

Learning in a Dynamic and Strategic World

Thesis Defense: Geelon So; on May 29, 2026

Committee: Sanjoy Dasgupta (co-chair), Yian Ma (co-chair), Russell Impagliazzo, Daniel Kane, Barna Saha

Roadmap

A. Motivation

B. Research Overviews

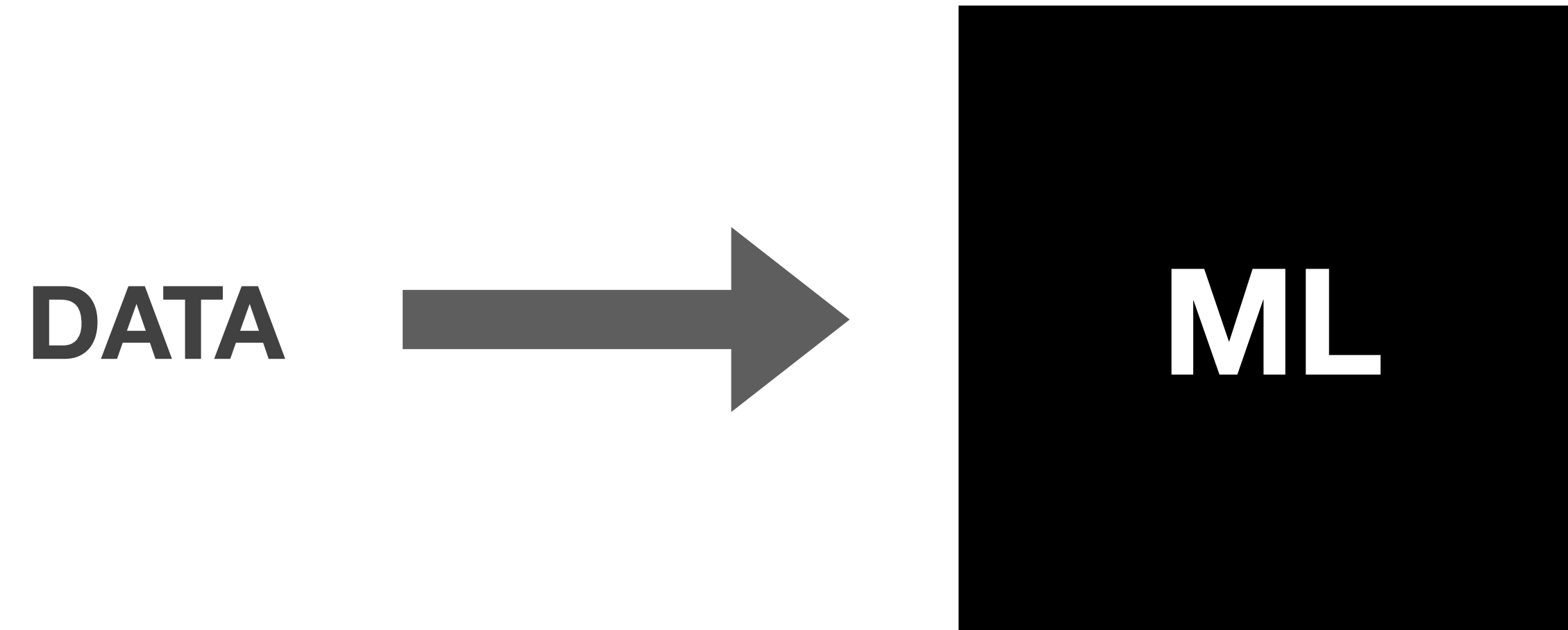
MOTIVATION

Two Views of Data

A functional description of machine learning



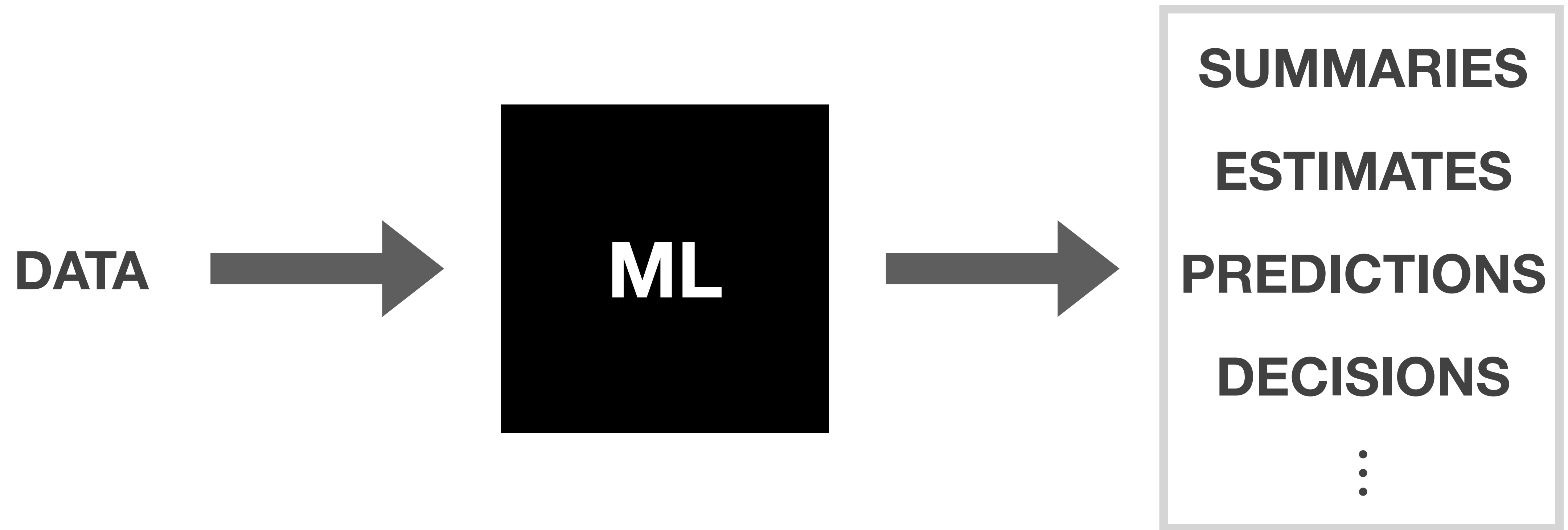
A functional description of machine learning



A functional description of machine learning



A functional description of machine learning



A functional description of machine learning



Where does the data come from?

A functional description of machine learning



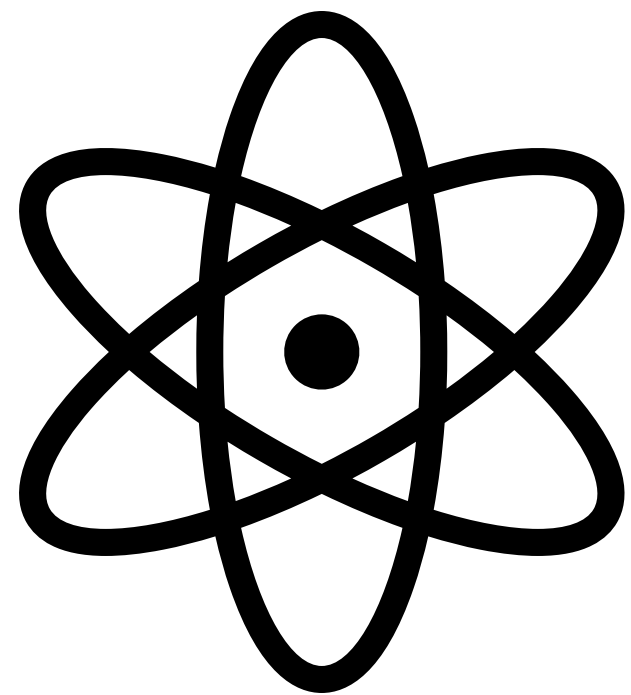
Where does the data come from?

Two broad settings.

