket attack and the 2021 hostage-taking at Congregation Beth Israel in Colleyville, Texas were also both influenced by extremist propaganda and antisemitic ideologies that circulated online.

The same pattern has emerged in more recent attacks motivated by anti-Israel extremism. The man charged with murdering Yaron Lischinsky and Sarah Milgrim outside the Capital Jewish Museum in Washington, D.C. on May 21, 2025 published writings calling to “bring the war home” and, after the attack, repeatedly shouted “Free Palestine” while stating he acted “for Gaza.” A few days later, Mohamed Sabry Soliman, the man accused of carrying out the firebomb attack against a Run for Their Lives gathering in Boulder, Colorado shouted “Free Palestine” during the attack, which ultimately claimed the life of Holocaust survivor Karen Diamond.

These cases span different ideologies and different online spaces, but they point to the same troubling pattern: online ecosystems can reinforce hatred, normalize violence, and accelerate radicalization before that hatred manifests in the real world.

It is for this reason that StopAntisemitism conducted a yearlong review of reports submitted to TikTok, Meta, and X between July 2025 and July 2026. Throughout the reporting period, our team submitted documented reports on a rolling weekly basis identifying accounts that appeared to violate each platform’s publicly stated community standards.

Each report was supported by multiple documented examples of the account’s conduct and included evidence of potential violations involving hate speech, promotion of terrorist organizations or extremist ideologies, harassment, incitement to violence, glorification of violence, or other prohibited content. By applying a consistent methodology across all three platforms, this review provides a comparative assessment of how major social media companies enforce their own policies when confronted with documented instances of digital hate.

“ *This report does not evaluate whether content is merely offensive or controversial. It evaluates whether content that violated each platform's own published policies was enforced consistently after being reported.*

STOPANTISEMITISM, POLICY TEAM

This report analyzes content enforcement outcomes for reports submitted to TikTok, Instagram (Meta), and X between July 2025 and July 2026. All assessments were based on each platform's publicly available policies in effect during the study, including TikTok's Community Guidelines, Meta's Community Standards, and X's Hateful Conduct Policy.

Throughout the yearlong review, StopAntisemitism submitted reports directly to each platform on a weekly basis. Each submission focused on a single account and included an average of five documented examples of content that appeared to violate the platform's published Community Guidelines or Terms of Service. Every report cited the applicable policy provisions and explained why the reported content appeared to meet the criteria for enforcement.

Selection Criteria

Accounts were selected for reporting based on repeated or egregious examples of content that appeared to violate one or more platform policies. In prioritizing reports, StopAntisemitism also considered factors such as follower count and potential audience reach, recognizing that accounts with larger audiences have a greater capacity to amplify harmful content. However, audience size was never the sole criterion for selection, and accounts with both large and small followings were reported based on the nature and severity of the apparent violations.

Enforcement Outcomes

For the purposes of this report, platform responses were categorized using the following enforcement outcomes:

- **Account Removal:** The reported account was permanently suspended or removed.
- **Content Actioned:** One or more reported posts were removed, restricted, de-monetized, or otherwise actioned while the account remained active.
- **No Action:** The platform determined that no violation had occurred or declined to take visible enforcement action.
- **Repeat Offender:** An account that continued posting substantially similar content following an enforcement action.
- **Ban Evader:** An account created to circumvent a previous suspension or removal.



SECTION III KEY FINDINGS

This yearlong review found that all three platforms enforced their published policies against antisemitic content to varying degrees, but enforcement outcomes differed significantly depending on the platform. Some platforms demonstrated a greater willingness to permanently remove repeat offenders, while others primarily removed individual posts or declined to take action altogether. The findings below summarize the most significant enforcement trends observed during the reporting period.

