

AMENDED Exhibit 243

PLAINTIFFS' OMNIBUS OPPOSITION TO DEFENDANTS' MOTIONS FOR SUMMARY JUDGMENT

Case No.: 4:22-md-03047-YGR

MDL No. 3047

In Re: Social Media Adolescent Addiction/Personal Injury Products Liability Litigation

Message

From: [REDACTED] [/O=THEFACEBOOK/OU=EXCHANGE ADMINISTRATIVE GROUP (FYDIBOHF23SPDLT)/CN=RECIPIENTS/CN=[REDACTED]]
Sent: 9/13/2021 7:09:40 PM
To: Miki Rothschild; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]
CC: [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; [REDACTED]; Sayed Otaru
Subject: Re: Sensitive Content Controls and Safe Search

Hi Miki,

On #3 we've done some early experimentation with the Drebbel team on [additional enforcement for users who we detect to be in a rabbithole](#) on Feed Recs. TL;DR we did see a meaningful drop in exposure with targeted demotion, but it came with a clear engagement cost and the effects did not last after intervention.

This was a learning experiment - we would need to invest further in order to refine for a launch candidate including Eng/DS lift, driving PXFN alignment around rabbit hole definition & allowing stronger enforcement, as well as negotiating budget to support engagement regressions from launching at scale.

Longer term we think the right holistic solution to this & similar challenges will be evolving toward integrity-aware engagement metrics, reducing the engagement vs integrity tension by better aligning incentives upstream by focusing on good engagement. We are kicking off discussion between RI & AI Integrity team on this as part of our broader Integrity Personalization workstream.

thanks,

From: Miki Rothschild [REDACTED]
Sent: Sunday, September 12, 2021 2:26 PM
To: [REDACTED]
Cc: [REDACTED] Sayed Otaru
Subject: Sensitive Content Controls and Safe Search

Hey folks, with the recent PR fire TikTok has seen around inappropriate content for teens (see [here](#)), we were chatting within Youth XFN about our own deficits when it comes to discoverability, namely: (1) it's still very easy to find borderline and sometimes violating content/accounts in search, e.g. try typing "drug" or "love" (2) if you start following borderline accounts or interacting with such content, our recommendations algorithms will start pushing you down a rabbit hole of more egregious content.

Long term, the solution here, which we've discussed in the past, is to start measuring both actual and synthetic (based on stress tests simulation of both queries and interactions) *discoverability* of both violating and non-recommendable content/accounts, and then drive that down similarly to what we did for *prevalence* or such content. Once we have this type of measurement in place, we'll find the various root causes, which will lead to the right types of projects. My guess is we need to think deeply about how we treat such seeds in training data and optimization functions for recommendations, and will also need to think more deeply about search intents (I can share the journey from FB Search).

Short term, I think we should revisit the priority of implementing "safe search", as part of our Sensitive Content Control "Limit More" state. Specifically, I'm thinking that in "Limit More" state we could filter non-recommendable tiers from search as opposed to what we do now which is to down rank them. Additionally, I'd be interested to explore some short term solutions to the rabbit holing problem, one idea I had was to drop our precision thresholds when we are confident that someone already follows non-recommendable accounts/content, since we know we are more likely to recommend such content at that point.

Would love your thoughts on the above ideas and broader context of where we are on this problem space. In particular:

1. What's the current state of discoverability measurement? How are we thinking about implementing this and by when? [REDACTED]
2. What do you think about adding Safe Search into Sensitive Content Control "Limit More" mode? Can we tee this up as proposal? [REDACTED]
3. How are we thinking about non-rec/violating Rabbit holes in recommendations? Do we have any specific approaches to prevent this? [REDACTED]

Thanks!
Miki