

AMENDED Exhibit 240

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=EXTERNAL
(FYDIBOHF25SPDLT)/CN=RECIPIENTS/CN=CDBDC555C0894B67BAFC70709D3FBA31]
Sent: 4/22/2020 12:25:01 AM
To: [REDACTED]; Karina Newton [REDACTED]; [REDACTED]; [REDACTED]
Subject: Message summary [{"otherUserFbId":null,"threadFbId":2346200088816187}]
Attachments: 33434104_2326762057350551_2134873477990055936_n.gif

[REDACTED] (4/21/2020 16:16:36 PDT):

>Hey product friends. We just chatted on the policy side about sensitive content control. We think it makes sense to draw the line at likely violating with the agreement that the control will *not* have any impact on entity-level filtering we have in place. This would mean that our account-level enforcement for SSI (e.g. you are not Explore eligible if you've been checkpointed in the prior 30 days for SSI, including for non-violating admission content) *should* solve the SSI issue -- does this sound right from a product standpoint? Is there an additional "general SSI" filter we use today on Explore that I might not be aware of? That means the only remaining open Q would be holocaust denial, which as will mentioned in his prior email is more of an academic debate given we dont have enforcement for it today.

[REDACTED] (4/21/2020 16:19:02 PDT):

>If this sounds right, we're going to socialize with Monika first to get her alignment

[REDACTED] (4/21/2020 16:21:35 PDT):

>Does anyone have flags with the accuracy of the above? Or other flags before we chat with monika?

[REDACTED] (4/21/2020 16:23:28 PDT):

>we do have an SSI filter on Explore, I'm trying to confirm how significant it is

[REDACTED] (4/21/2020 16:23:57 PDT):

>Thanks - and it's a non-violating filter?

[REDACTED] (4/21/2020 16:24:30 PDT):

>In general, I'm on board with likely violating being the line. I would appreciate your support in advocating for the iterative approach to get there vs. the cold turkey approach.

[REDACTED] (4/21/2020 16:24:48 PDT):

>If you're going to respond on that thread

[REDACTED] (4/21/2020 16:24:59 PDT):

>Yes, agree on the iterative approach -- we definitely can't cold turkey without proper violating classifiers

[REDACTED] (4/21/2020 16:25:14 PDT):

>not going to respond on the thread yet - first we're going to socialize with Monika

[REDACTED] (4/21/2020 16:25:38 PDT):

>this line is more "north star" than "day one"

[REDACTED] (4/21/2020 16:26:11 PDT):

>Ok. Please say that on the thread when someone responds bc i think the current convo is conflating the two.

[REDACTED] (4/21/2020 16:27:20 PDT):

>One comment: I wonder if there is a better way to frame the entity level thing you are trying to get at. I'm worried that it won't resonate with Adam bc I thinks we should have a higher bar for entity level stuff and give people more control bc it's more of a blunt instrument. I see what you are trying to do but maybe there is a different way to frame it around CO actions vs. entity level stuff. For example, account tiering takes into account CO actions but also non-rec classifiers.

[REDACTED] (4/21/2020 16:28:05 PDT):

>Yes, that's a good point. Are we going to consider allowing people to opt out of filtering based on tiering though?

[REDACTED] (4/21/2020 16:28:27 PDT):

>I can say "entity level filtering based on CO actions"

[REDACTED] (4/21/2020 16:28:30 PDT):

>to be more precise

[REDACTED] (4/21/2020 16:28:38 PDT):

>but again, not responding on the thread yet to be clear

[REDACTED] (4/21/2020 16:29:25 PDT):

>Under what policy does CO action these things? It does muddle the line a bit.

[REDACTED] (4/21/2020 16:29:44 PDT):
>SSI Admission

[REDACTED] (4/21/2020 16:30:31 PDT):
>This content has sensitivity screens on it right? If so, that could be another way to draw the line.

[REDACTED] (4/21/2020 16:30:32 PDT):
>What we're saying is at the CONTENT level, you can opt into seeing non-rec stuff so long as it is NOT likely violating. However, the control does NOT allow you to opt in to see any content from entities that are being filtered at the account level

[REDACTED] (4/21/2020 16:30:58 PDT):
>so the distinction is that the control applies strictly to content level filtering, not entity level filtering

[REDACTED] (4/21/2020 16:31:40 PDT):
>I'm afraid if we get into details on the entity piece, we'll derail again. So if we can agree on this as the scope of the control today, we can revisit in the future if we want to bring entity level in scope and, if so, what lines we draw

[REDACTED] (4/21/2020 16:31:46 PDT):
>Is the basis of the distinction that the entity filter is a strikes-based penalty. like how these accounts also can't go live?