An idealized view of data

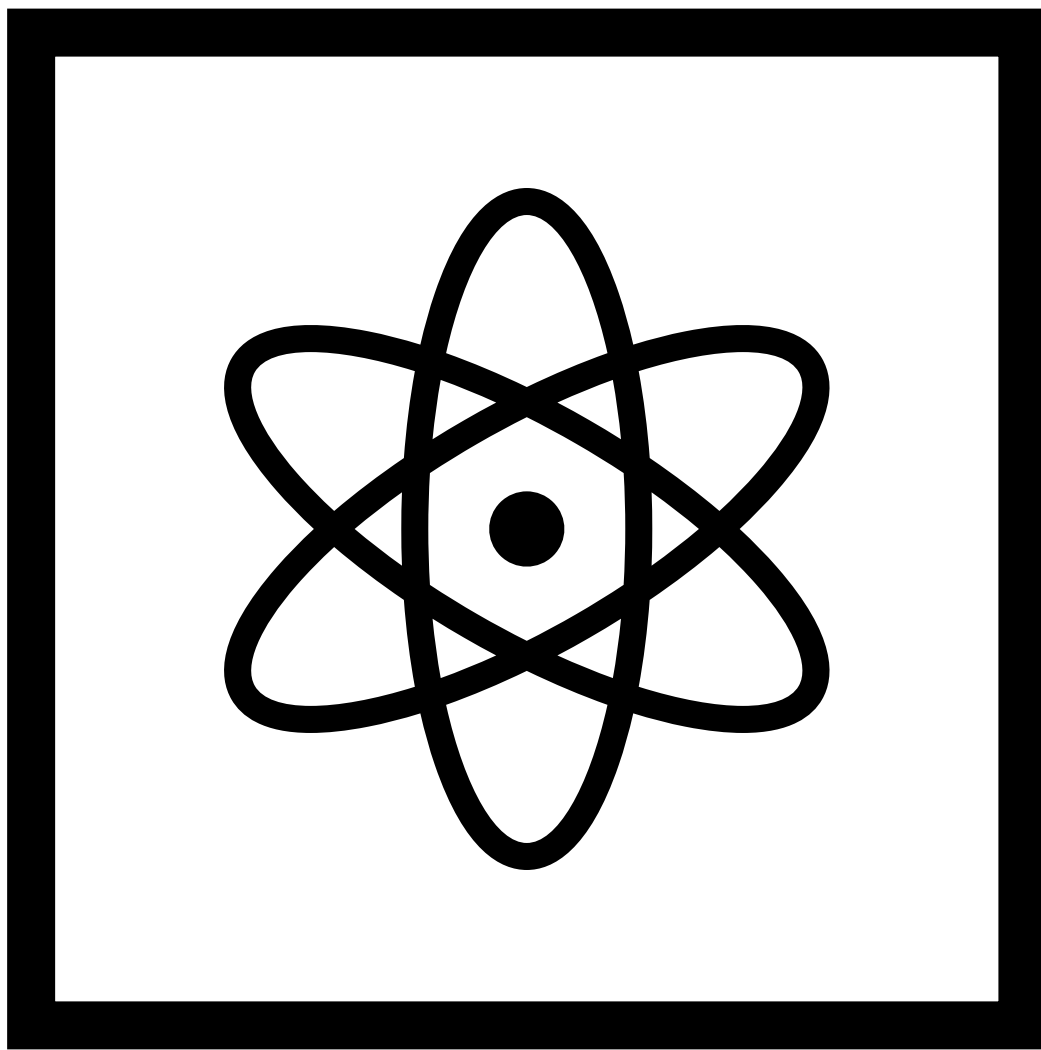
An idealized view of data: learning in the lab

An idealized view of data: learning in the lab



Interesting
phenomenon

An idealized view of data: learning in the lab

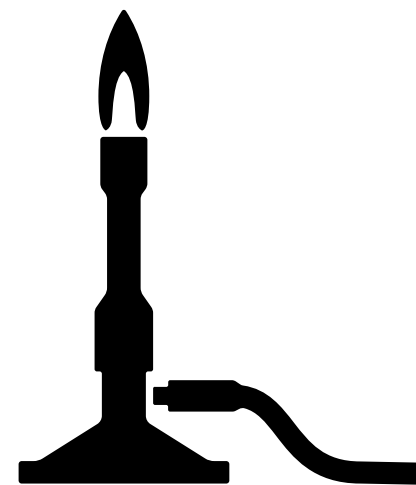
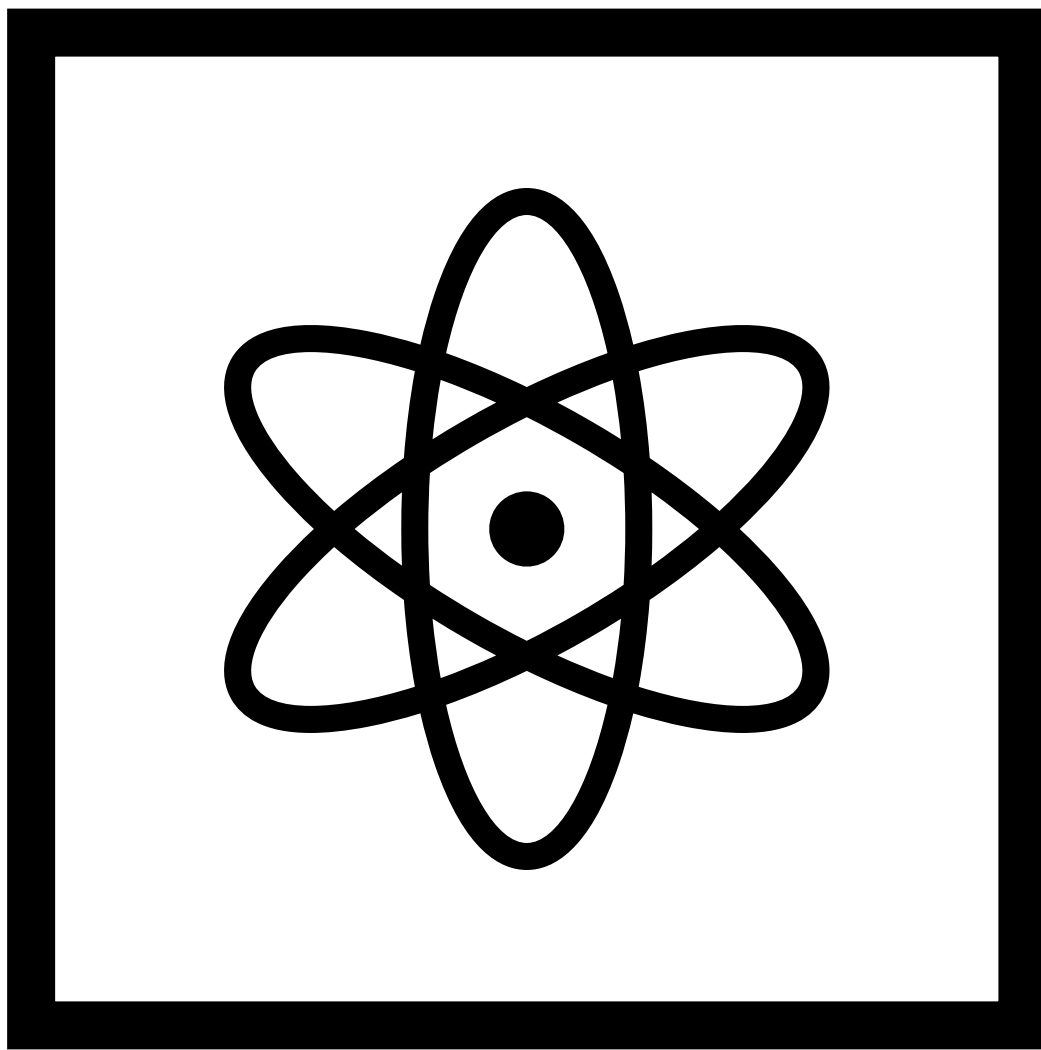


