

AMENDED Exhibit 388

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

METAMAAG-001-00076377

Metadata

All Custodians	[REDACTED]	SEMANTIC
Confidentiality	HIGHLY CONFIDENTIAL	SEMANTIC
Filename	Product Growth HPM Integrity Jun 4, 2021	SEMANTIC
From	[REDACTED]@fb.com>	SEMANTIC
MDL BEGBATES	META3047MDL-003-00075698	SEMANTIC
MDL ENDBATES	META3047MDL-003-00075707	SEMANTIC



Message

From: [REDACTED] [/O=THEFACEBOOK/OU=EXCHANGE ADMINISTRATIVE GROUP (FYDIBOHF23SPDLT)/CN=[REDACTED]]
Sent: 6/4/2021 4:07:56 PM
To: [REDACTED]
BCC: [REDACTED]

[REDACTED] Denise Moreno [REDACTED]

[REDACTED] Shilpa Mody [REDACTED]; Diego Castaneda [REDACTED]
 [REDACTED]
 [REDACTED]

Subject: Product Growth HPM | Integrity | Jun 4, 2021

Attachments: image (86).png; Screen Shot 2021-06-04 at 8.33.22 AM.png; Screen Shot 2021-05-07 at 6.28.20 PM.png; image (87).png; Screen Shot 2021-06-04 at 8.38.07 AM.png; image (88).png; image (89).png; image (91).png

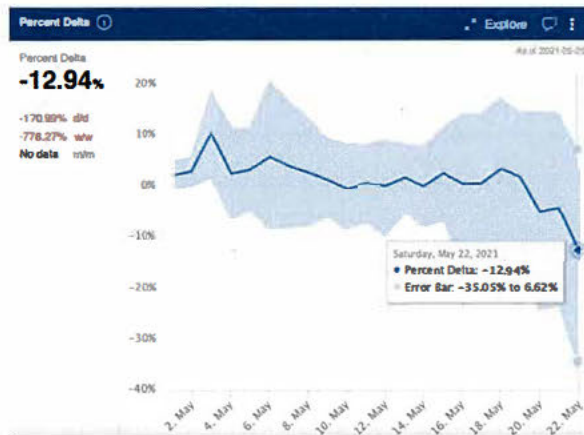
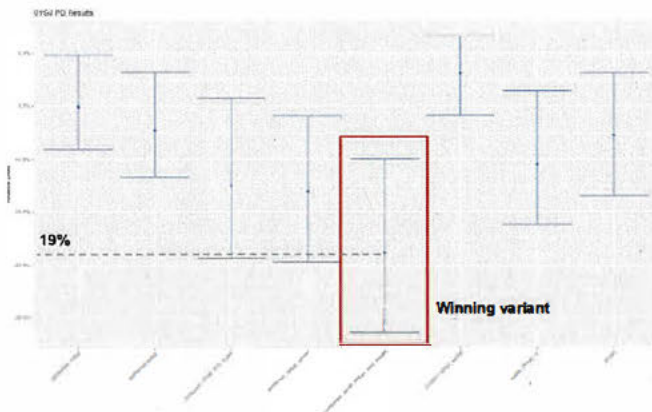
With the long Memorial Day weekend, we've combined last week & this week into one mega HPM! Sending for Ferrari while he's on PTO.

Highlights

[REDACTED] Recommendation Integrity] Demoting borderline Groups on GYSJ reduced GYSJ non-rec prevalence by 19% (+ [REDACTED])

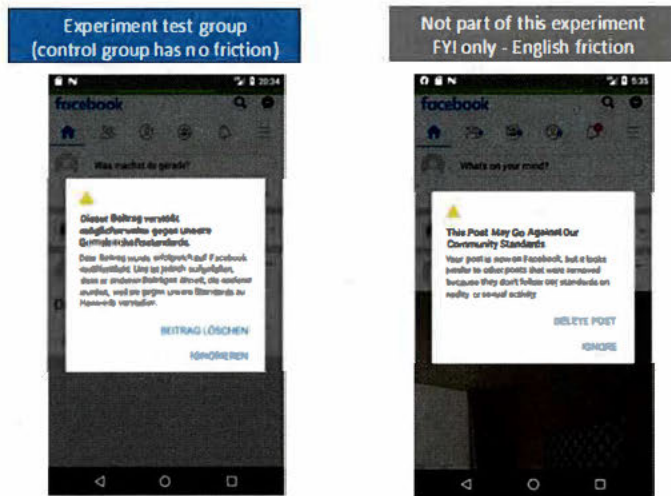
- As part of the effort to reduce Entity non-rec prevalence, Deamplification and Recommendations Integrity partnered together to launch v1 demotions in GYSJ.
- We demote Groups using strikes, severe strikes, misinfo strikes, borderline nudity model and EAU non-rec model.
- Prevalence delta labeling result showed **19%** reduction on GYSJ non-rec prevalence, which translate to **~15%** reduction on overall Entity non-rec prevalence.
- Tradeoff on growth/engagement metrics:
 - 0.2% regression in CEDAU (Community Engage Daily Active Users)
 - 3.9% regression in GYSJ conversions overall