[REDACTED] (4/21/2020 16:32:08 PDT):
>I get that but the reason we do the entity filtering is to filter out more content so im not sure that argument holds up. since it's more of a blunt instrument too.

[REDACTED] (4/21/2020 16:32:09 PDT):
>it's not really strikes based -- SSI admission deletions don't give strikes, but it's similar

[REDACTED] (4/21/2020 16:32:32 PDT):
>The idea that we would do entity enforcement based on take downs of violating content seems fine to me bc that's the violating bar.

[REDACTED] (4/21/2020 16:32:33 PDT):
>Ok, so if we don't think entity level is a good option, then we're going to need to discuss further

[REDACTED] (4/21/2020 16:32:53 PDT):
>it's not a violating bar though

[REDACTED] (4/21/2020 16:32:56 PDT):
>But it's the entity filtering that's at some higher bar that seems a bit tricky to rationalize

[REDACTED] (4/21/2020 16:32:56 PDT):
>SSI Admission is not violating

[REDACTED] (4/21/2020 16:32:58 PDT):
>that's the problem

[REDACTED] (4/21/2020 16:33:10 PDT):
>Yeah. Is SSI Admission the only example here?

[REDACTED] (4/21/2020 16:33:18 PDT):
>It's the one we're trying to solve for, yes

[REDACTED] (4/21/2020 16:33:26 PDT):
>So entity level filtering based on CO action

[REDACTED] (4/21/2020 16:33:30 PDT):
>could be the line

[REDACTED] (4/21/2020 16:33:40 PDT):
>we don't allow you to opt in to seeing entities we've filtered based on a CO action

[REDACTED] (4/21/2020 16:34:05 PDT):
>but if it's based on something else, like high volume of non rec content, theoretically, we'd allow opt out of that

[REDACTED] (4/21/2020 16:34:10 PDT):
>I'm willing to go for it. Just telling you I see the hole in the argument. But might be the best one.

[REDACTED] (4/21/2020 16:34:23 PDT):
>What is the hole if we use the above ^

[REDACTED] (4/21/2020 16:34:52 PDT):

>That the filtering by CO is done with a bar higher than the violating/likely violating bar.

[REDACTED] (4/21/2020 16:36:13 PDT):

>Yes so we could say "we allow you to opt out of filtered content under our Rec policy so long as it is not likely violating NOR CO actioned (e.g. MAS or MAD screen). The same applies to entity-level filtering (e.g. if your account was recently SSI checkpointed for admission content, it will remain filtered):

[REDACTED] (4/21/2020 16:36:31 PDT):

>Does this close the gap?

[REDACTED] (4/21/2020 16:37:55 PDT):

>the line: you can't opt of action taken under our Community Guidelines

[REDACTED] (4/21/2020 16:38:03 PDT):

>opt out of*

[REDACTED] (4/21/2020 16:40:36 PDT):

>I think that's a fine line to take. It's a bit complex when you dig in but not sure anyone will.

>

>The other options are:

>(1) Make the bar something a bit higher than "likely violating" -- more like "borderline" where it's a bit higher than our violating line but not significantly higher like much of our rec guidelines -- and make the case that SSI admission is very close to the violating content in this category.

>(2) Let people control SSI admission. Adam might be ok with that in a world where the default is that it's out and I think it's in the spirit of what we committed to.

[REDACTED] (4/21/2020 16:41:11 PDT):

>borderline is what we'll be doing almost all filtering on -- borderline nudity = sexually suggestive

[REDACTED] (4/21/2020 16:41:24 PDT):

>so if we dont allow people to opt out of that then they essentially get to opt out of nothing today

[REDACTED] (4/21/2020 16:42:07 PDT):

>Maybe that's not the right word to use in this context. In the nudity example I mean something that is so close to nudity but is cut off or covered up in the right place.

[REDACTED] (4/21/2020 16:42:25 PDT):

>That'd mean we need to create a third policy line

[REDACTED] (4/21/2020 16:42:30 PDT):

>and define it

[REDACTED] (4/21/2020 16:42:33 PDT):

>and build classifiers for it

[REDACTED] (4/21/2020 16:43:08 PDT):

>That's basically what our classifiers for violating would catch right?

[REDACTED] (4/21/2020 16:43:23 PDT):

>?

[REDACTED] (4/21/2020 16:43:34 PDT):

>So how is it different from "likely violating"?