Interesting
phenomenon



isolate it

An idealized view of data: learning in the lab



Interesting
phenomenon

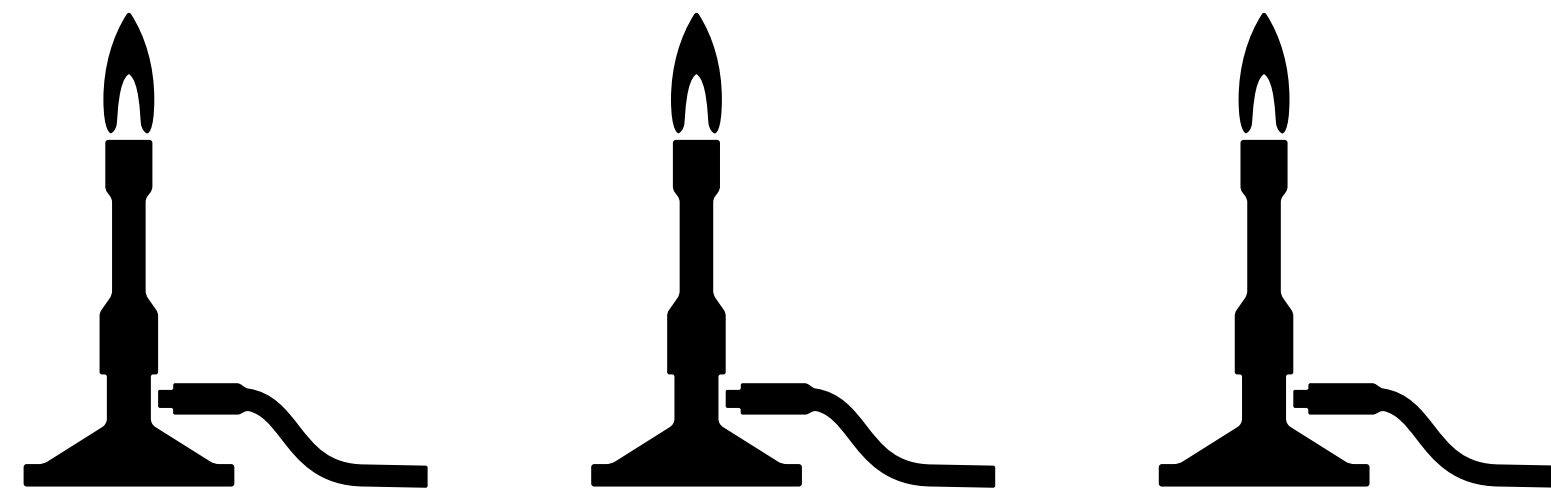
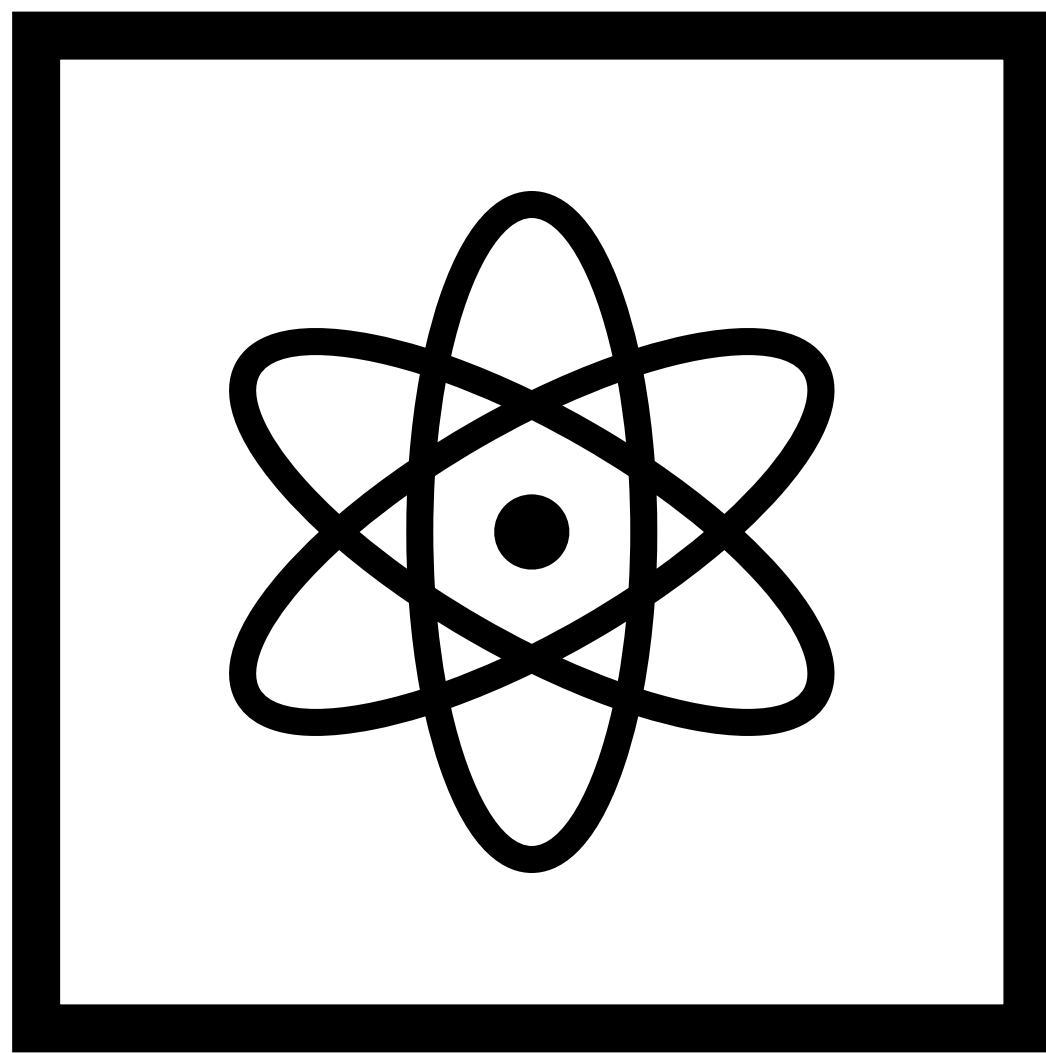


