

Contents

List of Figures	xiii
Preface	xix
Overview	xx
Who is this book for?	xxiv
Acknowledgments	xxv
Prologue	xxvii
1 Introduction	1
<i>From its roots, machine learning embraces the anything goes principle of scientific discovery. Machine learning benchmarks become the iron rule to tame the anything goes. But after decades of service, a crisis grips the benchmarking enterprise.</i>	
The iron rule	3
The ImageNet era	4
The LLM era	6
2 Populations and predictions	13
<i>The mathematical foundations of machine learning follow the astronomical conception of society: Populations are probability distributions. Optimal predictors minimize loss functions on a probability distribution.</i>	
2.1 Prediction	15

VIII CONTENTS

2.2	Risk minimization	22
2.3	Errors and metrics	24
2.4	Model training	26
	Notes	29
3	Detecting differences	31
	<i>A single statistical problem illuminates much of the mathematical tools necessary for benchmarking. The key lesson is that sample requirements grow quadratically in the inverse of the difference we try to detect.</i>	
3.1	Model comparisons from samples	32
3.2	Coin tossing	34
3.3	Distances between distributions	38
3.4	Concentration inequalities	41
3.5	From coin tosses back to benchmarking	44
	Notes	48
4	Holdout method	51
	<i>The holdout method separates training and testing data, anything goes on the training data, iron rule on the testing data. Not all uses of the holdout method are alike.</i>	
4.1	Testing on the training set	51
4.2	Generalization	52
4.3	The holdout method	54
4.4	What's the holdout method for?	59
4.5	Variants of the holdout method	62
4.6	Error bars and confidence intervals	65
	Notes	70
5	Test set reuse	75

Statistics prescribes the iron vault for test data. But the empirical reality of machine learning benchmarks couldn't be further from the prescription. Repeated adaptive testing brings theoretical risks and practical power.

5.1	Test set reuse in machine learning benchmarks	75
5.2	The holdout method under adaptivity	79
5.3	Alternatives to the holdout method	84
5.4	Freedman's paradox	85
	Notes	89

6 Scientific crisis 91

A replication crisis has long gripped the empirical sciences. Statistical practice is vulnerable for fundamental reasons. Under competition, researcher degrees of freedom outwit statistical measurement.

6.1	The replication crisis in the statistical sciences	91
6.2	Propensity of false positives	95
6.3	Perspectives on the crisis	98
6.4	Goodhart's law	102
	Notes	103

7 Replication in machine learning 109

The preconditions for crisis exist in machine learning, too. And yet, the situation in machine learning is different. While accuracy numbers don't replicate, model rankings replicate to a significant degree.

7.1	The preconditions for crisis	109
7.2	Replication in machine learning	111
7.3	The trouble with absolute benchmark numbers	113
7.4	Model rankings in the ImageNet era	117
7.5	Measurement versus ranking	124
	Notes	125

8 Forces against crisis 129

If machine learning thwarted scientific crisis, the question is why. Some powerful explanations emerge. Key are the social norms and practices of the community rather than statistical methodology.

8.1	Beating the previous best	130
8.2	Biases and heuristics	136
8.3	Rip Van Winkle’s replication problem	139
8.4	The touch of the Blenheim Spaniel	140
8.5	Code and collaboration	143
8.6	Kaggle versus science	144
8.7	From benchmarking to scientific progress	146
	Notes	148

9 Labeling and annotation 149

If the holdout method is the greatest unsung hero, data annotation is not far behind. But conventional wisdom clouds the subtle role that annotation plays for benchmarking.

9.1	Annotator errors and annotator agreement	150
9.2	Labeling as prediction	155
9.3	Effects of label errors on model comparisons	158
9.4	Quantity versus quality	160
9.5	Resilience of rankings to label errors	165
	Notes	166

10 Generative models 169

The ImageNet era ends as attention shifts to powerful generative models trained on the internet. The new era also marks a turning point for machine learning benchmarks.

10.1	Language models	171
10.2	Scaling	178
10.3	NLP’s benchmark culture before transformers	186

10.4 CLIP and a final look at ImageNet	192
Notes	195
11 Evaluating language models	197
<i>After training, alignment fits models to human preferences. Part of the post-training pipeline, alignment transforms evaluation results. How post-training makes such a difference brings new challenges for benchmarking.</i>	
11.1 Post-training methods	199
11.2 Generative evaluation	206
11.3 Confounded evaluations	214
11.4 Model comparisons and rankings	221
Notes	226
12 The problem of aggregation	231
<i>Multi-task benchmarks promise a holistic evaluation of complex models. An analogy with voting systems reveals limitations in multi-task benchmarks. Greater diversity comes at the cost of greater sensitivity to artifacts.</i>	
12.1 Multi-task benchmarks	232
12.2 Problems of aggregation and voting systems	237
12.3 Ranked voting	239
12.4 Rated voting	243
12.5 Empirical trade-offs in multi-task benchmarks	247
12.6 Latent factors in benchmark performance	250
Notes	251
13 When the model moves the data	255
<i>Models deployed at scale always influence future data, a phenomenon called performativity. Performativity breaks evaluation and creates the problem of data feedback loops. Dynamic benchmarks try to make a virtue out of it.</i>	
13.1 Morgenstern's prophecy about prediction	257

XII CONTENTS

13.2 Performative prediction	259
13.3 Repeated risk minimization	265
13.4 Data feedback loops	270
13.5 Dynamic benchmarks	274
Notes	277
14 Evaluation at the frontier	279
<i>As models gain in capabilities, human supervision increasingly becomes a bottleneck. The hope is that models will supervise and evaluate each other, but there are limits to automatic evaluation.</i>	
14.1 LLM as a judge	281
14.2 Debiasing evaluations	287
14.3 Restricted model evaluation strategies	292
14.4 Evaluation in the real world	298
Notes	302
15 Epilogue	305
References	311
Index	341

— 1 —

Introduction

From its roots, machine learning embraces the anything goes principle of scientific discovery. Machine learning benchmarks become the iron rule to tame the anything goes. But after decades of service, a crisis grips the benchmarking enterprise.

“The only principle that does not inhibit progress is: anything goes,” the philosopher Paul Feyerabend proclaimed half a century ago. Taking on the giants of scientific method, like Popper and Kuhn, the mantra became a point of debate in philosophy circles and at dinner parties. Feyerabend advocated for flexibility and creativity in knowledge production to a degree he deemed *anarchic*. No single way of doing science is superior, he insisted.

The early days of machine learning, then called pattern recognition, certainly resonated with Feyerabend’s words. Pattern recognition was a child of the cybernetics era, when science and science fiction came close. Emerging concepts from control, communication, and computation blended in ambitious research programs aimed at building intelligent machines.