At a Glance

3 PLATFORMS STUDIED	126 ACCOUNTS REPORTED	762 DOCUMENTED EXAMPLES
16.3M+ FOLLOWERS REPRESENTED	50–100M ESTIMATED REACH	12 MONTHS OF REPORTING

PLATFORM	DOCUMENTED EXAMPLES	ACCOUNTS REPORTED
TikTok	324	49
X	263	41
Meta	175	36
Total	762	126

TikTok Demonstrated the Strongest Account-Level Enforcement

Among the three platforms analyzed, TikTok permanently removed the highest percentage of reported accounts. Nearly half of all reported TikTok accounts were ultimately removed, representing a significantly higher account-level enforcement rate than either Meta or X.

TikTok also demonstrated enforcement across multiple areas of the platform, including account removals, content moderation, demonetization, and, in several cases, enforcement against antisemitic merchandise sold through TikTok Shop. The platform additionally removed certain antisemitic sounds and audio used to amplify hateful content, illustrating a broader approach to enforcement than simply removing individual videos.

SECTION III KEY FINDINGS

Supporting Statistics

Accounts Reported	49
Account Removals	24
Account Removal Rate	49%
Documented Content Examples Submitted	324
Individual Pieces of Content Actioned	156
Content Removal Rate	48%

Meta Frequently Removed Content While Leaving Problematic Accounts Live

Meta’s enforcement strategy differed markedly from TikTok’s. Rather than removing repeat offenders outright, Meta more often removed individual posts, reels, or stories while allowing the account to remain active.

This approach reduced the visibility of violating content but frequently allowed repeat offenders to continue posting similar material after enforcement actions.

Supporting Statistics

Accounts Reported	36
Account Removals	12
Account Removal Rate	33%
Documented Content Examples Submitted	175
Individual Pieces of Content Actioned	70
Content Removal Rate	40%

X Demonstrated the Lowest Rate of Account-Level Enforcement Across the Board

Among the platforms analyzed, X removed the smallest proportion of reported accounts. In many cases, reported content was determined to be “political speech” or otherwise non-violative, even when reports documented repeated examples of antisemitic rhetoric, praise for terrorist organizations, or calls supporting extremist groups.

SECTION III KEY FINDINGS

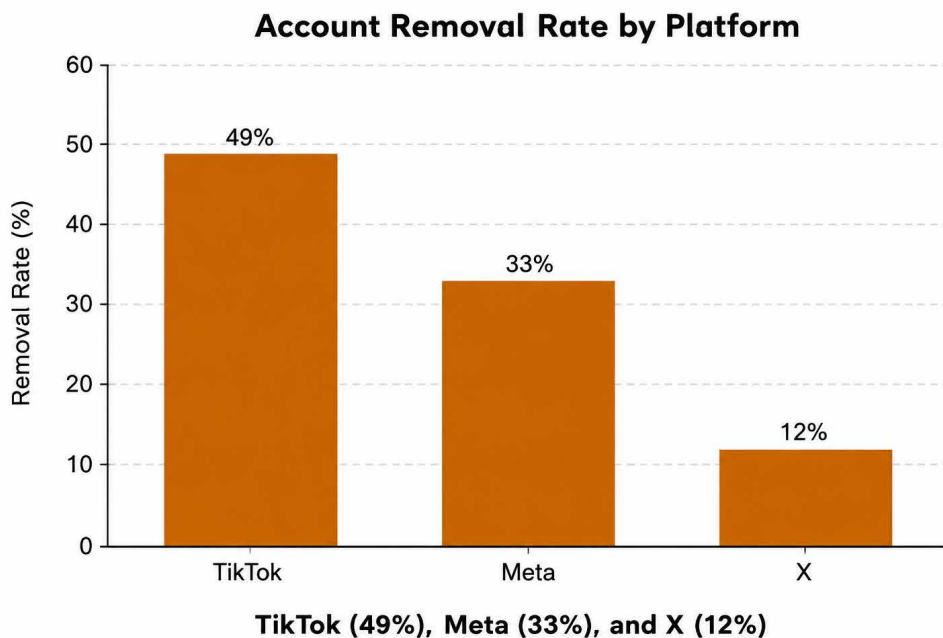
While X did remove individual posts and suspend a limited number of accounts, account-level enforcement occurred far less frequently than on TikTok or Meta.