isolate it



experiment on it

An idealized view of data: learning in the lab



Interesting
phenomenon

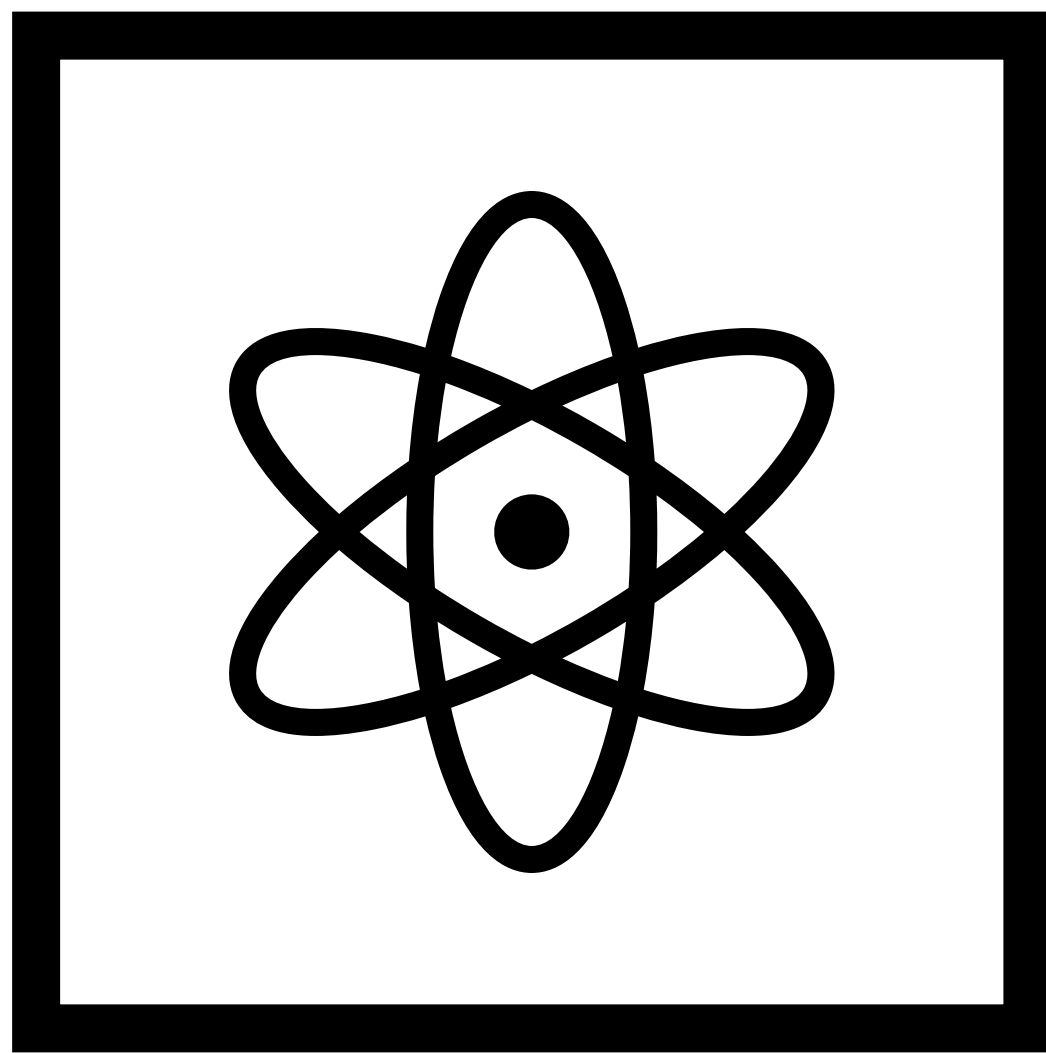


isolate it



experiment on it
repeatedly

An idealized view of data: learning in the lab



Interesting
phenomenon



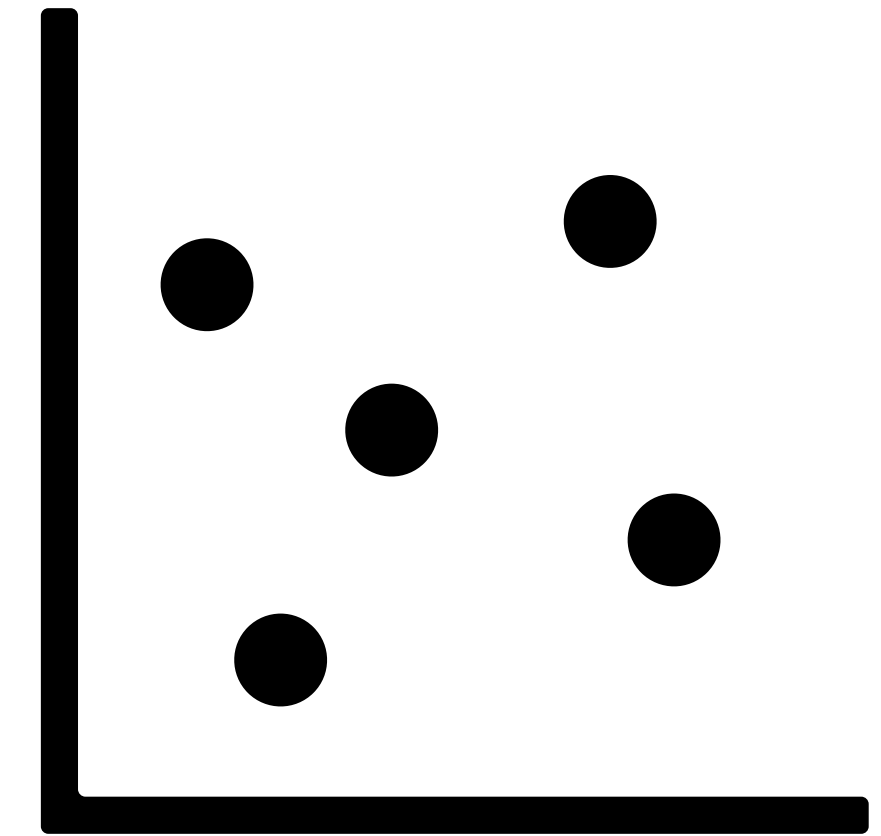
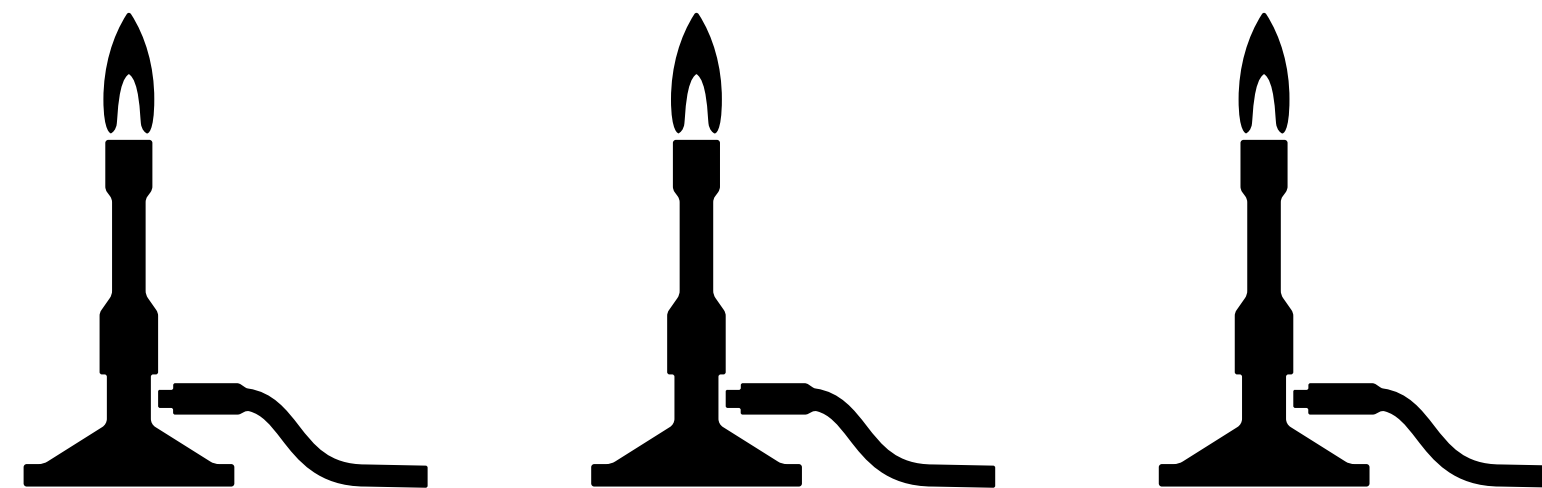
isolate it



