

CS388: Natural Language Processing

Lecture 10: Evaluation Principles, Dataset Artifacts





Announcements

- ▶ HW 2 due today
- ▶ Quiz 2: next Tuesday!



Recap

- ▶ **Pretraining (BERT):**
 - ▶ Train a big model to fill in masked-out words, then adapt it to other tasks. Led to big gains in **question answering** and **NLI** performance. BART/T5, GPT-3, etc. push this further and extend it to other tasks
- ▶ **Decoding methods:** nucleus sampling > greedy for open-ended tasks
- ▶ **Two tasks we'll focus on today: Question answering (QA)...**
 - ▶ "What was Marie Curie the first female recipient of?"
-> "The Nobel Prize" (find this span in a document)
- ▶ **...and NLI**
 - ▶ "But I thought you'd sworn



Today

- ▶ Evaluation in NLP: benchmarks and generalization
- ▶ Spurious correlations / dataset artifacts
- ▶ Debiasing

Cross-Dataset Evaluation

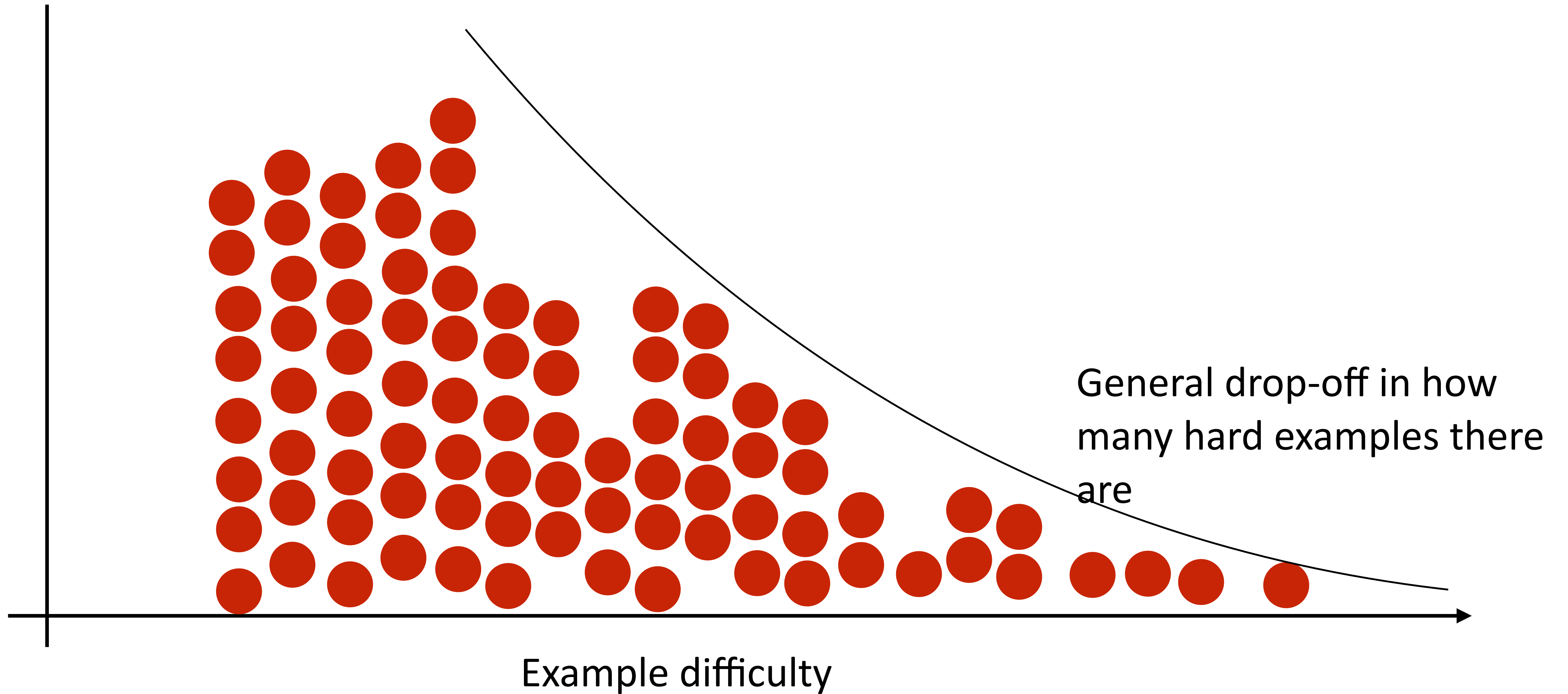


Principles of Evaluation Suites

- ▶ Training and testing on i.i.d. data with big neural models often yields very high performance
- ▶ “Solving” a task (getting human-level performance) may be useful, but often can’t tell us about our models more broadly
- ▶ A parable of single-task evaluation: SWAG