Frank Rosenblatt, the inventor of perceptrons—a precursor to modern artificial neural networks—was emblematic of the time. His work on perceptrons was unconstrained by scientific orthodoxy; instead, he embraced exploration, intuitive leaps, and interdisciplinary thinking. A psychologist by training, Rosenblatt took inspiration from, among many others, the ideas of the economist Friedrich Hayek about how intelligence emerges in distributed computation. Hayek, too, opposed what he called *scientism*, the

excessive application of scientific methods and principles to domains where they may be inappropriate. It's hard not to see Rosenblatt's 600-page tome, *Principles of Neurodynamics*, as a testament to the unconstrained scientific enterprise Feyerabend demanded.

In some significant ways, today's machine learning research has stayed true to its roots. To this day, the field doesn't prescribe how to get to a result. It does not insist on rigor, method, or science in the way that researchers derive what they propose. If tomorrow a breakthrough result appeared that was somehow inspired by child development, quantum chemistry, or cell biology, no one would raise an eyebrow. Researchers are completely unconstrained in what techniques they may apply. Any model architecture, any optimizer, and any tweak is fair game.

As a result, it's been difficult to build a science of the things that people actually do in practice. The model builders always seem to outpace all attempts to theorize their doings. In a search space with infinitely many degrees of freedom, no single thing is necessarily fixed. Any particular design choice can be compensated elsewhere. Hundreds of model architectures came only to give way to a different architecture within months. What seems necessary at any point in time—different kinds of regularization or weight normalization—looks dated eventually. A flurry of papers promising to reveal “all you need” only culminated in the consensus that “all you need is all you need”. This truism characterizes the craft. In some sense, machine learning is never more than whatever set of tricks you currently need.

There's an exhilarating freedom to “anything goes” that has surely attracted many to the field of artificial intelligence. But the early days already experienced the tragedy of anything goes. The initial excitement about a machine that could separate squares from triangles created the field of pattern recognition. The explosively growing field—brimming with activity—rapidly spawned myriad ways of attacking different pattern recognition problems. Research teams claimed all sorts of advances, illustrated by specific experiments and eclectic demonstrations. It became hard to tell what worked best in an ocean of possibilities. For a practically minded field, not knowing a clear answer to what works best is a torturous predicament.

One ingredient was missing to tame the “anything goes”.

The iron rule

To tame the “anything goes” there had to be some kind of test. The test had better be empirical and quantitative. Aesthetics, theory, and subjective opinion ought not to play a role in the test. Since the goal of pattern recognition was to do things like classify objects in a scene, it made sense to score an algorithm by how often it succeeded in doing so. Classification accuracy was a natural target. Researchers quickly realized, however, that any model did a lot better on data points it had encountered during training than those it hadn’t. So, they agreed to separate training and test cases. Soon, model builders began to compete over who could achieve the highest test accuracy. This common sense agreement among researchers was the starting point of machine learning benchmarks.

In developing benchmarks, pattern recognition discovered its own instance of the iron rule of modern science. A term coined by the philosopher Michael Strevens, the iron rule asserts that all disputes between scientists must ultimately be settled by competitive empirical testing. In this view, modern scientific communities organize around empirical protocols that lay out the rules of scientific competition. These are a lot like the rules in a sporting competition. Scientists are free to think and do whatever they want, but for the purposes of scientific competition, they stick to the rules.

The iron rule makes a virtue out of what might seem like a problem: relentless competition among scientists. By making empirical testing the objective, scientists accumulate knowledge as they compete. Scientific institutions—funding agencies, journals, and universities alike—reinforce the rule by rewarding those who come out ahead in the metrics. Deciding who gets what via empirical testing lowers friction in the gears of science, as it avoids drawn-out debate and keeps personal opinions in check. What results, Strevens argues, is an efficient knowledge machine that powers modern science.

Benchmarking is the iron rule of machine learning research and a radically simple contract at that: Anything goes on the training set, competitive ranking on the test set. The recipe is simple. What’s surprisingly hard is to explain why and when it should work as an engine of progress. Strevens calls the idea that science should run on the iron rule unreasonable on the face of it. Likewise, academics have made a good case against benchmarking.

Critics argue that competitive testing promotes narrow research objectives that stifle more creative scientific activities. They say benchmarks incentivize gaming and overfitting, ultimately misleading scientists about model capabilities. Besides, benchmarks give a structural advantage to industry labs that can build the biggest models.

The critics aren't wrong. The scientific case *for* benchmarks is weak. And researchers might've indeed dismissed the idea of benchmarking as unreasonable, if it hadn't been so successful.

The ImageNet era

Benchmarks emerged from little more than common sense and intuition. They appeared in the late 1950s, had some life during the 1960s, hibernated throughout the 1970s, and sprang to popularity in the late 1980s when machine learning began to look like pattern recognition. Today, benchmarks are so ubiquitous, we take them for granted. And we expect them to do their job. After all, they have in the past.

The deep learning revolution of the 2010s was a triumph for the benchmarking enterprise. The ImageNet benchmark was at the center of all the cutting-edge advances in image classification with deep convolutional neural networks. Despite massive competitive pressure, it reliably supported model improvements for nearly a decade. Throughout its long life, a sprawling software ecosystem grew around ImageNet, making it ever simpler to develop and test models on the benchmark.

Even its tiny cousin, CIFAR-10, did surprisingly well for itself. Model builders often put CIFAR-10 into the development loop for architecture search. Once they found a promising candidate architecture, they would then scale it up to ImageNet. Folklore has it that some of the best ImageNet architectures were first developed on CIFAR-10. Even though CIFAR-10 features only 10,000 tiny pixelated test images from ten classes, such as frogs, trucks, and ships, the dataset was any model builder's Swiss Army knife for many years. The platform PapersWithCode counts more than 15,000 papers published with CIFAR-10 evaluations. This does not count the numerous evaluations that went into every single one of these papers. It also doesn't count the enormous amount of engineering work that used the dataset in one way or another.