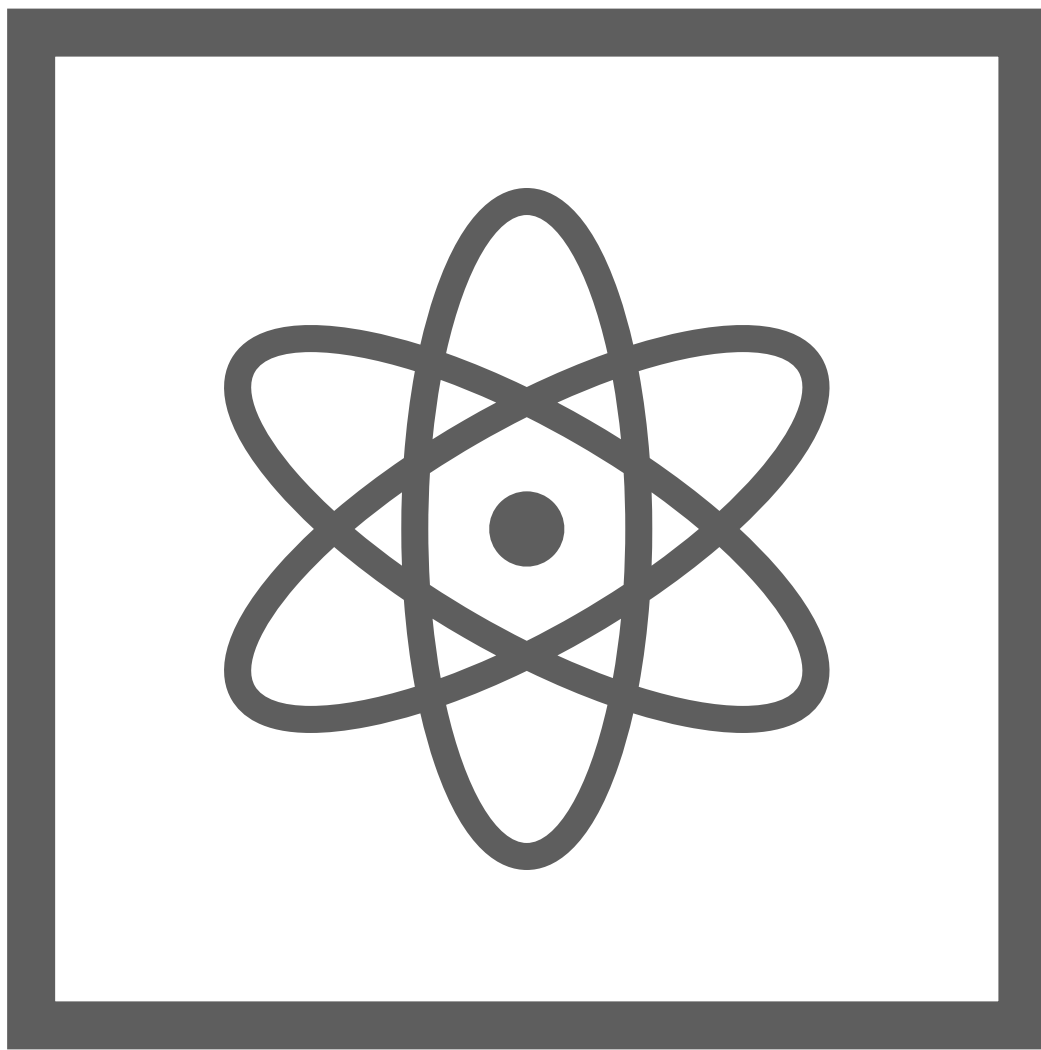
experiment on it
repeatedly



perform
statistical
analysis

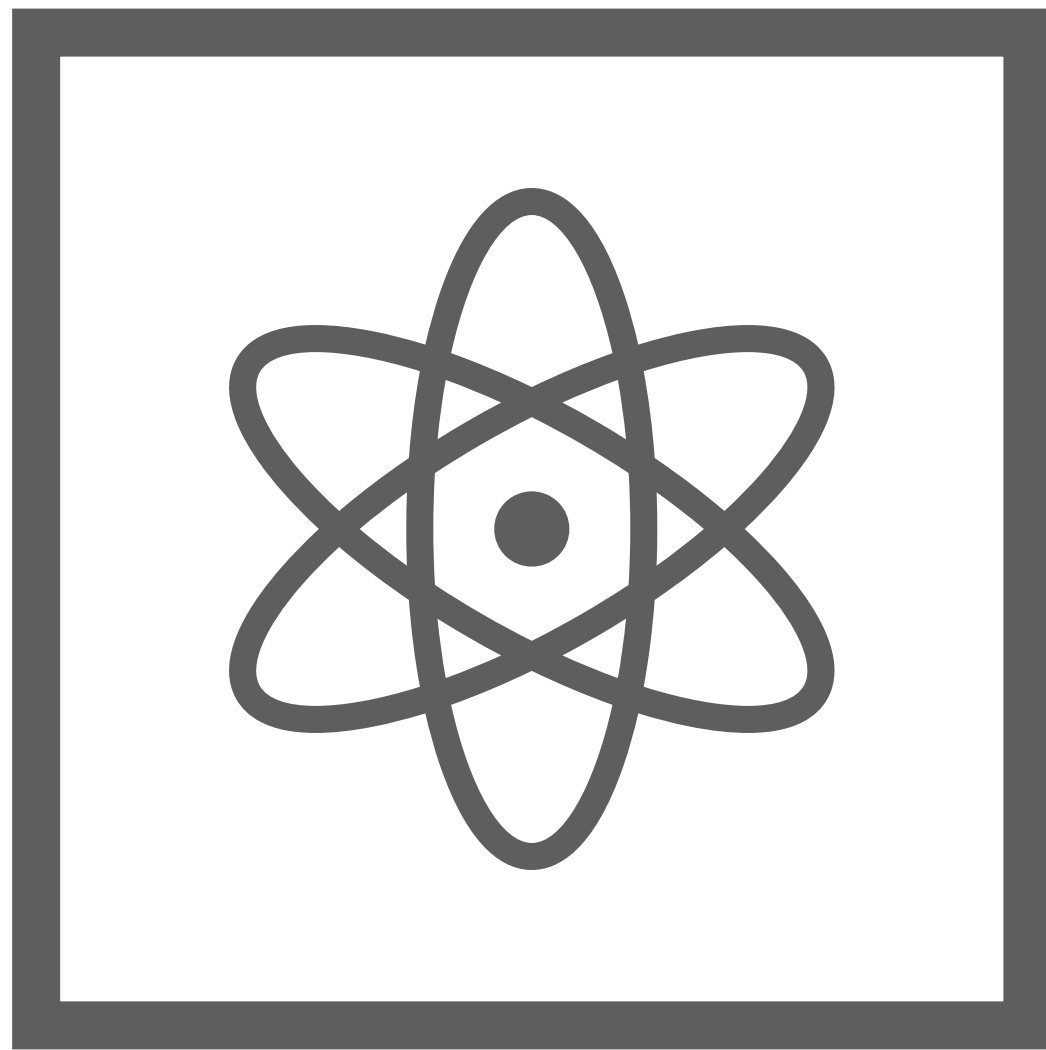


Learning in the lab



Rigorous foundation
for modern science

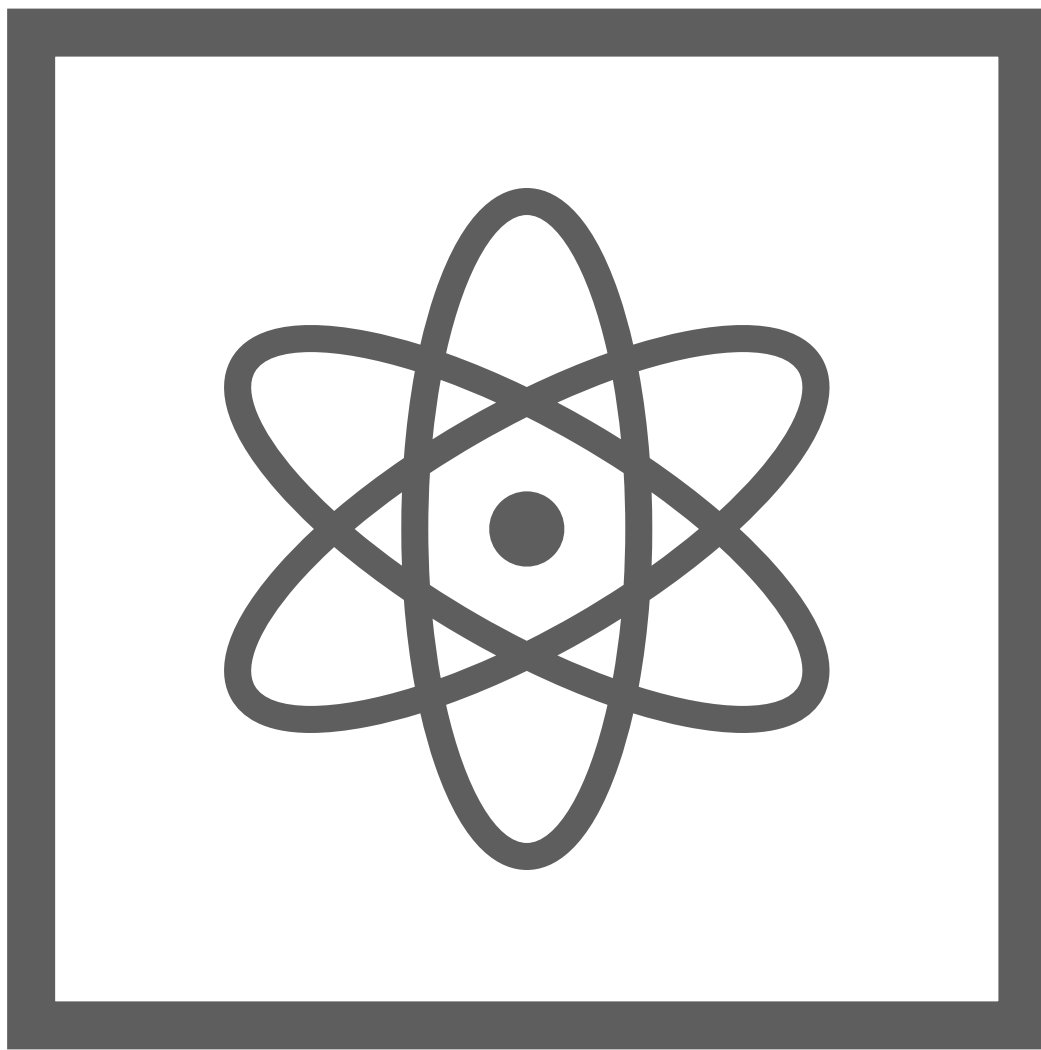
Learning in the lab



Rigorous foundation
for modern science

- **Curated, high-quality data**
 - Repeated trials, independent observations
 - Ability to intervene/experiment
- **Controlled environments**
 - The world is fairly stationary
 - The observer is outside of the system