[REDACTED] (4/21/2020 16:43:41 PDT):

>AND how do we explain it externally?

[REDACTED] (4/21/2020 16:44:18 PDT):

>I guess in the case of SSI we would have to say admission is so close to violating that it falls into the "likely violating" bucket even if it's not violating.

[REDACTED] (4/21/2020 16:44:26 PDT):

>it seems pretty straightforward to me to say: We filter things based on whether they are likely violating our Community Guidelines or have been actioned under our Community Guidelines. YOU cannot opt out of that filtering. We also filtering things based on this Rec policy. You can opt out of those.

[REDACTED] (4/21/2020 16:45:10 PDT):

>This feels more straightford than the path I led us down.

[REDACTED] (4/21/2020 16:45:23 PDT):

>I'm also a big fan of this phrasing

[REDACTED] (4/21/2020 16:45:58 PDT):

>Can you point to somewhere in the public facing community guidelines that talks about SSI admission?

[REDACTED] (4/21/2020 16:47:32 PDT):

>https://www.facebook.com/communitystandards/suicide_self_injury_violence

[REDACTED] (4/21/2020 16:49:11 PDT):

>This is when we checkpoint an account: "We provide resources to people who post written or verbal admissions of engagement in self injury, including:
>-Suicide
>-Euthanasia/assisted suicide
>-Self-harm
>-Eating disorders"

[REDACTED] (4/21/2020 16:50:25 PDT):

>+1, I think this framing makes sense.

>
>In terms of SSI risk if we were to do this, we do have a non-rec SSI filter in place on Explore today targeting depiction. It's an old classifier with low precision and low filter volume so unlikely to materially change ecosystem but may result in some SSI admission content appearing if we removed it

[REDACTED] (4/21/2020 16:52:24 PDT):

>Given that we wouldn't go cold turkey, however, I'd expect we'd wait to remove this until we have likely violating mitigations in place

[REDACTED] (4/21/2020 16:52:41 PDT):

>"We provide resource to people...and btw prevent their content from getting distribution"

[REDACTED] (4/21/2020 16:54:00 PDT):

>ok so generally are people aligned with this approach? I'm going to let people vote via thumbs up on the options:

[REDACTED] (4/21/2020 16:54:01 PDT):

>aligned

[REDACTED] (4/21/2020 16:54:02 PDT):

>not aligned

[REDACTED] (4/21/2020 16:54:08 PDT):

>haha

[REDACTED] (4/21/2020 16:54:21 PDT):

>(they created a workchat polling option for this ^)

[REDACTED] (4/21/2020 16:55:28 PDT):

>The more I think about this, the more I feel like the right approach is to just not filter SSI admission for people with the setting off. I'm not going to fight that battle but a pure version of the explanation is the only one that really lands well to me.

Karina Lynn Newton (4/21/2020 16:55:44 PDT):

>i am really not ok with that

[REDACTED] (4/21/2020 16:55:53 PDT):

>I agree that the option of setting the bar higher than "likely violating" is the worst option

Karina Lynn Newton (4/21/2020 16:55:59 PDT):

>and this is my problem with describing this as about borderline content

Karina Lynn Newton (4/21/2020 16:57:09 PDT):

>i think the holocaust denial example is an academic example that applies to very little content

[REDACTED] (4/21/2020 16:57:37 PDT):

>I feel like you're just making the case to not provide this control then. Which is an argument that also resonates with me. I'm just worried that we are rationalizing this strange middle ground that will be too hard to defend.

[REDACTED] (4/21/2020 16:57:58 PDT):

>that being said, I will fully back your proposal outside of this thread.

Karina Lynn Newton (4/21/2020 16:58:05 PDT):

>ha

Karina Lynn Newton (4/21/2020 16:59:21 PDT):

>my issue is this: people who are suffering from depression and self-harm go down IG rabbit holes, and explore functionality compounds this issue. this is why adam published a national op-ed in the uk to say we won't recommend ssi content anymore. this isn't an academic example. and it's not responsible to give young people a way to opt out of this.

[REDACTED] (4/21/2020 17:00:35 PDT):

>I'm also totally fine not providing this setting for people under 18 if that helps. Issues could be that we don't have age for many and could create incentives to lie about age.