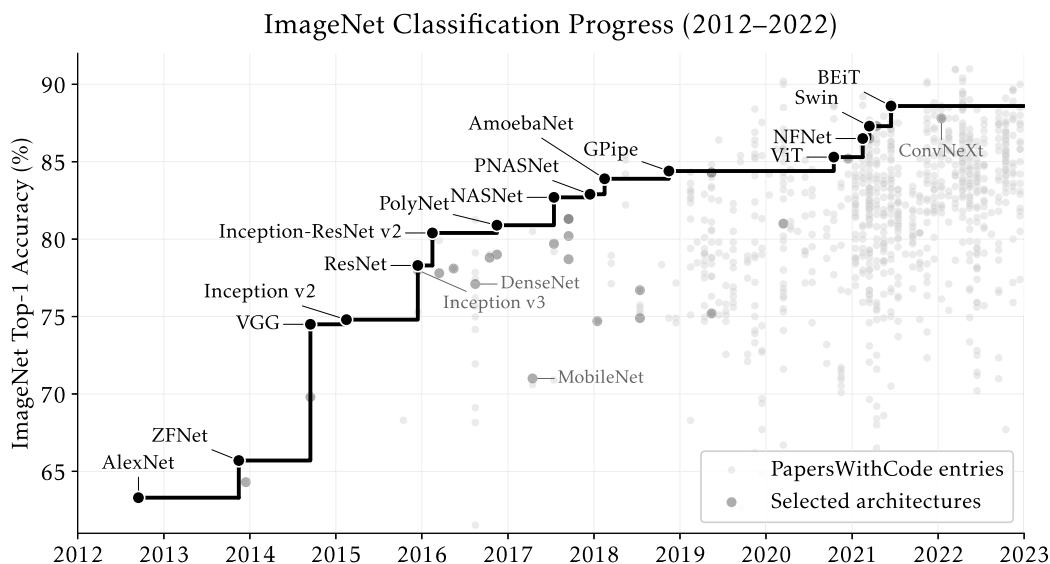


Figure 1.1: A decade of progress on the ImageNet ILSVRC 2012 benchmark.

It might seem reckless that so much work should turn on so little data. With the benefit of hindsight, though, we can even verify that the ranking of popular models on CIFAR-10 largely agrees with the ranking of the same models on ImageNet. Model rankings on ImageNet in turn transfer well to many other datasets. Researchers even created a toy dataset called ImageNot, full of noisy web crawl data, designed to stray as far as possible from ImageNet, while matching only its scale and diversity. Retraining all key ImageNet era architectures on ImageNot from scratch, model rankings turn out to be the same.

The stability of model rankings is a robust empirical fact of the ImageNet era. Numerous papers show that model accuracies change erratically from one benchmark to the next. At the same time, relative comparisons between models turn out to be surprisingly reliable. If one model beats another in one benchmark, it's likely to beat the same model in a different context, too. There's evidently a certain kind of robustness to beating the previous best—the key mechanism of the iron rule. This wasn't clear to begin with. It's just something researchers got used to.

When rankings work, benchmarks free you from second-guessing yourself: Suppose you take the highest-ranked model off the shelf and you spend a few months adapting it to your application. You find out that it doesn't

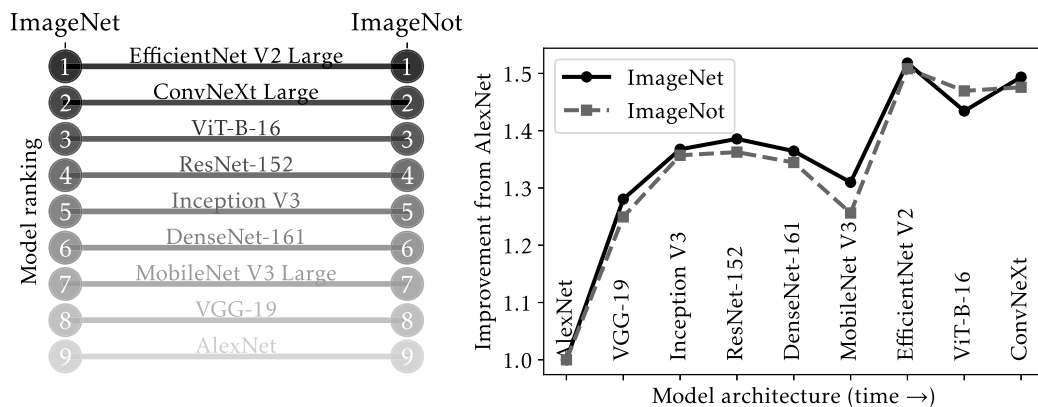


Figure 1.2: Stability of ImageNet model rankings (left) and relative accuracy improvements over AlexNet (right).

work. A good benchmark should reassure you that probably no other model would’ve worked either. This saves you a lot of time, since you don’t have to try out the other few hundred models on the shelf. This is the powerful promise of a good ranking. The benchmark may not anticipate how well a model will work for your application. But it can tell you what’s the best model to start from, thus reducing trial and error.

Faith in the top of the leaderboard was the religion of the deep learning revolution. As is often the case, though, faith is strongest just before the crisis.

The LLM era

Eventually, attention shifted from image classification to natural language processing (NLP), as the new transformer architecture scored a victory over the sluggish recurrent neural networks that had long been the workhorse for sequence problems of all kinds. Transformers were much easier to scale up and quickly took over. A simple training objective—repeatedly predict the next token in a sequence of text—avoided costly human data annotation. Companies quickly scraped up any sequence they could find on the internet, from chat messages and Reddit rants to humanity’s finest writing. New scaling laws suggested empirical relationships between the dataset size, the number of model parameters, and pretraining compute that jointly minimize training loss. For a while, it seemed the only thing left to do was to sit back and watch new capabilities emerge in explosively growing models.

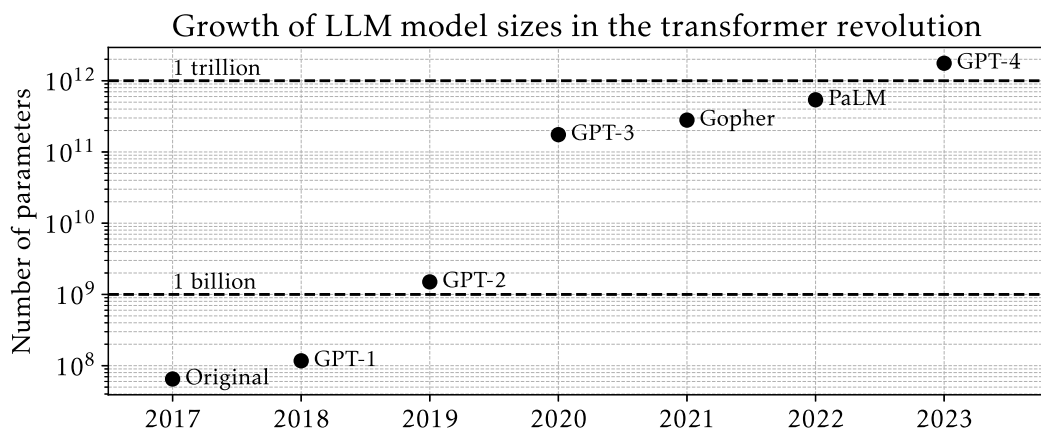


Figure 1.3: Model size growth in the transformer revolution. GPT-4 model size according to rumors.

