

FILED UNDER SEAL

EXHIBIT 21

Message

From: [REDACTED] [/O=EXCHANGEELABS/OU=EXCHANGE ADMINISTRATIVE GROUP (FYDIBOHF23SPDLT)/CN=RECIPIENTS/CN=[REDACTED]]
Sent: 8/13/2020 3:09:56 AM
To: Vaishnavi J [REDACTED]@fb.com]; [REDACTED] [REDACTED]@fb.com]; [REDACTED] [REDACTED]@fb.com]; [REDACTED]
[REDACTED]@fb.com]; [REDACTED] [REDACTED]@fb.com]; [REDACTED] (CI) [REDACTED]@fb.com]; [REDACTED]
[REDACTED]@fb.com]; [REDACTED] [REDACTED]@fb.com]; Diego San Martin [REDACTED]@fb.com]; [REDACTED]
[REDACTED] [REDACTED]@fb.com]; [REDACTED] [REDACTED]@fb.com]; [REDACTED] [REDACTED]@fb.com]; [REDACTED]
[REDACTED]@fb.com]; [REDACTED] [REDACTED]lly@fb.com]; [REDACTED] [REDACTED]@fb.com]; [REDACTED]
[REDACTED]@fb.com]; [REDACTED] [REDACTED] [REDACTED]@fb.com]; [REDACTED] [REDACTED]@fb.com]; [REDACTED]
[REDACTED]@fb.com]
CC: [REDACTED] [REDACTED]@fb.com]; [REDACTED] [REDACTED]@fb.com]; [REDACTED] [REDACTED]@fb.com]; [REDACTED]
[REDACTED] [REDACTED]@fb.com]
Subject: Re: Heads-up and request for advice: CEI on IG

PRIV

From: [REDACTED]
Date: Wednesday, August 12, 2020 at 8:03 PM
To: Vaishnavi J, [REDACTED] (CI)", [REDACTED]
 [REDACTED] Diego San Martin, [REDACTED]
 [REDACTED]
Cc: [REDACTED]
Subject: Re: Heads-up and request for advice: CEI on IG

PRIV

From: Vaishnavi J
Date: Wednesday, August 12, 2020 at 5:45 PM
To: [REDACTED], [REDACTED], [REDACTED], [REDACTED], [REDACTED] (CI)", [REDACTED], [REDACTED]
[REDACTED], Diego San Martin, [REDACTED], [REDACTED], [REDACTED], [REDACTED], [REDACTED]
[REDACTED], [REDACTED], [REDACTED], [REDACTED], [REDACTED], [REDACTED]
Cc: [REDACTED], [REDACTED], [REDACTED], [REDACTED], [REDACTED]
Subject: Re: Heads-up and request for advice: CEI on IG

Hi [REDACTED]

Thanks for sharing this, and I agree that the two measures you've called out would be helpful here to mitigate some of the eSafety Commissioner's concerns. @██████████ and @██████████ are the best folks to share more about the implicit minor sexualisation policy. There is an experiment currently running on the profile-level reporting option and a final writeup with recommendations will be shared at the end of the week - @██████████ (CI) and @██████████ may have more details on how we follow up from this experiment and next steps.

When do you anticipate that the eSafety Commissioner will want to meet?

CONFIDENTIAL
CONFIDENTIAL

META3047MDL-034-00106792
METANMAG-002-02073003

Cc: [REDACTED]
Subject: Re: Heads-up and request for advice: CEI on IG

Thank you @ [REDACTED] and @ [REDACTED] for the really detailed update and explanation. It is really helpful to have an understanding of where some of the current gaps are in our system (but also reassuring knowing that there is so much work going on in the backend to address these). Please let me know if there's anything I can offer from the LE Outreach perspective to support the great work you're doing!

Cheers,

On 11/08/2020, 11:51, '██████████' <██████████@fb.com> wrote:

+1, well summarized, @

Already happening (calibrate IICT2 classifier)

- Calibrating the existing IIC2 classifier on IGDMs. This means we'll be able to identify users who are at various likelihoods of being violating. This unblocks work like "block X functionality for users who are Y+% likelihood of being IIC bad actors"

Potential next steps (improve coverage/performance)

- Improve IIC2 classifier on IGDMs. This requires ongoing labeling capacity for us to gather new training data, iterate on the classifier, and recalibrate the improved classifier. Upstream detection improvements (i.e. age determination work from core dimensions) will potentially have the same impact (improved classifier performance, but requiring re-calibration)
 - Specific next steps if we want to improve classifier performance here are somewhat unclear until we get the results of this calibration completed, but we currently don't have available reviewers so we'd need to reprioritize within existing review requests or we'd need to secure additional review capacity here. For context, we're currently also not doing ongoing reviews for IIC2 in Messenger
- Start detecting/enforcing on IIC outside of Messenger and IGDMs. We're still ironing out policy guidance on what we want to do with sexualizing engagement from adults to children in posts/comments. This is also something that has fallen below the cusp for eng prioritization. While we think this will help with enforcing on people who do violations in messages, we don't think this takes priority over the work we've committed to for H2 (which is more CEI E2EE focused)