Learning in the lab

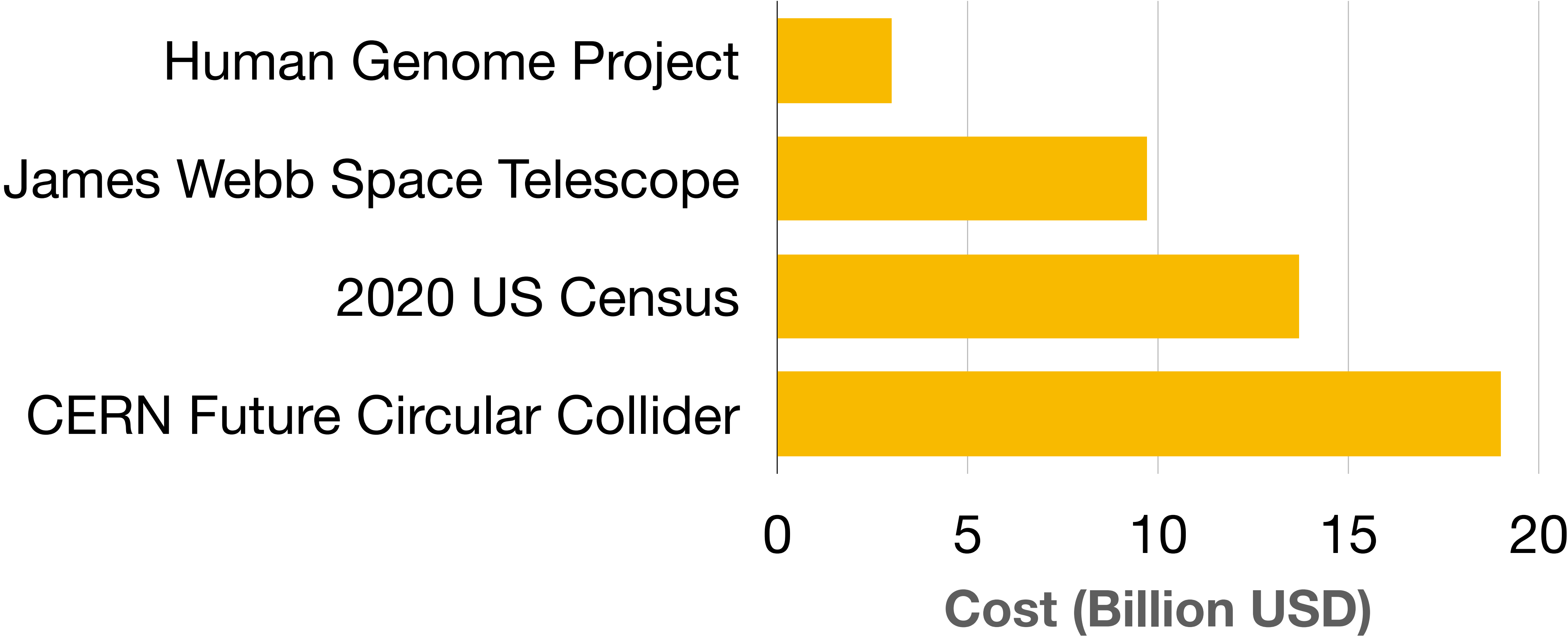


Rigorous foundation
for modern science

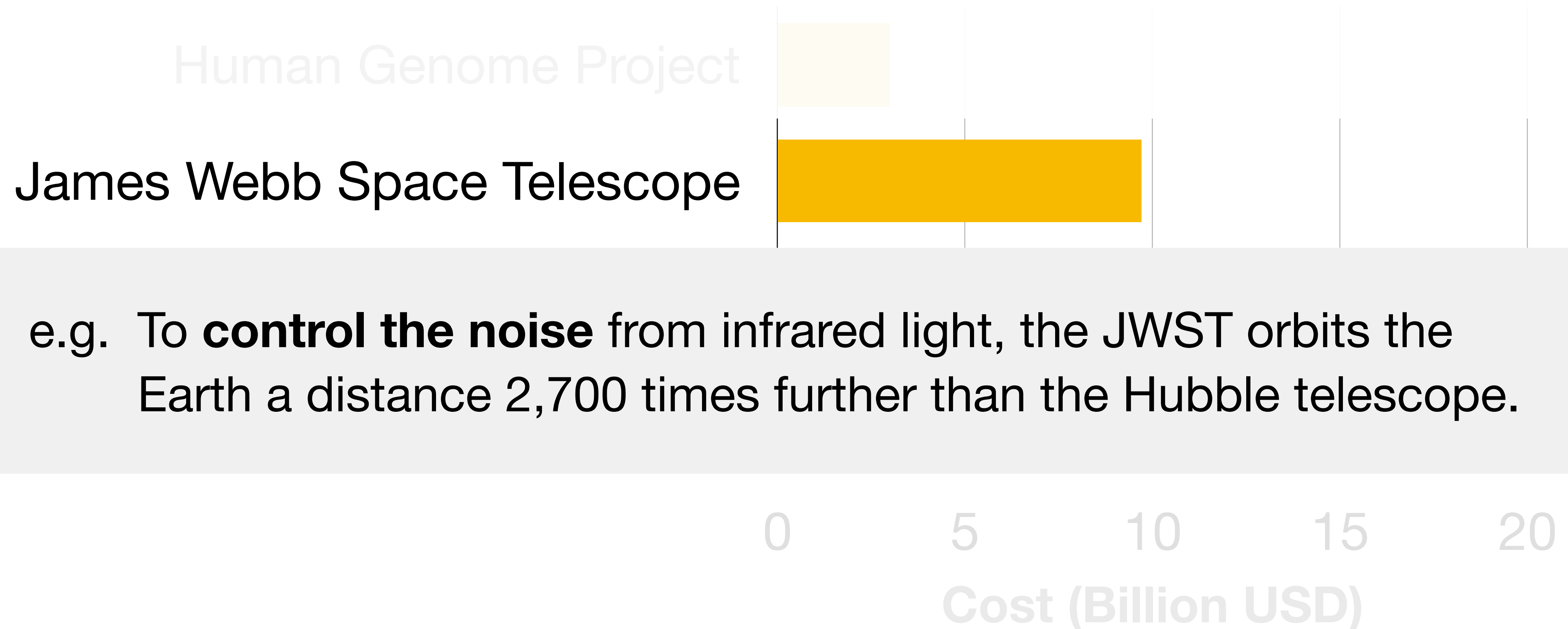
- **Curated, high-quality data**
 - Repeated trials, independent observations
 - Ability to intervene/experiment
- **Controlled environments**
 - The world is fairly stationary
 - The observer is outside of the system
- **Well-specified/specific tasks**
 - Classification, regression, hypothesis testing
 - Clear goal of high accuracy

Generating high-quality data is expensive

Generating high-quality data is expensive



Generating high-quality data is expensive



Generating high-quality data is expensive



e.g. To **control the noise** from infrared light, the JWST orbits the Earth a distance 2,700 times further than the Hubble telescope.

Aside: How much of the cost of data is due to strong assumptions of statistical methods?

DATA



ML



TASK

But, lots of interesting ML problems are not driven by clean data.

A street view of data

A street view of data: learning in the wild

A street view of data: learning in the wild

There are ongoing processes...

A street view of data: learning in the wild



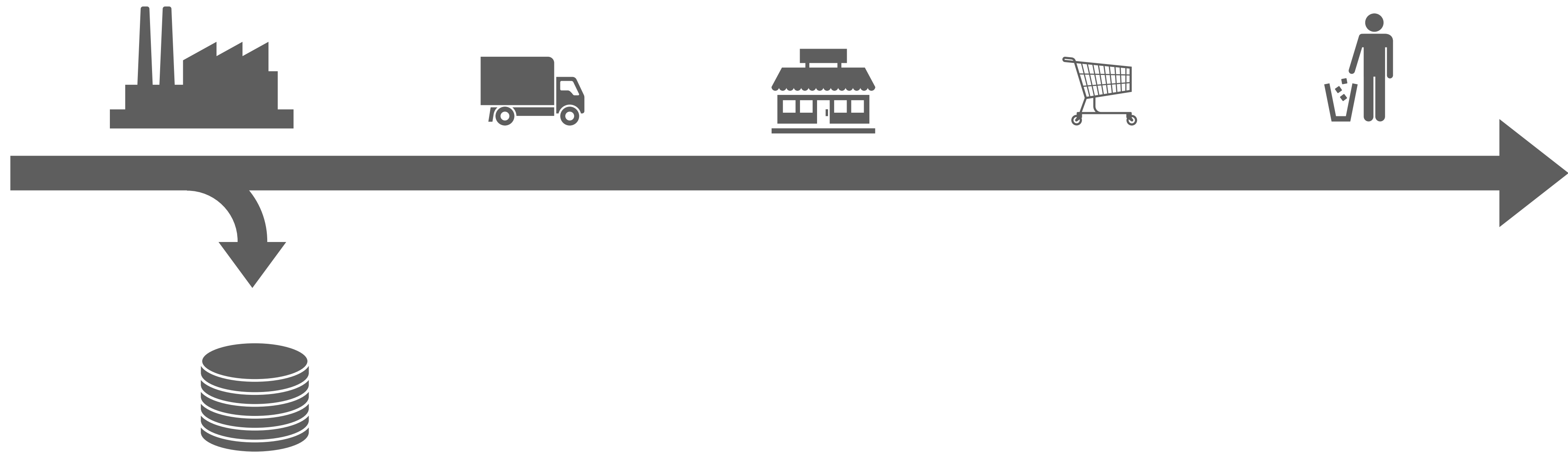
There are ongoing processes...