Unlike in image classification, in NLP there has never been a single central benchmark task. Translation, parsing, entailment, and sentiment are only a few of the many tasks we associate with language understanding. Each has all sorts of different benchmarks. Most of these test tasks are a leap removed from how the model trains by predicting the next token in a never-ending stream of text. This divergence between training objective and test task spelled trouble. Right after training, large language models tend to ramble on, displaying disregard for utility, truthfulness, or manners. Hallucination is the default rather than the exception. It's part of what makes these models so intriguing: They can make up stuff. Yet, they are still reluctant to follow instructions and solve any concrete task reliably. Something was missing.

The paradigm of alignment became the catch-all cure for whatever problems models have right after training. Alignment fine-tunes pretrained models in one way or another on data such as human demonstrations, comparisons, or thumbs-up clicks from chatbot platforms. Computationally, alignment costs only a tiny fraction of the resources allocated to training. Yet, it has a transformative impact on benchmark performance and subjective human evaluations. Alignment is now part of a broader post-training pipeline that aims to make a model useful for concrete use cases. Post-training is what you do with benchmark performance in mind.

If language has no single most important task, what benchmark should we pay attention to? An exploding task plurality directly threatens the iron rule

of modern science. What exactly should scientists compete over when there are hundreds of options on the menu? If decathlon isn't exactly a crowd favorite at the Olympics, imagine watching "centathlon"—an eclectic and somewhat arbitrary collection of a hundred events with no clear winner in the end.

The growing unease with language benchmarks fueled the desire for one authoritative ranking compiling all available evaluation data into one. Researchers hoped that piling up more and more tasks would uncover a *true* ranking. This is the premise of multi-task benchmarks. But aggregation of diverse rankings is a notoriously tricky problem—the bane of all voting systems—and there are no perfect solutions. Evidence soon emerged that aggregation breaks the stability of model rankings we've come to expect from our ImageNet upbringing.

And that was only the beginning of the trouble.

As multi-task benchmarks provided no clear solution to a growing evaluation crisis, much of the competition began to gravitate toward only a few benchmarks. Among them was the MMLU benchmark, which stands for Massive Multitask Language Understanding. MMLU consists of thousands of college-level multiple-choice questions from numerous subjects. In a departure from traditional benchmarks, there is no conventional training dataset; after all, large language models are trained by predicting next tokens on some chunk of the internet.

A general knowledge question in MMLU might ask:

As of 2016, about what percentage of adults aged 18 years or older were overweight? A: 10%, B: 20%, C: 40%, D: 80%.

A question from high school computer science could look like this:

Let $x = 1$. What is $x \ll 3$ in Python 3? A: 1, B: 3, C: 8, D: 16.

Scoring high on MMLU hinges on knowledge the model acquires during training, as well as its ability to understand the prompt. The latter turns out to be surprisingly tricky. Progress on the benchmark catapulted from barely better than random guessing to over 90% in just a few years. What mesmerized observers is that accuracy gains only picked up as models reached a certain size. This sudden increase in accuracy became known as *emergent abilities*, fueling narratives about sudden unpredictable AI advances.

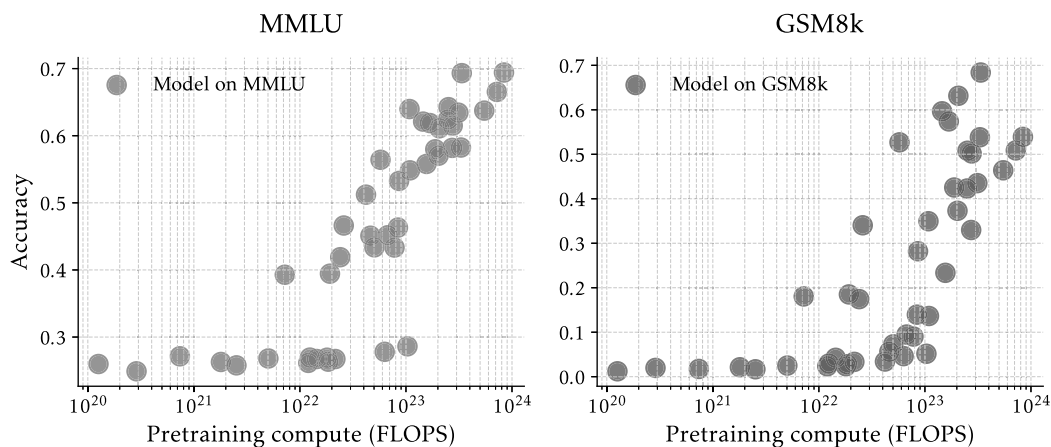


Figure 1.4: Compute versus accuracy on MMLU and GSM8k (Grade School Math 8k), a popular math benchmark.

The fact that models have reached 90% accuracy on MMLU may be an indication that progress on the benchmark has saturated and we should stop using it. However, once established, benchmarks have a long life. Researchers continued to use CIFAR-10 actively well after the 90% mark had been surpassed. Model comparisons don’t necessarily need wide margins. Even small signals can indicate that one model is better than another.

It’s not like people think that CIFAR-10 and MMLU are particularly faithful representations of the real-world challenges you’ll encounter. Rather, these benchmarks create a point of reference for the community. They lay out the rules of the game. The right question is not how to make benchmarks more realistic. The better question is how to set up a game whose winner we should care about. CIFAR-10 turned out to be a good game to play. But there was a deeper reason why MMLU could never become CIFAR-10.

Unlike CIFAR-10, MMLU has no training data. It’s just a test set like most other LLM benchmarks. Providing training data for such benchmarks might seem pointless anyway, given that models train on the internet. In the ImageNet era, researchers primarily compared models trained on the same training data. The LLM era quietly departed from this implicit agreement by making data part of the competition. Each competitor can now choose training data that maximally improve benchmark performance—a problem the AI community dubbed *benchmaxxing*. Even without explicitly training on the test set, model providers can still train on the *test task* by choosing data that best anticipates the test set.