- We are monitoring H1 holdout to confirm the prevalence impact. So far we have observed **13% reduction** for Entity non-rec prevalence. We plan to increase labeling volume to narrow the confidence interval.



[REDACTED], IX Enforcement] Expanding Hate Speech Friction to 13 more locales is estimated to reduce topline violations by **-97K/year (-125K, -68K)** ([REDACTED])

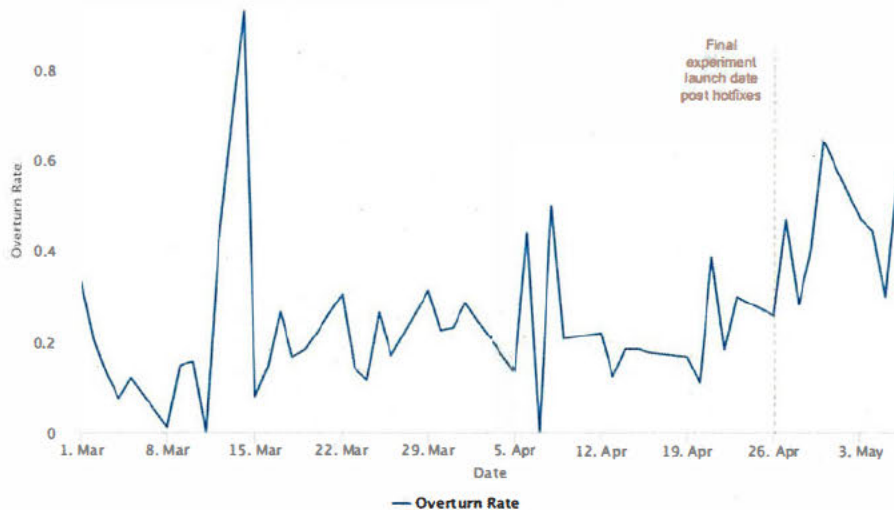
- Hate Speech Friction showed significant prevalence reduction in the past but has been limited to English locales
- We expanded coverage to 13 more locales, including Arabic, Portuguese and French
- Experiment results are in line with expectations:
 - Hate speech violations reduced by -7.1% [-9.1%, -5.0%] when compared to control group.
 - Topline reduction of -0.13%[-0.17%, -0.09%], or -97K [-125K, -68K] less violations per year.
 - There is no significant regression in compact core metrics.
- We will continue to expand coverage to more platforms (such as FBLite) and to operationalize our friction coverage metric to track expansion progress.



[REDACTED], Clustering] Recently launched p(FP) ranker in HIPO drives ~5x False Positive IV and ~2x baseline overturn rate! (+ [REDACTED])

- High Impact False Positive Override (HIPO) currently runs in the US with ~120 reviews/day of content we proactively address for concerns of over-enforcement. To maximize the impact of cluster-based content restoration and increase False Positive IV (FP IV; measured through 7-day restored VPVs) we trained and tested a ranker that outputs predicted False Positives (p(FP)) scores leveraging a combination of cluster signals, WPIE inputs, previous violations from a user, and historical job data.
- Given the lack of job volume for a true A/B test, we chose to launch the p(FP) ranker and monitor FP IV impact through a pre-post. Between the two timeframes we saw a ~109% increase in overturn rate (from 17.1% to 35.7%) alongside a ~387% increase in FP IV (from 805 to 3,923 per review).
- The increase in FP IV came from both an increase in restored VPVs on the object (89 to 317; ~256%) and restored VPVs from fanouts (717 to 3,606; ~403%) with the latter accounting for a larger portion of the lift.
- Table below highlights results on key metrics with the line chart outlining the sustained increase in overturn rate we've seen since launching the ranker. More details [here](#).