A street view of data: learning in the wild



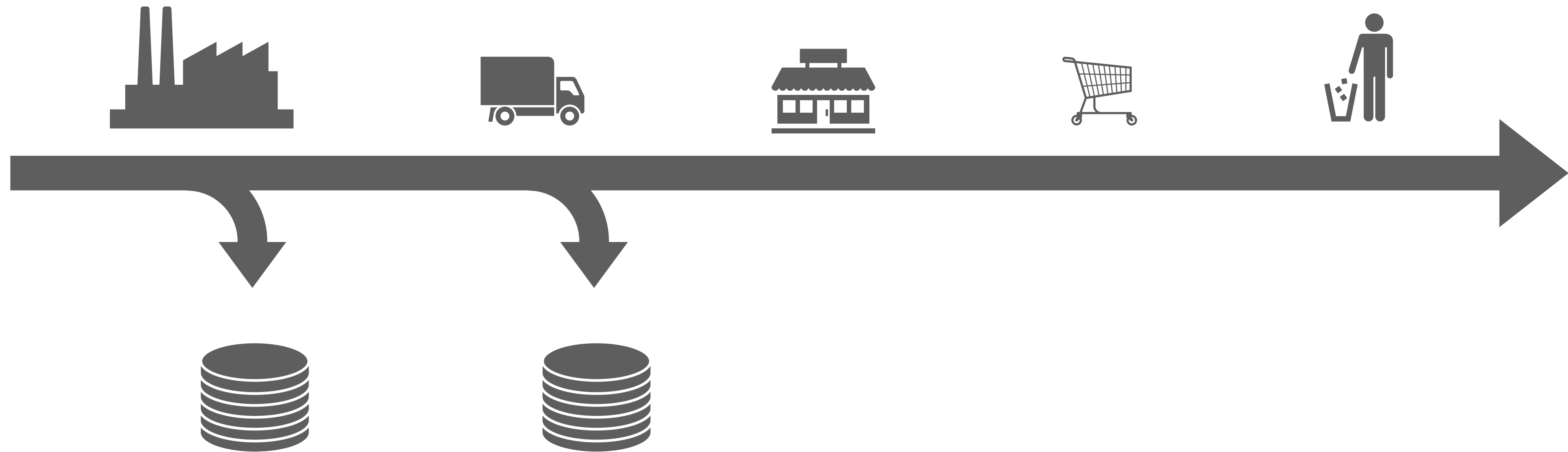
There are ongoing **social** processes...

A street view of data: learning in the wild



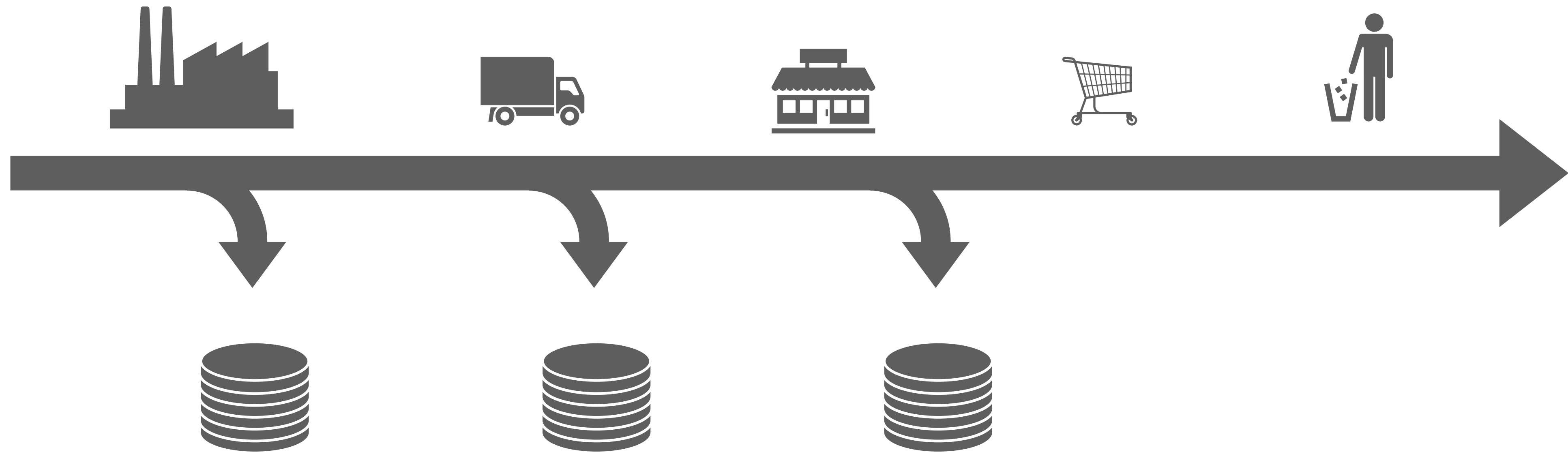
There are ongoing social processes that **emit data as a byproduct.**

A street view of data: learning in the wild



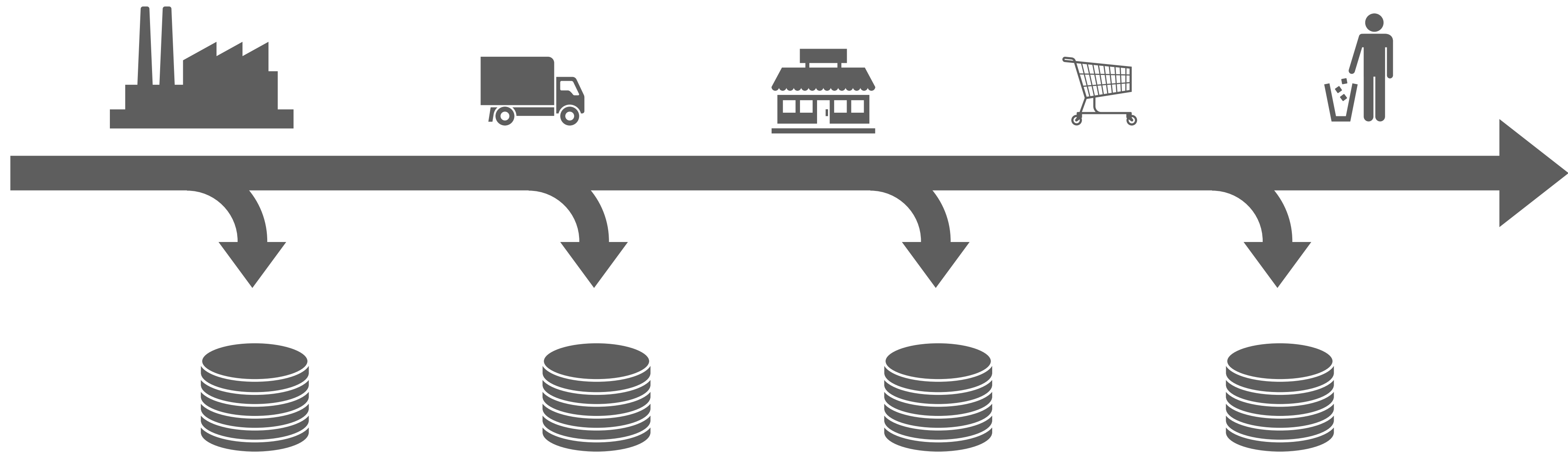
There are ongoing social processes that **emit data as a byproduct.**

A street view of data: learning in the wild



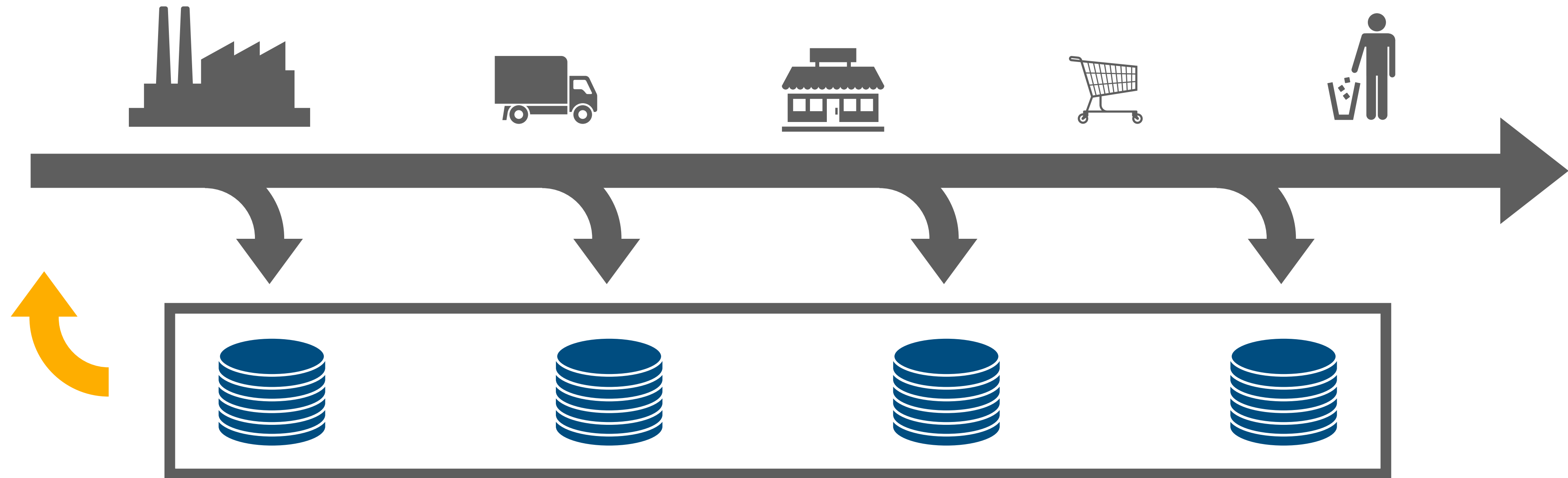
There are ongoing social processes that **emit data as a byproduct.**