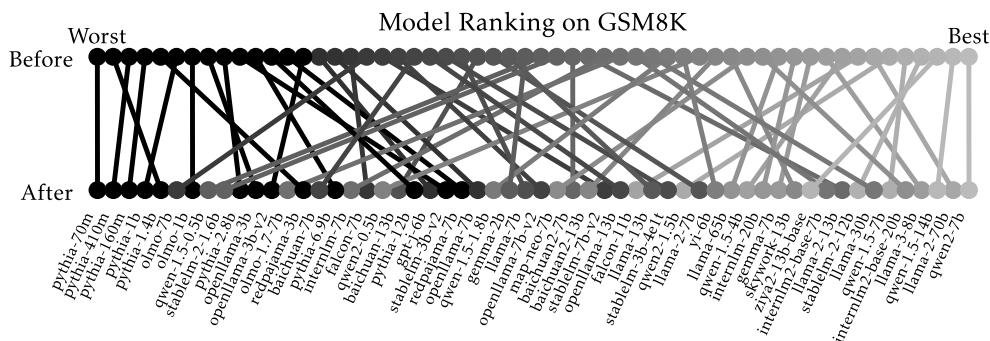


Figure 1.5: Model rankings on the GSM8k benchmark before and after giving each model the same preparation for the benchmark.

The lack of a shared training set has profound consequences for model comparisons and benchmarking more broadly. Every competitor now has their own data recipes, often as closely guarded a secret as the Coca-Cola recipe. When we compare two models on a test set, we therefore can’t know how much each model specifically prepared for this benchmark. The top model on the leaderboard may have simply studied more for this specific test than others. Looking at it one way, this is fair game. Why shouldn’t knowledge about the downstream test inform your training practices? From a scientific perspective, however, it breaks the no second-guessing guarantee of a good benchmark. A lower-ranking model may well have been your best choice, but it showed up less prepared for the test.

The validity of model rankings was an empirical phenomenon of the ImageNet era, but it’s not clear that we should get this lucky again. Indeed, Figure 1.5 illustrates on the GSM8k grade school math benchmark what is true more broadly across different benchmarks: Giving each model the same task-specific preparation results in rankings that differ strongly from those we get when we directly evaluate the model as it is. Which model ranking we should trust is a sticky question in the LLM era.

What further troubles evaluation in the chatbot era is that models deployed at scale influence future data, a phenomenon called *performativity*. It’s a century-old problem that has gained new urgency. Performativity challenges model evaluation, in particular, since there is no longer model-independent ground truth. Data—from text to labels—are now contingent on the model. Many worry that ChatGPT will drown the internet in its own text generations, leading to a vicious cycle of data and model degradation.

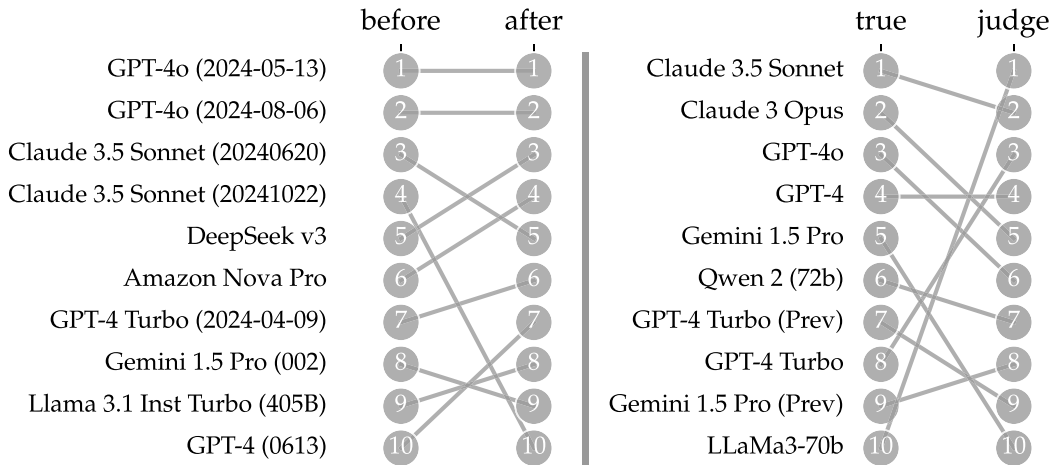


Figure 1.6: Instability of model rankings. Left: Effect of adding weak models to the HELM (Holistic Evaluation of Language Models) multi-task benchmark. Right: Effect of using an LLM (Llama 3) as a judge.

The more we use ChatGPT, the more text will look like ChatGPT. Research on performativity sheds light on this problem of data feedback loops that many see as a fundamental risk to the data ecosystem.

Where some see a problem, others see opportunity. Dynamic benchmarks try to make a virtue out of data feedback loops by creating tests that evolve as models improve. Adding examples where models fail, the benchmark steadily becomes more difficult. Whether model-data feedback loops will be vicious or virtuous is uncertain, but further progress depends on the answer.

The final problem benchmarking now faces is an existential one. As model capabilities exceed those of human evaluators, researchers are running out of ways to test new models. There's hope that models might be able to evaluate each other: Take the best available model and use it to judge new candidates. But this idea of using models as judges runs into serious hurdles. LLM judges are biased, unsurprisingly, in their own favor. Attempts to de-bias evaluations ultimately point back at what's missing in the first place: a reliable benchmark.

Will the iron rule—the old engine of progress in machine learning—grind to a halt?

In a moment of crisis, we tend to accelerate. We do more of the same, hoping that the problem will go away on its own. What if instead we step back

and ask why we expected benchmarks to work in the first place—and what for? This question leads us into uncharted territory. For the longest time, we took benchmarks for granted and didn't bother to work out the method behind them. We got away with it mostly by sheer luck, but we might not be as lucky this time again. Over the last decade, however, a growing body of work has started to map out the foundations of a science of machine learning benchmarks. What emerges is a rich set of observations—both theoretical and empirical—raising intriguing open problems that deserve the community's attention.

If benchmarks are to serve us well in the future, we must put them on solid scientific ground.

— 2 —

Populations and predictions

The mathematical foundations of machine learning follow the astronomical conception of society: Populations are probability distributions. Optimal predictors minimize loss functions on a probability distribution.

The philosopher Ian Hacking called it the *astronomical conception of society*. It's the idea that we can think of populations as probability distributions and find regularities in them in the same way that an astronomer finds patterns by observing the universe.

The idea goes back to the 19th century Belgian astronomer Adolphe Quetelet, one of the founders of quantitative social science, who called it *social physics*. Quetelet believed that there were laws to be discovered in society, such as stable crime rates in France, and that we could find these laws by fitting statistical models to data. Modern statistics adopted the perspective and the discipline built the formalisms around it that students now learn in virtually every class that touches on data. By the end of the 1930s, statistics had settled on its foundations. Scientists are so used to the idea of populations as probability distributions that they no longer question it. Machine learning, in particular, follows the same astronomical conception, especially when it comes to the bulk of work on model building and evaluation, from linear regression to ChatGPT.