Supporting Statistics

Accounts Reported	41
Account Removals	5
Account Removal Rate	12%
Documented Content Examples Submitted	263
Individual Pieces of Content Actioned	46
Content Removal Rate	17%

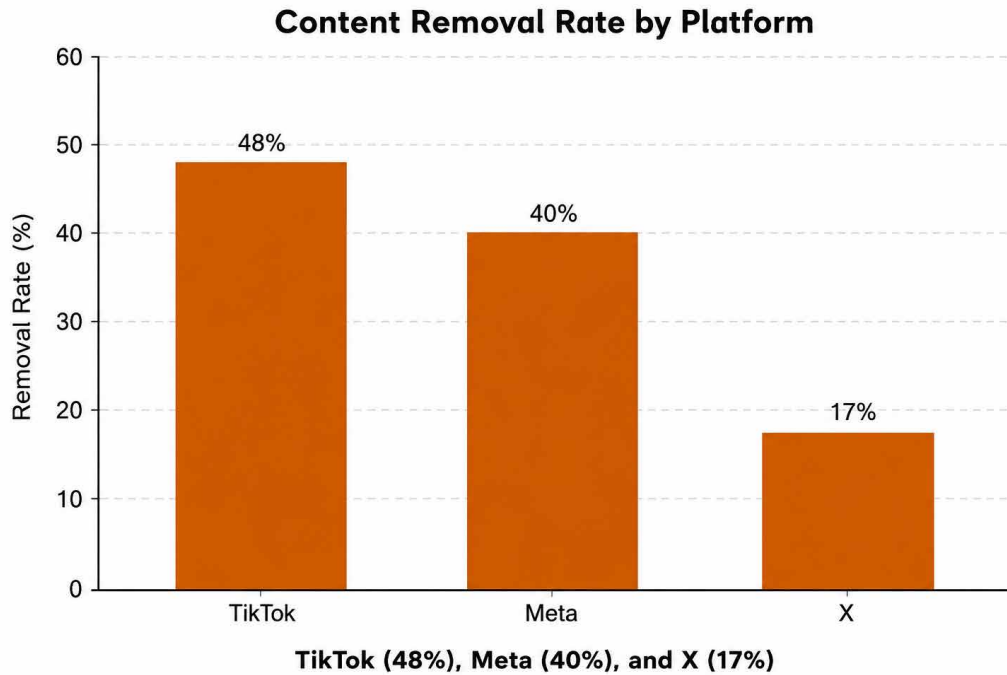
Platform Enforcement Varied Significantly

The chart below compares permanent account removal rates across TikTok, Meta, and X.



SECTION III KEY FINDINGS

The chart below compares permanent individual content removal rates across TikTok, Meta, and X.



Key Finding: Ban Evasion Remains a Persistent Challenge

Account removal does not necessarily end policy violations. Across TikTok, Meta, and X, multiple creators returned under new usernames after being suspended or permanently removed. Some rebuilt substantial audiences before additional enforcement action occurred, while others remained active despite repeated reports documenting ban evasion.

Although several repeat offenders were ultimately removed following subsequent reports, the findings suggest that existing systems for identifying and preventing ban evasion remain inconsistent across all three platforms.

Supporting Statistics

Identified cases of ban evasion	9
Subsequently removed after additional reporting	4
Ban evasion removal rate	44%

One of the clearest findings of this analysis is that major social media platforms do not interpret or enforce hate speech policies uniformly. While each platform publicly prohibits hateful conduct such as Holocaust denial and content promoting violence or terrorism, enforcement outcomes varied considerably.

The following sections examine each platform individually, highlighting overall enforcement trends, notable successes, recurring challenges, and representative case studies.