Test Group	Jobs Reviewed	Jobs Overturned	Overturn Rate	Total Restored VPVs (per Review)	Restored VPVs on Object (per Review)	Restored Fanout VPVs (per review)
Pre	533	91	17.1%	805	89	717
Post	575	205	35.7%	3,923	317	3,606
Lift vs Pre Timeframe			108.8%*	387.3%*	256.2%*	402.9%*



Progress

[REDACTED], Recommendation Integrity] RI Non-rec V4 filtering protocol reduced non-rec VPVs on IFR by 84%, and we are working on getting launch approval (+ [REDACTED]).

- In V4, we added new signals and updated existing signals which will detect more not-recommendable contents and therefore reduce non-rec VPVs and prevalence on IFR.
- We tested two variants with slightly different signal threshold, and the winning group showed **84%** reduction in non-rec VPV and **2%** reduction in reported VPVs.
- Precision of the V4 protocol for IFR is **78%**, which is substantially higher than guardrail of 50%.
- We already got the green light in IFR review. Next week, we will work on getting the Facebook Integrity Build with Care approval as a required next step.
- In parallel, we are setting up the content holdout and will kick off prevalence delta labeling.



[REDACTED], Clustering] GTM labeling shows that we can safely migrate ~74% of current HIPO volume to cheaper and higher supply reviewers. Findings allow us to scale across markets in H2! [REDACTED]

- To scale the impact of cluster-based content restoration by (a) sourcing more review capacity within the US and (b) expanding to additional markets, we explored the use of Scaled Review (general

content moderation reviewers) in place of the lower supply - but higher accuracy - Subject Matter Expert (SME) reviews. In evaluating trade-offs, we looked at consistency of both sets of reviewers using ground-truth labels (GTM).

- Findings showed top-line **consistency with GTM is ~75% for SME and ~78% for Scaled**. Scaled had a higher overall consistency due to **improved consistency on benign content (~78% vs ~54%)** but **lower consistency on truly violating content (~79% vs ~91%)**. The latter is of higher concern as we **re-introduce these violations** back into the system after having previously taken content down.
- Based on results, the key trade-off that we evaluated was as follows: **out of a 100 HIPO reviews, Scaled introduces 7 net violations** with the benefit of (a) **cheaper review (\$30/hr vs \$120/hr)** and (b) **review that's easier to expand across markets** (currently 10x capacity in the US alone). Cutting results by cases where we **only restore when we get two consecutive ignore decisions** from Scaled (and don't require a third tie-breaking review) showed we can go from **~79% to ~84% consistency on violations** and go from **7 → 4 net violations** while restoring **~74% of current volume**.
- Given findings, we've aligned on the following two-stage strategy for scaling HIPO:
 - [Short-term - Q3] **Move to Scaled but only overturn when reps consecutively agree to restore.**
 - [Longer-term - Q4/H1 '21] As we consolidate HIPO with X-Check (another mistake prevention lever), use **X-Check market reps to review the jobs we're carving out.**
- More details [here](#).

IG Mental Well-being] Hike like controls (Daisy) launched worldwide May 26th!

Shilpa Mody,

- Context:**
 - Prior to the start of new Daisy Controls tests, ~20% of MAU was in a Pure Daisy experience where they do not see like counts, while ~80% had not experienced Daisy.
 - In the Pure Daisy V1 tests in December 2020, we observed an Ad Revenue Impact of $-1.0\% \pm 0.9\%$ and a TS impact of $-0.290\% \pm 0.235\%$.
 - Daisy Controls experiments allow (1) 80% of users who are new to Daisy to opt-in to 'Hide Likes' and (2) 20% of users who are already in Pure Daisy V1 to 'See Likes'.
 - In addition to a consumer control (see or don't see likes on any content in feed) we introduced a producer control (show or don't show likes for anyone for this post).
 - Our expectation was that for users in the New to Daisy Cohort we will see a negative ecosystem impact bounded by the previous Pure Daisy V1 test results while the Pure Daisy Cohort would see a positive ecosystem impact based on the same logic. As a formula: [net number of 'hide like' users] * [effect per user].
- Results Summary:**
 - Ecosystem Impact**
 - New to Daisy Cohort (~80% of MAU, expected negative impact):** Test results show (Column D) that the ecosystem impact is inline with our original expectations of lower negative impact than Pure Daisy V1 tests (Column D).
 - Note that historical data for which users were in Daisy V1 has expired so some V1 users are mixed in to our 'new to daisy' cohort, likely explaining the net positive ecosystem impacts. (I used pocket deltoid to split out countries where daisy v1 was rolled out to 100%.)