What was radical about the idea when it came up is that human populations clearly are not probability distributions. In a sense, they couldn't be further. The world we live in is a complex, interconnected dynamical system of people, institutions, infrastructure, and nature. What we observe changes

over time and in response to our actions. What we see even changes in response to the models that we deploy in the world. It's important to keep this in mind to understand why sometimes things fail. Datasets collected at different times or different locations generally don't follow the same distributions. Machine learners call this *distribution shift*, but the expression starts from a wrong premise. There is no distribution to begin with. The dynamic nature of society contradicts the stationarity of the astronomical conception. When things go wrong in machine learning, it's often fundamentally this contradiction in disguise.

Yet, the astronomical conception is vindicated by its relentless utility. From its early applications in insurance pricing to recommender engines on digital platforms and AI chatbots, fitting statistical models to datasets has always made someone rich. It's hard to argue with success. It's also hard to argue with its simplicity. The formal setup needed for much of this book fits in a few pages and is well worth knowing in your sleep.

What makes the astronomical conception so useful is that it gives us a simple way to make predictions. A prediction is an educated guess about something we don't know for sure given information we do have. Bernoulli called it the *art of conjecturing*. In this sense, prediction is not just about the future. It also applies to things that have already happened, but we're uncertain about them. A camera took an image, but we're unsure from the pixels whether we're looking at a mouse or a hamster. A company filed their quarterly report with the Securities and Exchange Commission, but we're unsure from the numbers if financial fraud has occurred. In both cases, there's a definitive truth value about which we're unsure.

Prediction also applies in cases where there is no definitive answer. To give an example, suppose you know somebody's age, but nothing else. Make a guess about whether the person has a valid driver's license. Statistics lets you do that, once you fix a population. For any given age, say 18, you look up if the majority of 18-year-olds in the population have a driver's license. If so, you guess yes, otherwise no. This common sense rule turns out to be the optimal predictor if the goal is to maximize the probability of a correct guess. Here the probability refers to a random draw of an individual from that fixed population.

The example makes it clear that your best guess is specific to a population. It depends on whether you consider the population of all citizens of

the United States of America or the population of Tokyo residents. It also depends on time. Your best guess in 2025 may not be your best guess in 1998.

2.1 Prediction

To formalize prediction, we start from a distribution $\mathcal{D}$ called *data-generating distribution* or *population*. A draw from the data-generating distribution gives us one data point. A labeled data point is a pair (x, y) , where $x \in \mathcal{X}$ is an array of feature values that could describe a row in a table, a text sequence, or an image. The value $y \in \mathcal{Y}$ is the label assigned to x . The set $\mathcal{Y}$ contains the possible labels that can occur. The set $\mathcal{X} \times \mathcal{Y}$ of all possible data points is called the *universe*. The support of the data-generating distribution is some subset of the universe. A *predictor* is a function $f: \mathcal{X} \rightarrow \mathcal{Y}$ that maps an input x to a *prediction* $f(x)$.

It's often convenient to write the data-generating distribution as a pair (X, Y) of jointly distributed random variables denoting a random draw from $\mathcal{D}$. This lets us write things succinctly, like $\mathbb{P}\{f(X) = Y\}$, the probability of a correct guess. The random variable Y is then called *target variable* as it is the target of our predictor. Throughout, I omit mentioning the population when it is clear from context.

There are two common settings of prediction:

- In a classification problem, the set $\mathcal{Y}$ contains categorical values, such as {hamster, mouse}. These are the *classes* of the classification problem. In the context of classification, a predictor is also called *classifier*. The most well-studied setting in learning theory is binary classification, where there are two classes that we can write as 0 and 1. Sometimes it's mathematically convenient to instead work with the numbers -1 and 1 for the two classes. A classification problem with more than two classes is a multi-class problem.
- In a regression problem, the label set $\mathcal{Y} = \mathbb{R}$ is the real line. Rather than predicting a discrete value, our goal is to estimate quantitative information, such as the income of a person, the temperature of an engine, or the stock price of a company share.

In a classification problem, we try to be correct, in a regression problem, we try to be close. Here are some examples of typical prediction problems:

(continued...)

Index

- A/B test, 301
- accuracy, 17. *See also* classification error
 - top-5, 117
- accuracy drop, 113, 119
- adaptive
 - analyst, 77
 - tree, 79
- adaptive data analysis, 77, 89
- adaptivity, xxi, 77. *See also* adaptive data analysis
- aggregation, 232
- agreement rate, 152
- AlexNet, xxvii, 117
- algorithmic stability, 71
- alignment, 7, 198
- Amazon Mechanical Turk, 119
- anchoring effect, 138
- annotation, 149
- annotator
 - disagreement, 150
 - fatigue, 152
- answer matching, 213, 282, 293
- anything goes, 1, 110, 171, 306
- arithmetic coding, 176
- Arora, Sanjeev, 139
- Arrow's impossibility theorem, 243
- art of conjecturing, 14
- astronomical conception of society, 13, 255
- attention
 - dot-product, 170
- Attention Is All You Need, 170
- attention mechanism, 170
- automatic evaluation, 206, 279
- autoregressive model, 172
- availability heuristic, 138
- bandwagon effect, 258
- base model, 199
- base rate, 23
 - fallacy, 24
- Bayes optimal, 273
- Bayes' rule, 23
- benchmark, xix, 3, 206
 - design, 146
 - dynamic, 11, 274
 - multi-task, 8, 231
 - saturation, 192
- benchmarking, 3
 - as process, 308
- benchmaxxing, 9
- Benjamini–Hochberg procedure, 101
- Bernoulli distribution, 34
- BERT, 178
- best-of-N, 205
- bias
 - dataset, 165
 - judge, 283
 - self-preferencing, 284
 - style, 284
 - verbosity, 205, 284
- BIG-Bench, 232
- binomial distribution, 37, 41
- bits-per-byte, 176
- bitter lesson, 116
- black box, 31, 54