[REDACTED] (4/21/2020 17:00:57 PDT):

>I'd actually let you make the call there independent of this other convo

[REDACTED] (4/21/2020 17:01:02 PDT):
>If you have an opinion

[REDACTED] (4/21/2020 17:01:10 PDT):
>It's something ive brought up with [REDACTED]

Karina Lynn Newton (4/21/2020 17:01:49 PDT):
>we also only have age for 55ish percent of people

Karina Lynn Newton (4/21/2020 17:02:18 PDT):
>but we will have new teen/non-teen classifiers we could use to gate this control

Karina Lynn Newton (4/21/2020 17:02:20 PDT):
>in july

Karina Lynn Newton (4/21/2020 17:03:02 PDT):
>but i like something on age, despite the incentives to lie

[REDACTED] (4/21/2020 17:03:21 PDT):
>Not sure we should do it on a classifier and make the classifier prediction so obvious in the product. But totally fine with doing it on stated age

[REDACTED] (4/21/2020 17:03:59 PDT):
>I totally agree with you. Do you think people who are struggling are so likely to want to go down this rabbit hole that they will see out the setting and turn it off?

Karina Lynn Newton (4/21/2020 17:04:30 PDT):
>yes, sadly

[REDACTED] (4/21/2020 17:05:08 PDT):
>hmm, ok, if we feel so strongly about this one issue, i agree that [REDACTED] proposal is the best one.

[REDACTED] (4/21/2020 17:05:28 PDT):
shared: 33434104_2326762057350551_2134873477990055936_n.gif

[REDACTED] (4/21/2020 17:05:33 PDT):
>:)

[REDACTED] (4/21/2020 17:05:37 PDT):
>To state on the record, I still remain more in the camp of either doing this and being ok with the implications being kind of extreme or not doing it at all.

[REDACTED] (4/21/2020 17:07:04 PDT):
>I see everyone trying to water it down so much. Do you all want to actually fight for not doing it?

[REDACTED] (4/21/2020 17:07:35 PDT):
>I dont actualyl think it's watering it down to say that the topics we cover in CG are not opt-out-able

Karina Lynn Newton (4/21/2020 17:07:36 PDT):
>i think it's a non-starter to argue back on feed

Karina Lynn Newton (4/21/2020 17:07:56 PDT):
>and i think why this gets tricky is because we're not doing separate controls

Karina Lynn Newton (4/21/2020 17:08:08 PDT):
>which isn't to say that wasn't the right call

Karina Lynn Newton (4/21/2020 17:08:13 PDT):
>but it puts additional pressure

[REDACTED] (4/21/2020 17:08:13 PDT):
>You will be able to opt out of a ton of stuff: borderline nudity, borderline drugs, borderline guns, borderline spam, etc

[REDACTED] (4/21/2020 17:08:27 PDT):
>those things are all contentious

[REDACTED] (4/21/2020 17:08:35 PDT):
>and we're giving a user the control

[REDACTED] (4/21/2020 17:08:36 PDT):
>When I say not doing it at all, I mean making the control just about downranking in Home.

[REDACTED] (4/21/2020 17:09:07 PDT):
>And most of our "shadowban" claims are about those topics, not about SSI admission...

[REDACTED] (4/21/2020 17:12:40 PDT):
>BTW, I was wrong earlier. [REDACTED] helped confirm w/ eng on the SSI team and the SSI filter we're using is a violating classifier. So there would be no change in SSI classifiers on Explore if we took this approach

[REDACTED] (4/21/2020 17:12:54 PDT):
>HOORAY

[REDACTED] (4/21/2020 17:13:36 PDT):
>for future state, worth talking about how this would work for the future SSI borderline classifier? (it was planned for H2'20 might be pushed back due to COVID)

[REDACTED] (4/21/2020 17:13:40 PDT):
>I'm fine trying to land [REDACTED]'s proposal and the iterative approach (vs. cold turkey). If that doesn't land, let's reevaluate.

Karina Lynn Newton (4/21/2020 17:15:54 PDT):
>So...

>
>Option 1:
>Ineligible for opt-out: Content actioned by CO or eligible to be actioned by CO [MAD, MAS, SSI, or determined by classifier to be likely violating our community guidelines]

>
>Option 2:
>Option 1 + people under 18 don't get the control

[REDACTED] (4/21/2020 17:16:15 PDT):
>I like Option 2

Karina Lynn Newton (4/21/2020 17:22:25 PDT):
>i think option 2 releases some policy pressure and is a place that can help guy get to comfortability without overcomplicating

[REDACTED] (4/21/2020 17:25:01 PDT):
>the incentive to lie about age is the main concern I could see about restricting the control based on age.
>
>any other concerns we can think of or is it a no-brainer?