- **Pure Daisy Cohort (~20% of MAU, expected positive impact):** Test results show (Column E) that the ecosystem impact is in line with our original expectations of positive ecosystem impact compared to Pure Daisy V1 tests (Column E).
- **Daisy Controls - Blended Impact:** The blended impact (Column C) shows the impact of launching Pure Daisy V1 to all users.
- **Daisy Controls Adoption:**
 - **Consumption Control:** Non-Daisy users hiding likes is 0.72% (pre-test estimates were 1 to 5%), and Historical Daisy users choosing 'Show likes' is 64% (pre-test estimates were ~20%) driven mostly by interstitial QP (61% incremental show like rate) compared with megaphone QP (3% show like rate) [query]. Based on the adoption rates of Daisy consumption control we estimate that after 100% roll out of Daisy Controls 7.5% of MAU will have hidden like counts and 92.5% of MAU will be shown like counts.
 - **Production Control:** Looking at 'days since first exposure' view, 0.7% of users exposed will try the feature at least once within 21 days of exposure (2.2% among feed producers only) [query]. On a per day basis, 0.06% of users exposed for 21 days hid likes on at least one post and 0.15% of all posts had likes hidden. [query]
- **App start Regression:** In prior versions of the controls test we found significant app start regressions due to calculating social context for all users irrespective of whether they have opted into Daisy. Fixing this regression allowed us to ship with positive ecosystem impacts. Team will continue to iterate on perf post-launch.
- **Launch Press:** A key goal of launching Daisy as a control was to ensure people had access to a less pressurized feed without forcing users who rely on like counts for brand integration deals or who don't feel pressure from like counts. Following our [Good Morning America promotional segment](#) we've seen a slew of fair-to-positive articles (with [Tech Crunch's](#) being one representative sample).