- Black-Scholes, 186
- Blenheim Spaniel, 140
- BLEU score, 177, 279
- Bonferroni correction, 47, 101
- BookCorpus, 178
- boosting, 81
- bootstrap, 68
- Borda count, 242
- bounded rationality, 137
- Bradley–Terry model, 201, 204, 211, 237, 253
- bricolage, 149
- Brouwer’s fixed point theorem, 259
- byte-pair encoding, 171
- calibration, 20
 - group, 21
- Campbell’s law, 106
- causal inference, 264, 271
- chatbot, 7
- Chatbot Arena, 210, 284
- ChatGPT, 198
- Chebyshev’s inequality, 36
- Chernoff bound, 41
- Chinchilla scaling law, 183
- CIFAR-10, 4, 110, 119, 122
- CIFAR-100, 122
- classification, 15
 - binary, 15
 - fine-grained, 119
 - multi-class, 15, 141
- classification accuracy. *See* accuracy
- classification error, 17, 33
- classifier, 15
- CLIP, 192
- CodaLab, 72
- code re-execution. *See* reproducibility
- cognitive biases, 137
- Cohen’s κ , 153
- coin tossing, 33–37
- Common Crawl, 180
- common task framework, 71
- commonsense reasoning, 189
- competition, 72, 144, 306
- compression argument, 133, 139
- compute-optimal training, 184
- concentration inequality, 41
- conditional probability, 17
- Condorcet paradox, 240
- Condorcet winner, 240, 244
- Condorcet’s jury theorem, 156
- confidence interval, 65, 95
- confirmation bias, 137
- confusion matrix, 24
- context, 172
- convolutional neural networks, xxvii, 4
- corpus, 173
- counterfactual, 271
- credit scoring, 256
- cross-validation, xx, 62
 - k -fold, 63
 - leave-one-out, 65
- data
 - contamination, 194, 220
 - curation, 180
 - feedback loop, 11, 270–274
 - labeling, 149
 - leakage, 128, 220
- data split, 51
 - three-way, 60, 62
- data-generating distribution, 15, 51, 70
- dataset
 - bias, 113, 273
- Dawid-Skene model, 167
- debiasing, 280
- deep learning, xxvii, 4, 117
- deep learning revolution, xx, 4
- differential privacy, 85
- direct evaluation, 218
- direct sum theorem, 40
- disagreement rate, 46, 143
- distillation, 201
- distribution map, 259
- distribution shift, 14, 114, 124, 265
- domain adaptation, 115
- domain generalization, 115
- DomainBed, 115
- Donoho, David, 143
- DPO, 204

- Durand, William, 305
- Dynabench, 274
- dynamical system, 138
- echo chamber, 270
- effect size, 95
- Efros, Alyosha, 113
- Elo rating, 211, 237. *See also*
 - Bradley–Terry model
- emergent abilities, xxii, 8, 179, 186
- empirical risk, 27
 - minimization, 27
- ensembling, 80, 145
- entropy, 28
- equilibrium, 259
- error bars, 55
- evaluation frontier, 280
- expected value, 19
- experimental parameter variation, 305
- exploding gradients, 169
- external validity, 70, 112, 113, 222
- F*-test, 86
- F_1 -score, 25
- fairness, 271
- false discovery. *See* false positive
- false negative, 24
- false positive, 24, 94
- false positive rate, 24, 96
- feature, 15
- feedback loop, 76, 270
- few-shot learning, 179
- Feyerabend, Paul, 1
- filter bubble, 270
- fine-tuning, xxii, 188. *See also*
 - supervised fine-tuning
- fixed point, 259, 262. *See also*
 - equilibrium
- Fleiss’ κ , 154
- FLOPs, 180
- Forrester, Jay, 274
- foundation models, 195
- Freedman’s paradox, xxi, 85
- frictionless reproducibility, 143
- frontier model, 279
- FrontierMath, 298
- game theory, 257
- gaming, 4, 103
- garden of the forking paths, 101
- Gelman, Andrew, 101
- generalization, xxiv, 27
- generalization bounds, 56
- generalization gap, 27, 53, 80
- generative model, xxii
- ghost work, 150
- Gibbard–Satterthwaite theorem, 245
- Gigerenzer, Gerd, 139
- GLUE, 191
- Goodhart’s law, xix, 102, 110, 125, 128, 258, 306
- Google Brain, xxvii
- GPQA, 234
- GPT, 178
- GPT-1, 178
- GPT-2, 179
- GPT-3, 179
- GPT-4, 179
- GPU, 117, 140, 169
- gradient, 28
- gradient boosting, 145
- ground truth, xxiii, 10, 28, 155, 257, 264
- GRPO, 206
- Grunberg, Emile, 258
- GSM8k, 214
- Hacking, Ian, 13
- hallucination, 7, 197, 209
- Hayek, Friedrich, 1
- HealthBench, 295
- Hellinger distance, 39
- HELM, 235
- heuristics, 137, 139
- hindsight bias, 137
- Hinton, Geoffrey, xxvii, 117
- Hoeffding’s inequality, 43
- holdout method, xx, 51
- holdout set, 54
- homo economicus, 137
- Hugging Face, 145, 234
- human evaluation, 210, 279
- hyperparameters, xxvii
- hypothesis testing, 86, 92

- IIA. *See* independence of irrelevant alternatives
- iid assumption, 33
- ILSVRC, 117, 140
- image classification, 4
- ImageNet, xx, xxvii, 4, 110, 117–121, 145
 - V2, 126
- ImageNet era. *See* ImageNet
- ImageNot, 5, 122
- Imitation Game, 189
- in-context learning, 179
- incommensurability, 233
- independence of irrelevant alternatives, 243, 248
- indicator function, 20
- inequality, 307
- inference after selection, 88
- instruction tuning, 200
- inter-rater reliability, 152
- internal validity, 69, 112
- Ioannidis, John, 96
- iron rule, 3, 145, 178, 305
- iron vault, 58
- jackknife, 65
- judge model, 280
- Kaggle, 60, 144
 - competitions, 121
- Kahneman, Daniel, 98, 137
- Kendall's τ , 223
- KL divergence. *See* Kullback–Leibler divergence
- knowledge machine, 3
- Krippendorff's α , 154
- Krizhevsky, Alex, 117
- Kullback–Leibler divergence, 28
- label, 15
 - aggregation, 162
 - error, 158
 - noise, 150
- labeling. *See* data, labeling
- Ladder algorithm, 132
- LAION, 122, 193
- LAMBADA, 190
- language model, xxii, 7, 172, 173
- language modeling, 172
- latent construct, 165
- law of iterated expectation, 23
- leaderboard, xx, 6, 59, 118, 211
 - climbing, 80
- leaderboard error, 131
- leakage, 128
- learning rate, xxvii
- least squares, 86
- Lesley, Everett, 305
- Li, Fei-Fei, 118
- likelihood, 23, 172
- Llama, 198
- LLM-as-a-judge, 11, 213, 280
- lookup table, 52
- loss
 - cross-entropy, 28
 - function, 22, 132
 - logistic, 28
 - squared, 19, 22
 - zero-one, 22, 132, 159
- LSTM, 169
- Lucas critique, 258. *See also* Goodhart's law
- majority vote, 80, 155–158
 - iterated, 164
- Markov's inequality, 41
- MATH, 234
- McDiarmid's inequality, 49
- mean squared error, 19
- measurement, 61, 165
 - cardinal, 124
 - ordinal, 124
 - validity, 61
- Meehl, Paul, 98
- memorization, 52, 58
- minimax lower bound, 49
- Mixture of Experts, 235
- MMLU, xx, 8, 207, 300
- MMLU-Pro, 234
- MNIST, 112, 121
- model, 26
 - deployment, 260