TikTok OVERALL ASSESSMENT

TikTok demonstrated the strongest overall enforcement of the three platforms analyzed. Nearly half of all reported accounts were permanently removed during the reporting period, representing the highest account-level enforcement rate observed in this review. Beyond account removals, TikTok also enforced policies against monetization, including antisemitic merchandise sold through TikTok Shop, and, in several instances, removed antisemitic sounds and audio used to amplify hateful content.



Hiphop titanium steel pendant, bold, stylis...

From **\$7.99** \$15.99

A notable development during the reporting period was TikTok's creation of a dedicated Public Policy Manager for Hate Speech. While this report does not evaluate the internal staffing structures of other platforms, StopAntisemitism observed a measurable improvement in TikTok's enforcement of its Community Guidelines following the establishment of this role.

Although this report does not attribute causation, the observed improvement suggests that dedicated expertise in hate speech policy may enhance communication with trusted reporting organizations and contribute to more consistent enforcement outcomes.

TikTok's comparatively robust moderation systems also appeared to influence user behavior. Throughout the reporting period, StopAntisemitism observed creators increasingly relying on coded language, altered spellings ("algospeak"), symbols, and visual cues to avoid automated detection while communicating substantially the same messages. Common examples of algospeak include terms such as "unalive" instead of "kill," "Notsee" instead of "Nazi," "Israhell" instead of "Israel," and deliberately altered spellings, euphemisms, or coded references designed to evade automated moderation while conveying the same underlying message.

The findings also suggest that TikTok employed a more conservative threshold for identifying policy violations. In multiple instances, content that was removed or resulted in enforcement on TikTok remained available on Meta or X despite conveying substantially similar messages.

Case Study: USER A

StopAntisemitism reported the same creator and content to both TikTok and Instagram.

Platform Outcomes

- TikTok: Permanently Removed the user from the platform.
- Instagram: Several posts were actioned, but the account remains active with approximately 741,000 followers.

Representative Examples Reported

- Encouraging attacks on civilians: "I know you've been holding back and not aiming at civilians.. it's time to go all in."
- Antisemitic wordplay: "Jewicial system."
- October 7 conspiracy theory: "October 7th was an inside job."

TikTok's role in shaping public opinion is particularly significant given its influence among younger audiences. According to recent research, 64% of Gen Z users have used TikTok as a search engine, while 43% of U.S. adults under 30 now regularly get news from the platform, up from just 9% in 2020. As TikTok increasingly functions as both an entertainment platform and a source of information, consistent enforcement of policies governing hate speech, terrorist promotion, and incitement becomes increasingly important.

Although TikTok recorded the strongest overall enforcement, challenges remained. Several permanently removed creators returned under new accounts, requiring additional reports before further enforcement occurred. These cases demonstrate that improvements in ban-evasion detection and enforcement consistency remain necessary.

Key Strengths

- Highest account removal rate among all three platforms.
- Demonstrated responsiveness to reports.
- Consistently pursued account-level enforcement against repeat offenders.
- Enforced policies against monetization, including TikTok Shop listings.
- Removed antisemitic sounds and other platform-specific content.
- Demonstrated a willingness to remove several documented ban evaders following additional reporting.

Areas for Improvement

- Expand proactive detection of repeat offenders before they rebuild large audiences.

Meta OVERALL ASSESSMENT

Instagram demonstrated a mixed approach to enforcement. The platform frequently removed individual posts, reels, and other content found to violate its Community Standards but was significantly less likely to permanently remove the accounts responsible for repeatedly posting that content. This approach reduced the visibility of individual violations while often allowing repeat offenders to remain active.

Compared to TikTok, Meta appeared to prioritize enforcement against content promoting or glorifying designated terrorist organizations and violence over antisemitic hate speech. While posts praising designated terrorist groups or encouraging violence were frequently removed, content expressing anti-Jewish hatred or relying on antisemitic tropes was less consistently actioned, particularly when presented as political commentary or misinformation.