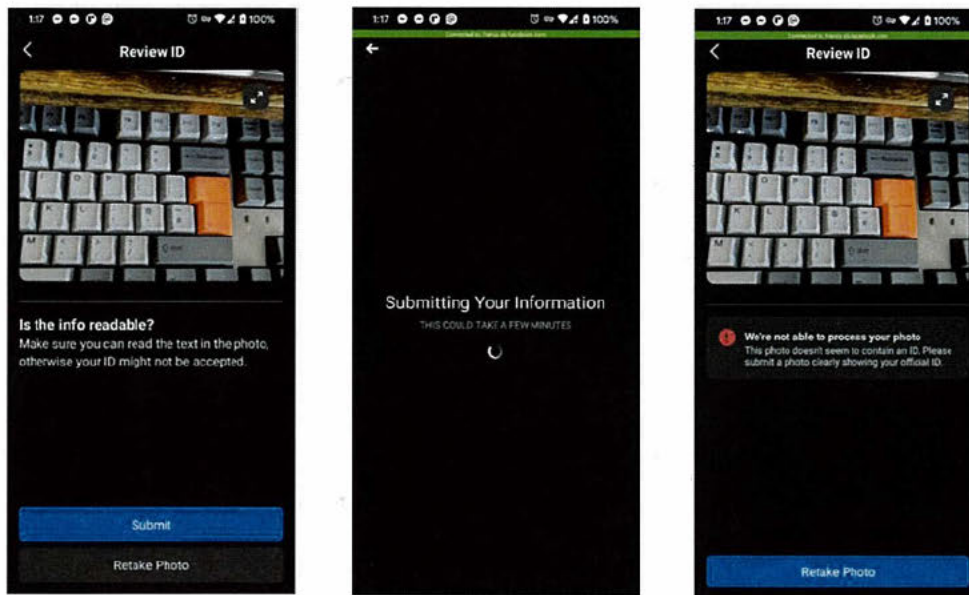
Launch post

Ads	12/11-12/25/19 network test	Daisy controls test 3 (fixed perf) 5/11 - 5/23		
	Pure Daisy V1	Daisy Controls (overall)	Daisy Controls (new to daisy)	Daisy Controls (in Pure Daisy V1)
IG Event Based Revenue	-1.0% ± 0.9%	can't use due to iOS signal loss		
IG Billing Based Revenue		0.107±0.040	0.101±0.086	0.19±0.18
Ad Impressions, Feed	-0.060% ± 0.189%	0.081±0.019	0.098±0.049	0.12±0.13
Ad Clicks, Feed	-0.923% ± 0.295%	0.132±0.082	0.109±0.088	0.29±0.22
Ads CTR, Feed	-0.7% ± 0.1%	0.023±0.065	0.009±0.082	0.17±0.20
Appwide				
IG DAP	-0.092% ± 0.067%	0.004±0.013	0.005±0.017	0.011±0.043
IG App Sessions	-0.281% ± 0.213%	0.000±0.014	0.007±0.040	0.10±0.10
IG App TS	-0.290% ± 0.235%	0.033±0.016	0.036±0.040	0.10±0.10
Consumption & Interaction				
Overall Likes	-1.162% ± 0.320%	0.051±0.078	0.021±0.083	0.27±0.23

Overall Comments	-2.132% ± 0.696%	0.04±0.19	0.06±0.19	-0.07±0.61
Home Feed Post Impressions	-0.022% ± 0.183%	0.051±0.041	0.055±0.044	0.03±0.11
Explore Media Impressions	-0.395% ± 0.233%	0.012±0.036	0.020±0.077	0.07±0.20
Outbound Follows	-0.989% ± 0.444%	-0.28±0.23	-0.33±0.25	0.13±0.66
Sharing				
Direct Sends	0.069% ± 0.649%	0.053±0.070	-0.00±0.13	0.03±0.37
Feed Media Created	2.615% ± 5.239%	-0.02±0.15	0.01±0.16	-0.21±0.37
Feed Producer DAP	-0.016% ± 0.423%	-0.01±0.13	-0.01±0.14	0.02±0.33
Feed Media Deleted (from A/B test)	-2.204% ± 1.813%	0.30±0.39	0.23±0.40	0.9±1.1
Feed Media Archived (from A/B test)	-2.202% ± 2.076%	0.08±0.48	-0.25±0.52	1.0±1.2
Notes				
	All users forced into 'Hide Likes' experience.	Fixing compute recovers negative perf impact.	Some Daisy V1 users mixed in due to data expiration of original tests so this is not 100% clean.	65% opt to see likes = positive ecosystem impact

[REDACTED], Identity Integrity] Real-time feedback during ID submission process leads to increasing ID submissions (ID submissions +16.6% or +538/day) and submissions with better quality ([REDACTED])

- **Background:** we launched an experiment on ID submission process on EAR Android.
 - Currently users receive no feedback during ID submission, and has to wait for FB to get back to them with a decision. As a result, if the ID is not accepted, users will have to go through the submission process again. Majority of ID rejections are due to usability reasons (i.e. too blurry, not actually an ID etc).
 - In this experiment, after the user uploads the ID (but before submission), our detectors will verify if the ID passes or fails usability checks and we will provide real-time feedback to enable the user to re-upload another ID.
- **Results:** real-time feedback leads to **higher submission volume** and submissions with **better quality**
 - Total submissions: + 16.6% or +538 per day.
 - Quality metrics (% of ID submissions with per-field OCR confidence > 0.8): increased by more than 10% for first name, last name, date of birth and ID expiry date fields.
 - Total jobs eligible for manual review: +10.2% or +172/day. This increase is small compared to total IDR jobs for manual review (85k/day). With the increase, we are still below the EAR guardrail for ID submission.
- [Post](#) with more details
- Experimented flow:



[REDACTED], Compromised Response] Better content leads to a 1% increase in Epsilon clear rate and ~1.5K fewer ID reviews a month (...and a bunch of broken attackers' scripts) [REDACTED]

- This test is part of [the Epsilon Funnel improvement project](#) that is aiming at increasing the checkpoint clear rate by addressing some of the known [UX issues](#). In this particular experiment, we are adding content to help people navigate through the checkpoint the challenges ([design mocks](#))
- Overall results are showing minor positive change to the checkpoint clear rate, but we see different funnel rate dynamics for different interfaces:
 - **native apps** are improving: +1% to the clear rate or +24K monthly clears. This increase is driven by the rise in challenge attempts (+0.6%)
 - meanwhile, **the web** is going down: -1% to the clear rate or -10K monthly clears. This drop is due to the jump in bounced sessions (+4%)
- Negative web performance is likely due to the broken adversarial scripts.
 - This is not by design, scripting is failing now because we are adding experiment components and that shifts the invisible order of the page elements including buttons and input fields. This shift messing up some attacker scripts but keeps everything the same from a human perspective
- The improved clear rate leads to the decrease in the number of users with ID review attempts by -4% constituting a decrease in IDR by -1.6K accounts monthly
- Work on improving Epsilon UX continues. But we are now also discussing opportunities to add components order scrambling as a permanent deterrent against attackers or as a method to remove attacker factor from experiment results
- As always full results in [the post](#)

