

Exhibit 22

Drivers: [HYPERLINK "mailto:yandeng@meta.com" \h] [HYPERLINK "mailto:yuzhuoran@meta.com" \h]

Core Ranking - H2 2023 Lookback

TLDR

- **Goals & Milestones**
 - Feed Core Ranking had a successful half exceeding all 3 metric goals*, delivering +0.34% **Gen Pop Cap15 Sessions (136% of goal)**, +0.16% **Teen Cap15 Sessions (107% of goal)** and +0.73% **High ME Post Imps (Incl IV) (146% of goal)**.
 - We completed our Can't Fails to significantly increase our model sophistication, compared to BOH, increasing **training data volume by 11X**, **model size by 200X**, and **expanding our feature pool** with a wealth of dense, sparse, and embedding features, which led to up to 8% model NE improvements.
- **Top Highlight | We launched 3 major model upgrades, bringing us to parity with SoTA models and landing 0.25% gen pop sessions and 0.09% teen sessions cumulatively from backtests.** We made significant model sophistication investments in improved model architecture, which increased model size by almost 200X, from 258MB at the end of 2022 to 50 GB at the end of 2023; increased model training data from 5% to 100% for Teens and 60% for gen pop; and enhanced our feature pool with more nuanced signals to enable better model prediction, such as teen breakdown features, content / user understanding features, cross surface features, and more.
- **Top Learning | Increasing connected feed inventory through sourcing can unlock value for Teens.** Teens' lack of inventory explains most of the difference in teen and adult sessions' impact from Feed ranking (0.4% teen sessions for every 1% adult sessions). Dynamic lookback windows launch based on the user's inventory level drove scroll depth and consumption on feed ([HYPERLINK "https://fburl.com/deltoid3/fn3hds96" \h], [HYPERLINK "https://fb.workplace.com/groups/1840194289545280/permalink/3670886636476027/" \h]), especially for teens (+1.9%) vs. non-teens (+1.4%). We believe this is a promising direction, and plan to improve Connected Feed inventory through sourcing to unlock Teen sessions.

Strategy

Our mission is to help people catch up on the latest and greatest, by accurately predicting user behaviors efficiently. In H2, we set out to significantly enhance our model sophistication in order to deliver model improvements for Teens and Gen pop. Our strategy focused on:

1. **Data:** Upsample training data, especially for Teens to enable better model training.
2. **Signal:** Enrich feature pool to better represent viewer, authors and content.
3. **Model Complexity:** Increase model size and complexity to enable more sophisticated predictions.
4. **Efficiency:** Make predictions faster, cheaper and more scalable.

We ended the half with remarkable progress on all fronts, bringing our model sophistication to parity with FBApp and SoTA models and teed up Efficient Generalization as our next 3-year direction to unlock growth.

Goals

Feed Core Ranking team exceeded all 3 metric goals*, delivering +0.34% Gen Pop Cap15 Sessions (136% of goal), +0.16% Teen Cap15 Sessions (107% of goal) and +0.73% High ME Post Imps (Incl IV) (146% of goal).

In addition to the outsized metric impact, we made rapid advancements in model sophistication, increasing training data volume by **11X**, model size by **200X**, and **expanding our feature pool** with a wealth of dense, sparse, and embedding features, which led to outsized model NE improvements.

Metric	Metric type	Joint Goal*	7d Delta @ 1/8/2024	% of Goal Achieved
Connected Feed Ranking				
Gen Pop Cap 15 Sessions	Goal	0.25%	0.34%	136%
Teen Cap 15 Sessions	Goal	0.15%	0.16%	107%
High ME Post Imps (Incl IV)	Goal	0.50%	0.73% ¹	146%

*Goals were measured at the Connected Feed level, between Feed Core Ranking and Feed Ecosystem Ranking teams.

Highlights

Top Highlight | We launched 3 major model upgrades, bringing us to parity with SoTA models and landing 0.25% gen pop sessions and 0.09% teen sessions cumulatively from backtests. We made significant model sophistication investments in improved model architecture, which increased model size by almost 200X, from 258MB at the end of 2022 to 50 GB at the end of 2023; increased model training data to 100% for Teens and 60% for gen pop; and enhanced our feature pool with more nuanced signals to enable better model prediction, such as teen breakdown features, content / user understanding features, cross surface features, and more.

Other highlights:

- Data | Upsampled Training Data to 100% for Teens and 60% for Gen Pop:** Previously, Connected Feed only used 5% of user traffic to retrain our models on a daily basis, which limits our ability to learn marginal user / teen user behavior. In H2, we increased daily training data volume 11X, which led to 100% sampling rate for Teens and 60% gen pop. The increased training amplified our model improvements, which in combination led to 1.6% to 4.6% NE improvement across prediction tasks, 0.07% teen sessions, and 0.1% gen pop sessions ([[HYPERLINK "https://fb.workplace.com/groups/773237757831636/permalink/884958616659549/" \h](https://fb.workplace.com/groups/773237757831636/permalink/884958616659549/)]).
- Signal | Scaled Feature Pool to Improve User and Content Representation:** We expanded our feature pool with a wealth of dense, sparse, and embedding features, resulting in more accurate predictions. We introduced age-based, time-based, cross surface-, realtime, and content understanding features (increasing from 8 to an impressive 50 features!), as well as pre-trained user embedding features such as DV365 and SUM (Scaled User Modeling). Each of these feature proposals demonstrated promising NE gains, with some achieving up to 4% NE improvement ([[HYPERLINK "https://fb.workplace.com/groups/204375858345877/permalink/650887637028028/" \h](https://fb.workplace.com/groups/204375858345877/permalink/650887637028028/)], [[HYPERLINK "https://fb.workplace.com/groups/204375858345877/permalink/650697390380386/" \h](https://fb.workplace.com/groups/204375858345877/permalink/650697390380386/)], [[HYPERLINK "https://fb.workplace.com/notes/689187433319057/" \h](https://fb.workplace.com/notes/689187433319057/)])!
- Model Architecture | Launched 7 Major Model Architecture Improvements:** We launched 7 major model architecture improvements to drive NE and efficiency. In particular, we implemented a model optimizer (Shampoo) specifically designed for DNNs, improved sparse feature architectures (including semantic pooling, sparse feature scaling), improved dense feature representations (including FANet and MHTA), and optimized task architectures (including ATT x GroupBalance and merging secondary organic predictors to main MTML). Each change resulted in NE improvements of 0.1 - 1% across tasks, with some changes driving 4%.