During the course of this review, StopAntisemitism also observed Instagram's recommendation systems frequently surfacing additional content expressing similar extremist themes after researchers interacted with or documented antisemitic material. This observation is consistent with emerging research on algorithmic recommendation pathways. A 2026 study found that newly created Instagram accounts following mainstream wellness and fitness content were algorithmically recommended antisemitic conspiracy-related content within days, despite never actively searching for it.

These findings underscore that effective moderation requires more than removing individual posts. Recommendation systems should also minimize the repeated amplification of misinformation, conspiracy theories, and extremist narratives.

Case Study: USER B

StopAntisemitism reported User B to Instagram.

Platform Outcomes

- Three posts were removed but the account remained active.

Representative Examples Reported

Instagram deemed these violative and removed them:

- "Me at the WWII museum explaining to all the other people why it's mathematically impossible to be 6 million."
- "When you tell her, it was only 271k and she says yeah but 6 million would have been so much better."
- German shop bans Jewish customers "We're so back baby."

While this post was deemed non-violative:

- “Jesus watching his followers support the same people who crucified him.”

Key Strengths

- Frequently removed individual pieces of violating content.
- Permanently removed several significant repeat offenders.
- Demonstrated responsiveness to reports.
- Consistently enforced Community Standards against specific posts.

Areas for Improvement

- Escalate repeat offenders to account-level enforcement more quickly.
- Strengthen detection of ban evasion.
- Increase transparency surrounding enforcement outcomes.

X OVERALL ASSESSMENT

X demonstrated the weakest overall enforcement during the reporting period. Compared with TikTok and Meta, X removed the lowest percentage of reported accounts and most frequently determined that reported content did not violate its policies or constituted protected political expression. Although the platform occasionally removed individual posts or suspended accounts, account-level enforcement remained comparatively rare.

Case Study: USER C

StopAntisemitism reported User C to X.

Platform Outcomes

- Several reported posts were deemed political speech or commentary.
- One reported post was deemed violative and had visibility filter placed on it.
- The account remained active.

Representative Examples Reported

X deemed these posts to fall under political speech or commentary:

- “ Hamas was the moral and ethical one here, they cared for their hostages. No rape No torture.”
- “ You start by noticing Jews. You end by hating what they’ve done to every civilization.”
- “ Palestinians are being destroyed, stop calling them Hamas. Everyone is Hamas, because you’ve left us all no choice but to be! Freedom to Palestine! Strength to Hamas, Houthis, Hezbollah and any other H boys that fight Zionists!”

While this post had a visibility filter placed on it:

- “ Correct, exactly like Oct 7th, put people in a corner by force and death and they will find a way to fight back. Proud of them!”

X demonstrated the broadest interpretation of protected political speech among the platforms analyzed. Throughout the reporting period, StopAntisemitism received numerous enforcement decisions stating that reported content did not violate the platform’s policies because it constituted “political speech” or protected opinion.

Like all major social media platforms, X allows eligible creators to monetize engagement. Timely and consistent enforcement is therefore particularly important, as delays may allow harmful content to continue expanding its reach and, where applicable, generating revenue before moderation occurs.

Key Strengths

- Demonstrated that enforcement is possible in the most egregious cases.
- Willingness to demonetize egregious accounts.

Areas for Improvement

- Improve consistency when evaluating alleged terrorist promotion and antisemitic rhetoric.
- Increase transparency regarding enforcement decisions, outcomes, and policy interpretation.
- Increase account bans and demonetization for repeat offenders.

THE SOCK PUPPET PROBLEM

Throughout the reporting period, StopAntisemitism routinely revisited previously removed accounts to verify whether enforcement actions remained in place. In multiple instances, accounts that had previously been suspended or permanently removed were once again active on their respective platforms.

When this occurred, researchers documented the reinstatement and submitted follow-up reports to the platform, often resulting in additional enforcement. This process required ongoing monitoring well beyond the initial report and demonstrated that meaningful enforcement frequently depends on continued oversight.