Cheers,

From: [REDACTED]
Sent: Monday, August 10, 2020 6:36 PM
To: [REDACTED]; Vaishnavi J.; [REDACTED] (CI); [REDACTED] Diego San Martin;
[REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED];
; [REDACTED]; [REDACTED]
Cc: [REDACTED]
Subject: Re: Heads-up and request for advice; CEI on IG

There is, in fact, a gap in IIC proactive detection in IGD threads vs. MSGR. We're currently working on correctly calibrating the IIC classifier for use on IGD and it's scheduled to wrap this week or next in time for Interop country tests. While this will likely improve its performance on IGD, it will likely still not perform as well as it does on MSGR. @██████████ can speak in more detail about what it might take to gain parity in performance beyond calibration.

Hi all

Thanks for your recent advice about the trend we are seeing of an uptick of CSAM on IG.

I wanted to add in @ [REDACTED] from our local LE team, as she is also seeing a large increase in CSAM material that LE is reporting on IG (some of which has been reported a large number of times before being escalated to us and hasn't been taken down).

- Some of this thread are also aware that we also recently met with Collective Shout who provided more than 500 accounts they had reported over the last few months. I know @Diego and others are progressing the specific investigation but thought helpful to surface for this group's information.

Would appreciate any update on the actions discussed earlier on this thread – profile-level reporting and other items. If there is anything we can do to help, please don't hesitate to let us know.

Cheers

[REDACTED]

From: Vaishnavi J

Date: Saturday, 20 June 2020 at 4:21 am

To: '[REDACTED] (CI)', [REDACTED] Diego San Martin, [REDACTED]
[REDACTED]
[REDACTED]

Cc: [REDACTED]

Subject: Re: Heads-up and request for advice: CEI on IG

PRIV

From: [REDACTED] (CI)

Sent: 19 June 2020 4:38 AM

To: [REDACTED]; Vaishnavi J; [REDACTED] Diego San Martin; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]

Cc: [REDACTED]; [REDACTED]; [REDACTED]

Subject: Re: Heads-up and request for advice: CEI on IG

@Vaishnavi J IGPR is the most significant reporting gap we know of. On whether that would be sufficient signal for us to prioritise without content reports, we'd need to test precision to be sure - we should definitely do that exercise, and should ideally do the same for comments (where there is also a reporting gap)

From: [REDACTED]

Sent: Friday, June 19, 2020 12:35 AM

To: Vaishnavi J; [REDACTED]; [REDACTED] (CI); Diego San Martin; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]

Cc: [REDACTED]; [REDACTED]; [REDACTED]

Subject: Re: Heads-up and request for advice: CEI on IG

PRIV

a couple of days. These appear to be pieces of content that are very clearly violating. I wonder if this means there is an AI / proactive detection issue.

- Secondly, they are concerned about the amount of content they are seeing that represents benign photos of minors that are sexualised in the comments. Given that we have a working group being led by the amazing Ayesha on this front, it sounds like this is probably a policy issue that is about to be imminently resolved (cc @ [REDACTED] FYI).

We asked Market Operations for the data and – although it's not currently possible to break down user reports by market because of the way that content moderation for CEI has been rejigged during COVID – it does appear that there has been an increase in the number of reports we've received on CEI from eSafety and law enforcement (see below).

I asked eSafety what they thought was driving the increase. They did admit it was partially due to a campaign from Collective Shout (...) but they have indicated there genuinely are a number of very explicit pieces of content that users have seen and reported. They are adamant it is significantly worse on IG – there's no data to prove it because they only sent me two specific examples and one was FB and one was IG – but they are convinced it is an IG problem. I will say the two pieces of content they sent me were both genuinely graphic – not borderline cases at all.

I know this is fairly ambiguous but I wanted to send up the flag to see if there's any reason we would be aware of the AI on IG being any less accurate on proactively detecting CEI?

Appreciate the advice and very happy to jump on a call to discuss if helpful – cheers

[REDACTED]

Stats on CEI reports by month

Month	e-Safety Commissioner	LERT escalation contact form
January	3 tickets (5 contents reviewed)	13
February	7 tickets (12 contents reviewed)	23
March	5 tickets (7 contents reviewed)	47
April	10 tickets (20 contents reviewed)	35
May	11 tickets (19 contents reviewed)	44

–

[REDACTED] [REDACTED]

Head of Public Policy, Australia

FACEBOOK

[REDACTED] | [REDACTED]@fb.com