Funnel Comparison (share of users on each step of the funnel)

segment based on interface of the exposure

Segment / KPIs	control	test	impact	Δ%	Signif.
<i>all users</i>					
bounced after captcha share	9.7%	9.8%	6,580	0.9%	×
challenge attempted share	78.6%	78.7%	8,803	0.1%	×
challenge success share	62.5%	62.6%	7,760	0.2%	×
flow cleared share	63.4%	63.6%	15,150	0.3%	×
attempted id review share	0.5%	0.5%	-1,658	-8.0%	×
login compromised share	2.2%	2.2%	5,798	3.4%	×
<i>native apps exposure</i>					
bounced after captcha share	8.2%	8.2%	-1,988	-0.4%	×
challenge attempted share	79.0%	79.3%	18,873	0.4%	✓
challenge success share	60.0%	60.4%	19,668	0.6%	✓
flow cleared share	61.2%	61.7%	24,198	0.7%	✓
attempted id review share	0.6%	0.5%	-813	-2.7%	×
login compromised share	2.1%	2.1%	2,630	2.3%	×
<i>www/m_site exposure</i>					
bounced after captcha share	13.4%	13.8%	9,965	3.2%	✓
challenge attempted share	77.5%	77.0%	-10,890	-0.6%	✓
challenge success share	68.1%	67.5%	-11,665	-0.7%	✓
flow cleared share	68.1%	67.8%	-8,628	-0.6%	×
attempted id review share	0.5%	0.4%	-885	-8.4%	×
login compromised share	2.4%	2.5%	3,310	6.1%	×

[REDACTED], Messenger Well-Being] Analysis into alternative ePD compliant trigger signals for Safety Notices supports reactivating them in EU countries.

- Safety Notices are in-thread warnings for users if they are engaging with an account we suspect could be harmful. We currently show around 1.5M Safety Notices per day.
- Safety Notices have been entirely disabled in EU countries since the introduction of the ePD legislation due to compliance issues.
- Analysis into alternative trigger signals showed that using friending signals and ePD compliant classifiers would give us similar target audiences.
- We are now implementing changes to these triggers for ePD countries and we will run an experiment to measure the impact on interaction rates, block rates, and reporting rates.

People

[REDACTED] Thanks [REDACTED] for the support on HIPO expansion sizing!

[REDACTED] Happy 1st Faceversary, [REDACTED]!!!

[REDACTED] Happy FV1, [REDACTED]!! 🎉🐶

[REDACTED] Congrats [REDACTED] on the birth of your baby!

[REDACTED] Thanks [REDACTED] for all your help this week in setting up the CSOM upsell trigger signals!

[REDACTED] Thank you, [REDACTED] for setting up the script for the comment covers experiment!

[REDACTED] Thanks [REDACTED] for all the support and hard work on backlog reduction!

█████ Thank you █████ for being such a good partner on Bootcamp pitching for CS

Me

█████ Finished watching The Innocent on Netflix. What a great show.

█████ Dogsitting two chihuahuas named Archibald and Mortimer this weekend

█████ Rooting for the Suns to go all the way this year, what a story.

█████ Enjoying the sunny and warm weather in San Jose (CA) this week!

[█████] Zion & Salt Lake — a plane! a wedding! Wild.

█████ The social calendar this weekend is close to a pre-pandemic level of normalcy!

█████ I'm looking forward to warm weather and ice cream this weekend!

[█████] Had a great long weekend!

█████ I'm going to a cricket batting bar this weekend! I feel like this is the last step for me to become completely English.

█████ I went to Edinburgh and Glasgow for the long weekend. It was so amazing and I would definitely go back!

█████ Planned a tiny pub crawl for this weekend. I think it's a positive side effect from the vaccine shot I got last Friday.

█████ Had a lovely long weekend and went cycling in New Forest and saw wild ponies!

█████ Training for the half marathon begins next week

[█████ on PTO]