The need to repeatedly verify that previously banned accounts remain inactive raises important questions about the durability of platform enforcement. If an account that was previously determined to violate a platform's policies later becomes active again, platforms should be able to identify how that occurred.

Key questions include:

- Was the account reinstated following an appeal?
- Was the suspension temporary rather than permanent?
- Did the user successfully evade enforcement by creating a new account or altering identifying information?
- Did automated systems fail to detect a repeat offender?
- Are enforcement decisions applied consistently regardless of an account's size, influence, or audience reach?
- Are there gaps in internal review or enforcement processes that allow repeat offenders to return?

Representative Examples

Throughout the reporting period, StopAntisemitism documented multiple instances in which permanently removed accounts later reappeared on the same platform or returned under newly created accounts.

USER D: Platform Tiktok

- Original Account: Permanently removed after repeated policy violations while the account had approximately 2.2 million followers.
- First Return: A new account was created and grew to approximately 66,000 followers before StopAntisemitism identified it, submitted additional reports, and the account was permanently removed once again.
- Second Return: The creator later returned under another account, which accumulated approximately 25,400 followers before another enforcement action was taken.
- Most Recent Status: At the conclusion of this reporting period, the creator had once again returned under a new TikTok account with approximately 50,000 followers, demonstrating the ongoing challenge of preventing repeat offenders from rebuilding audiences after enforcement.

USER E: Platform Instagram

- Original Account: The creator's original account was permanently removed at 500,000 followers.
- First Return: The original account was reinstated at 500,000 followers. Following additional reporting by StopAntisemitism, the account was banned once again.
- Second Return: The creator subsequently returned to Instagram under a new account, amassing approximately 70,000 followers within a few months. After additional reporting, that account was also permanently removed.
- Third Return: The replacement account was later reinstated and remained active until it exceeded 90,000 followers, at which point it was once again removed following further reporting.
- Most Recent Status: Since then, the creator has established several additional smaller accounts that continue to publish content appearing to violate Instagram's Community Standards.

USER F: Platform X

- **Original Account:** Following reporting by StopAntisemitism, the creator's account was suspended by X at 170,000 followers. The platform did not specify whether the suspension was permanent.
- **Reinstatement:** The creator's account was subsequently reinstated and resumed activity on the platform.
- **Additional Enforcement:** Following further reporting by StopAntisemitism, the creator was demonetized, preventing the account from generating revenue through X's creator monetization program.
- **Most Recent Status:** The account remains active. StopAntisemitism has not received confirmation as to whether the creator's monetization privileges remain suspended.



Recommendation 1: Expand Specialized Moderation Expertise

Major social media platforms should designate dedicated subject-matter experts within their Trust & Safety or Public Policy teams to oversee issues involving antisemitism and other forms of context-dependent hate.

This yearlong review demonstrated that some of the most difficult content to moderate was not the most explicit, it was the most contextual. Many of the examples submitted required historical and cultural knowledge that may not be readily identifiable through automated systems or general content review.

These results demonstrate the value of sustained subject-matter expertise. While automated systems and general moderation teams remain essential, platforms should consider establishing specialized trust and safety teams dedicated to antisemitism and other forms of context-dependent hate. These teams would be better equipped to recognize evolving rhetoric, misinformation, and coded language that may otherwise be overlooked.

Recommendation 2: Improve Detection of Repeat Offenders

Account removal is most effective when platforms can prevent repeat offenders from rapidly rebuilding their audiences.

Platforms should continue improving systems that identify potential ban evasion using behavioral signals, linked accounts, repeated content, and other indicators beyond usernames alone.

Recommendation 3: Evaluate Recommendation Systems

Platforms should regularly evaluate whether recommendation algorithms are unintentionally amplifying increasingly extremist or conspiratorial content.