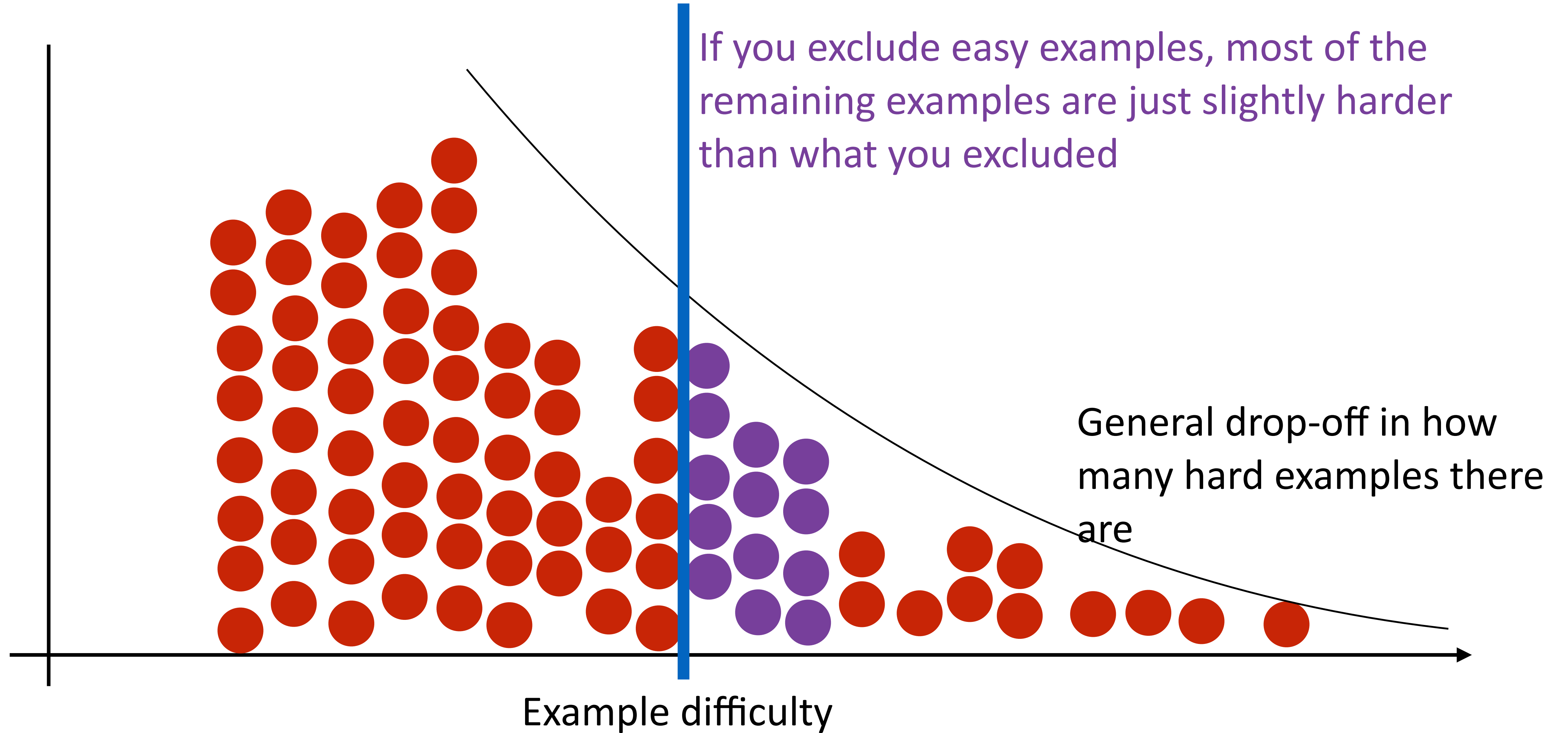


Intuition





Intuition





Principles of Evaluation Suites

- ▶ Training and testing on i.i.d. data with big neural models often yields very high performance
- ▶ “Solving” a task (getting human-level performance) may be useful, but often can’t tell us about our models more broadly
- ▶ **Designing a single, difficult task is really challenging!**
- ▶ To assess big models, we need **evaluation suites (benchmarks)** like GLUE
- ▶ What makes a good evaluation suite of tasks?