- evaluation, 26
- family, 27
- selection, 56
- training, 26
- model collapse, 272
- model comparison, 31
- model merging, 235
- model rankings, xxi, 5, 117, 119, 124, 307
 - agreement, 222
 - stability, 5
- model similarity, 143
- Modigliani, Franco, 258
- modus tollens, 99
- Moore's Law, 180
- Morgenstern, Oskar, 257
- multi-class, 143
- multiple choice evaluation, 206
- natural language processing, 6, 169
- nearest neighbor classifier, 52
- negative log-likelihood, 173
- Nelson, Leif, 99
- Netflix Prize, 145
- von Neumann, John, 257
- NeurIPS, 110
 - experiment, 110
- next-token prediction, 172
- NLP. *See* natural language processing
- normal distribution, 41
- normalization, 234
- null hypothesis significance testing, 92
- null model, 86
- ObjectNet, 122
- on the line phenomenon, 120
- one-hot encoding, 28
- Open Catalyst, 308
- open source, 111
- OpenAI, xxviii, 178, 299
- OpenLLM leaderboard, 234
- OpenML, 72
- operations management, 308
- optimization, 17
- outcome performativity, 265
- outcome variable, 260
- overfitting, xix, 4, 54, 129. *See also*
 - training on the test set
 - adaptive, 80, 121
- p-hacking, 88, 101
- p-value, 87
- paired comparison, 46
- pairwise comparison, 240
- PapersWithCode, 4
- parameter count, 178
- pattern recognition, 1
- peer review, 110
- Penn Treebank, 178, 306
- perceptrons, 1
- performative
 - effect, 264, 265
 - prediction, 258, 259
 - risk, 263
 - stability, 261, 262
- performativity, 11, 256
- perplexity, 174
- Plackett-Luce model, 253
- plurality vote, 141, 155
- polarization, 270
- Popper, Karl, 98
- population, 15, 70, 264
 - steering, 264
- positive predictive value, 25, 96
- post-training, xxii, 7, 199
- posterior probability, 18
- power law, 183, 185
- PPI. *See* prediction-powered inference
- pre-registration, 100
- precision, 25, 96. *See also* recall
- precision-recall trade-off, 25
- prediction, 15
- prediction-powered inference, 288
- predictive policing, 271
- predictor, 15
- preference learning, 201
- preference model, 201
- pretraining, 173, 199
- Principles of Neurodynamics, 2
- probability distribution, 13
- prompt, 172. *See also* prompt engineering

346 INDEX

- prompt engineering, 173, 179
- prompt injection, 287
- prompting, 198
 - few-shot, 179
- proper scoring rule, 29
- protein structure prediction, 308
- proxy score, 281
- publication bias, 97
- PyTorch, 145
- quality control, 308
- Quetelet, Adolphe, 13
- random variable, 15
- rationality, 137
- recall, 24. *See also* precision
- recency bias, 138
- recommender system, 256, 270
- recurrent neural networks, 6, 169
- reference score, 281
- regression, 15
 - after selection, 88
 - function, 19
 - linear, 13, 86
- reinforcement learning from human
 - feedback. *See* RLHF
- replication, 111–113
- replication crisis, 91
- representative agent, 137
- reproducibility, 111
 - code re-execution, 111
 - different conditions, 112
 - NeurIPS challenge, 111
- researcher degrees of freedom, xxi, 99
- residual networks, 109
- ResNet. *See* residual networks
- response function, 259
- retraining, 262, 265
- reward model, 201
- Rip Van Winkle, 139
- risk, 22
 - minimization, 22
- RLHF, 203
- RLVR, 206
- RNN. *See* recurrent neural networks
- ROC curve, 25
- Rosenblatt, Frank, 1
- rubrics, 213, 281, 295–297
- safety evaluation, 281
- sample average, 32
- sample size, 33, 55
- sample splitting, 88
- scaling law, xxii, 6, 181, 182, 196, 215, 225
- scatter plot, 120
- scientific crisis, 11, 100
- Scott’s π , 153
- selection frequency, 119
- self-attention, 170
- self-fulfilling prophecy, 16, 186, 256. *See also* performativity
- self-supervised, 192
- SFT. *See* supervised fine-tuning
- shortcut learning, 209
- sigmoid, 201
- significance level, 87
- Simmons, Joseph, 99
- Simon, Herbert, 137, 258
- Simonsohn, Uri, 99
- simulation, 273, 301
- SNLI, 188
- social algorithm, 136
- social choice theory, xxii, 238–239, 239, 247
- social network, 185
- social physics, 13
- softmax, 28, 201
- Spendolini, Michael, 308
- SQuAD, 188
- stable signal, 265
- standard deviation, 36
- state-of-the-art, 130
- stationarity, 255
- statistical power, 94
- strategic behavior, 244, 260
- strategic classification, 128
- Strathern, Marilyn, 125
- Stevens, Michael, 3, 305
- stride, 177
- subgaussian, 41

- SuperGLUE, 191
- supervised fine-tuning, 200
- supervised learning, xxiv, 26
- Sutskever, Ilya, xxvii, 117
- SWE-bench, 297
- SWE-Lancer, 299
- synthetic data, 272
- system dynamics, 273

- T*-statistic, 88
- tail bounds, 41
- target variable, 15, 149, 156
- temperature, 172
- TensorFlow, 145
- test error, 53
- test set, xix, 3, 53
- test set reconstruction, 119
- test set reuse, 75, 82
- test statistic, 86
- token, 171
- tokenization, 171
- Torralba, Antonio, 113
- total variation distance, 38
- training compute, 180
- training data, 27
- training error, 53
- training objective, 27
- training on the test set, 9, 58, 194
- training on the test task, 9–10, 217–219
- training set, xix, 3, 53
- transfer learning, 115
- transformer, 6, 170, 306
- true positive rate, 24, 94

- TruthfulQA, 208
- tune-before-test, 219
- Turing Test, 189

- underdog effect, 258
- union bound, 44
- universe, 15
- unsupervised learning, 26

- validation set, 63
- validity
 - external. *See* external validity
 - internal. *See* internal validity
- vanishing gradients, 169
- variable selection, 85
- variance, 36
- verifiable rewards, 205
- verification, 282, 297
- Vincenti, Walter, 305
- vocabulary size, 171
- voting
 - ranked, 239
 - rated, 243
- voting system, 163, 232, 238

- weather prediction, 308
- web crawl, 180
- WILDS, 115
- win rate, 235
- Winograd Schema Challenge, 189
- Winogrande, 190
- WordNet, 118

- zero-shot, 179, 194