Recommendation systems are inherently difficult to evaluate because they rely on highly personalized algorithms that continuously adapt to user behavior. However, when human reviewers, trusted reporting organizations, or internal investigations identify clusters of violating content repeatedly appearing within recommendation

pathways, those findings should be used to improve the underlying recommendation models. In other words, moderation decisions should not only remove individual pieces of content, they should also help train systems to reduce the future amplification of similar content.

Recommendation 4: Increase Transparency

Governments routinely publish detailed statistics on hate crimes to help policymakers, researchers, and the public understand trends, allocate resources, and evaluate prevention efforts. Online hate should be measured with a similar level of transparency.

Platforms should consider publishing annual transparency reports that include:

- The number of reports submitted involving alleged antisemitic content.
- The percentage of reports involving alleged antisemitic content compared to content involving other forms of hate and bigotry.
- The percentage of those reports resulting in account-level enforcement.

Recommendation 5: Develop Meaningful Subject-Matter Advisory Panels

Many forms of online hate evolve rapidly. Antisemitic rhetoric today often differs significantly from that of even a few years ago. All social media platforms should establish ongoing advisory relationships with modern day experts specializing in antisemitism, extremism, terrorism, and other forms of identity-based hate. Rather than relying solely on periodic consultations, these experts should provide routine feedback on emerging trends, evolving terminology, and recurring enforcement challenges, helping moderation policies adapt as online hate evolves.

The importance of this work extends beyond social media platforms themselves. Online content increasingly informs AI systems, and AI-powered tools are now widely used throughout K-12 education and higher education. As a result, antisemitic content that remains online has the potential to influence not only human audiences but also the technologies increasingly relied upon for learning, research, and information retrieval.

SECTION VII **CONCLUSION**

Over the course of one year, this review documented 126 accounts, 762 examples of alleged policy-violating content, and creators representing more than 16.3 million followers, with an estimated potential reach of 50 to 100 million users. These efforts resulted in 33 permanent account removals, representing approximately 4.5 million followers, in addition to dozens of content removals and other enforcement actions. The cumulative effect extends beyond the removal of individual posts or accounts by disrupting established channels capable of amplifying harmful content to millions of users.

The scale of that potential distribution is significant. Even under a conservative illustrative scenario, if just 1% of those 16.3 million followers shared a post, approximately 163,000 shares would result. If each share reached only 50 additional users, those shares alone could generate more than 8 million additional impressions, illustrating how quickly harmful content can spread across social media ecosystems.

The findings of this report demonstrate that meaningful moderation is possible. Across all three platforms, documented reports resulted in account removals, content removals, and other enforcement actions. The challenge is not whether platforms are capable of enforcing their policies, but whether those policies are applied consistently, if at all.

Perhaps the clearest lesson from this review is that specialized expertise matters. As online hate continues to evolve through coded language, memes, historical revisionism, and emerging technologies, effective enforcement requires ongoing subject-matter expertise alongside transparent reporting and durable enforcement systems.

As online platforms continue to shape public discourse and increasingly influence AI systems and educational technologies, the consequences of inconsistent moderation extend well beyond social media itself. The question is no longer whether meaningful enforcement is possible. This report demonstrates that it is.

About Us

StopAntisemitism is a grassroots watchdog organization dedicated to exposing groups and individuals who espouse incitement against the Jewish people and the State of Israel and engage in antisemitic behaviors.

Founded in 2018, StopAntisemitism was born in response to increasing antisemitic violence and sentiment across the United States.

StopAntisemitism has developed a commanding following, reaching over half a billion individuals monthly through social media, mailers, and its website. By publicly exposing antisemites, StopAntisemitism has created an environment where those who propagate hatred against the Jewish people are met with real-world consequences, including, but not limited to, arrests, job loss, school expulsions, and awards being revoked.



@StopAntisemites



@Stop_Antisemitism



StopAntisemitism



StopAntisemitismorg



StopAntisemitism.org



info@stopantisemitism.org