Principles of Evaluation Suites

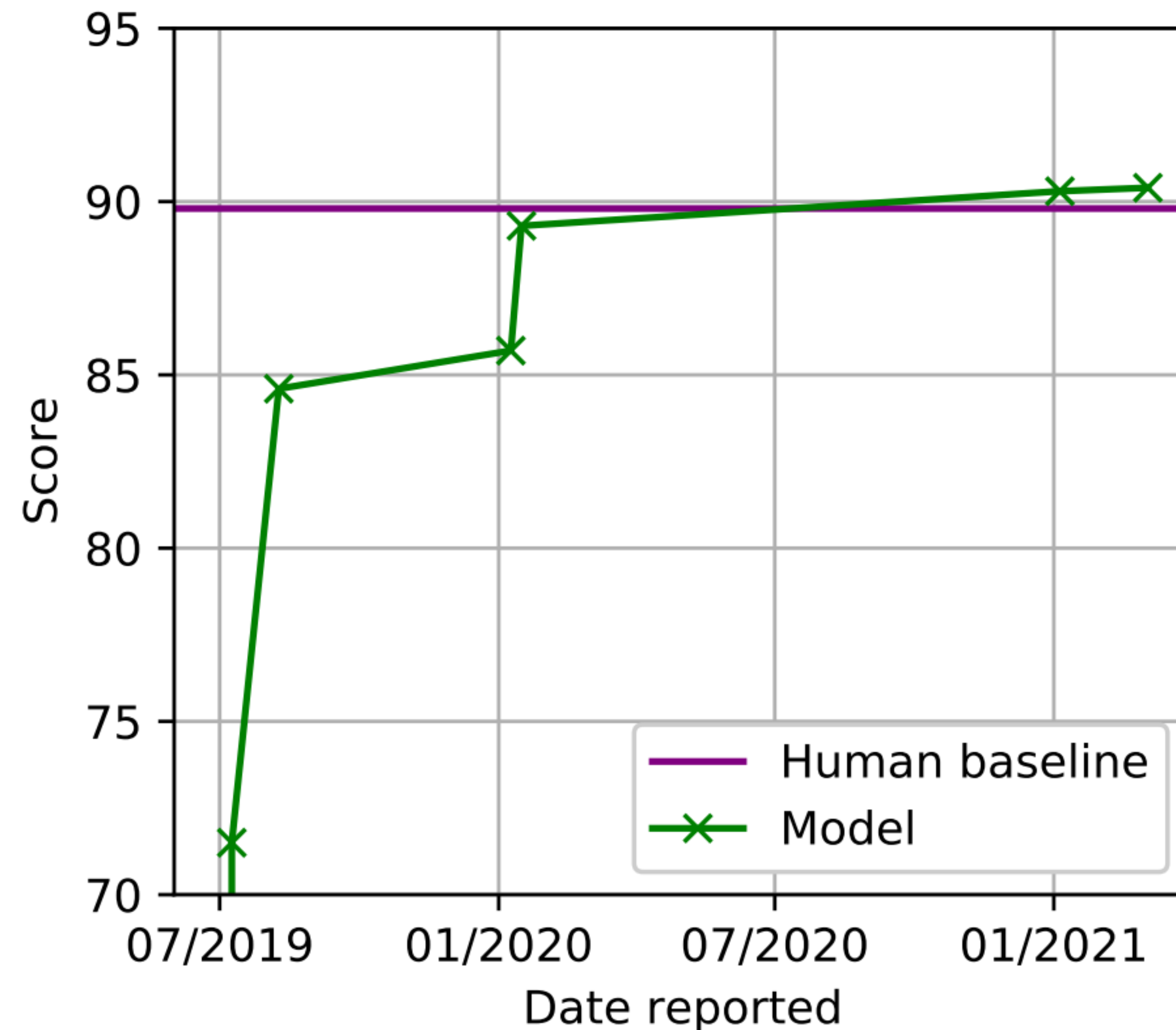
- ▶ Difficulty: even if some task can be solved by hand-engineering, it should be hard to solve all N tasks
- ▶ Diverse: doing well on it should say something useful
- ▶ Good “yardstick”: should understand where human performance is and what good performance on the task would mean
- ▶ GLUE was the first of these, but it wasn’t really diverse and it was too easy. Next step: SuperGLUE



SuperGLUE: Performance

As reported in BIGBench:

SuperGLUE state of the art over time





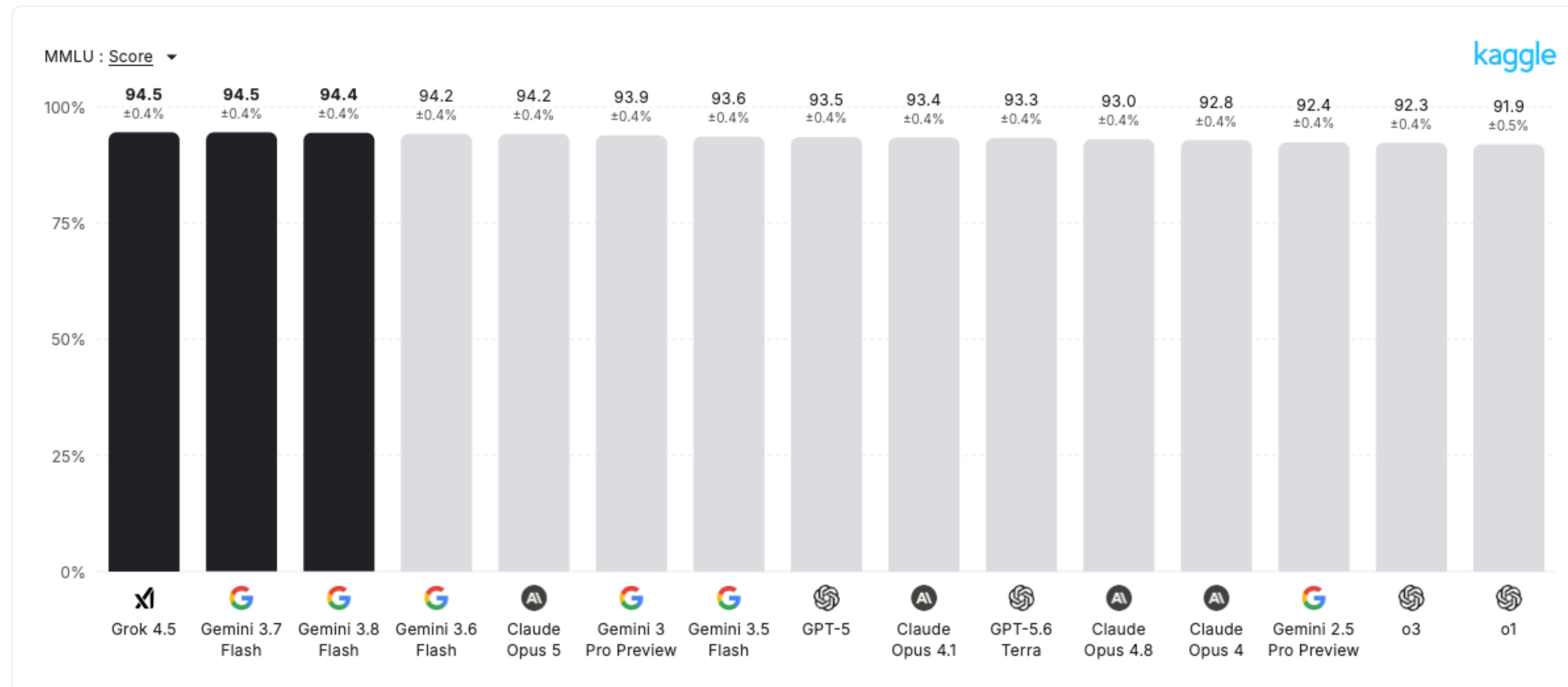
BIG-bench

- ▶ “Beyond the Imitation Game” — aim to learn more than what’s possible from model vs. human performance
- ▶ Particular emphasis on scaling
- ▶ Primarily for pre-trained models without fine-tuning. Therefore, not all tasks have large training (or even test!) sets





MMLU





The End of Benchmarks?

- ▶ Can we still trust benchmarks?
- ▶ Benchmaxxing aka Goodhart's Law
- ▶ Contamination
- ▶ Is this what LLMs are being used for?

Dataset	Split	#Total	#Online	#Total Contamination	#Input-only Contamination	#Input-and-label Contamination
ARC_c	Test	1172	372	336 (28.7%)	53 (4.5%)	283 (24.1%)
CommonsenseQA	Dev	1221	44	20 (1.6%)	3	

Evaluation Under Distribution Shift



Model Performance

- ▶ If models can be fine-tuned on each of n tasks in an evaluation suite and perform very well on the held-out test dataset, have we solved everything we want?
- ▶ What can go wrong?

Annotation Artifacts, Reasoning Shortcuts: QA

Annotation Artifacts, Reasoning Shortcuts: NLI