A street view of data: learning in the wild



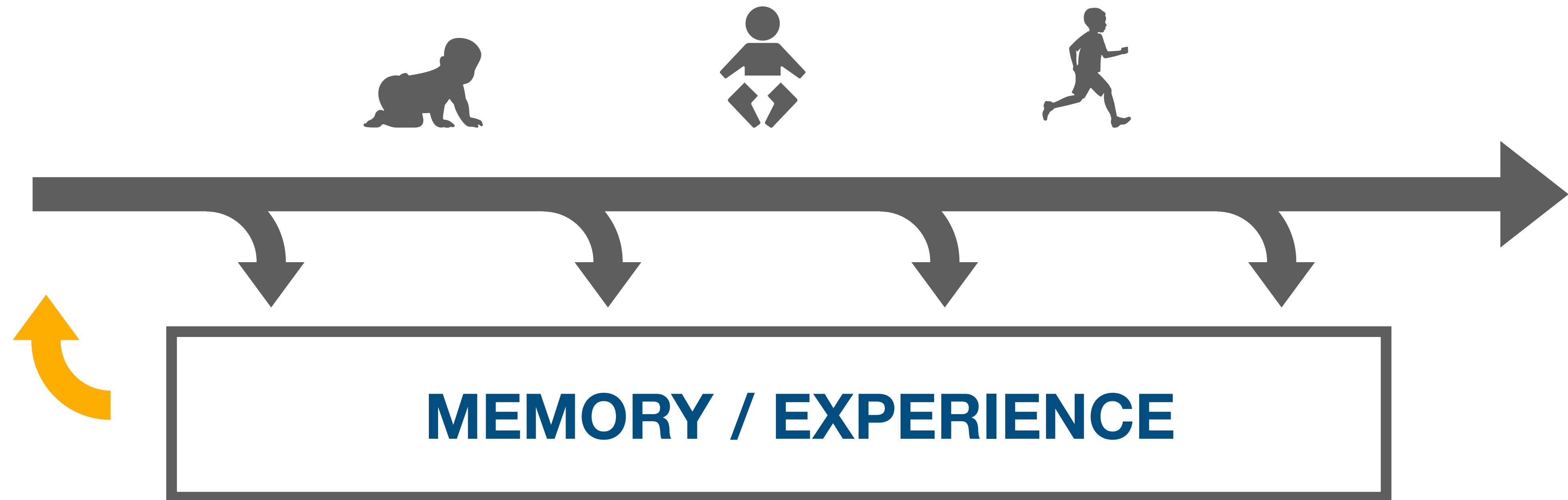
There are ongoing social processes that **emit data as a byproduct.**

A street view of data: learning in the wild

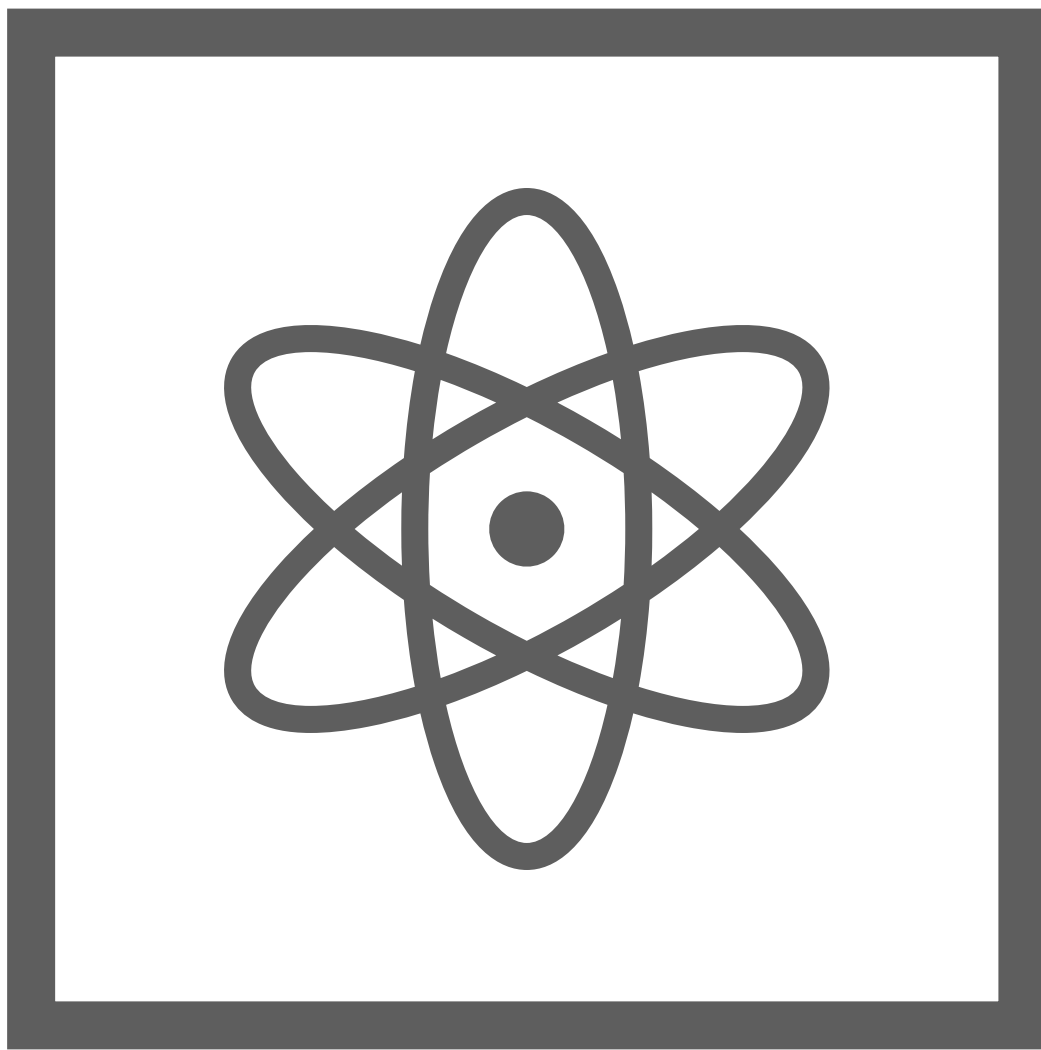


This data is often used to optimize the processes generating it.

A street view of data: learning in the wild

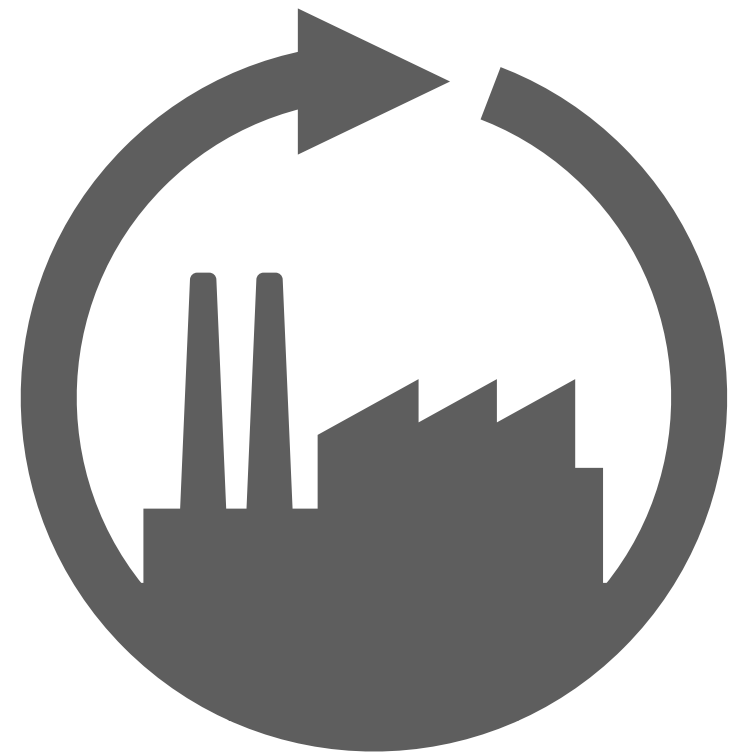


Learning in the lab



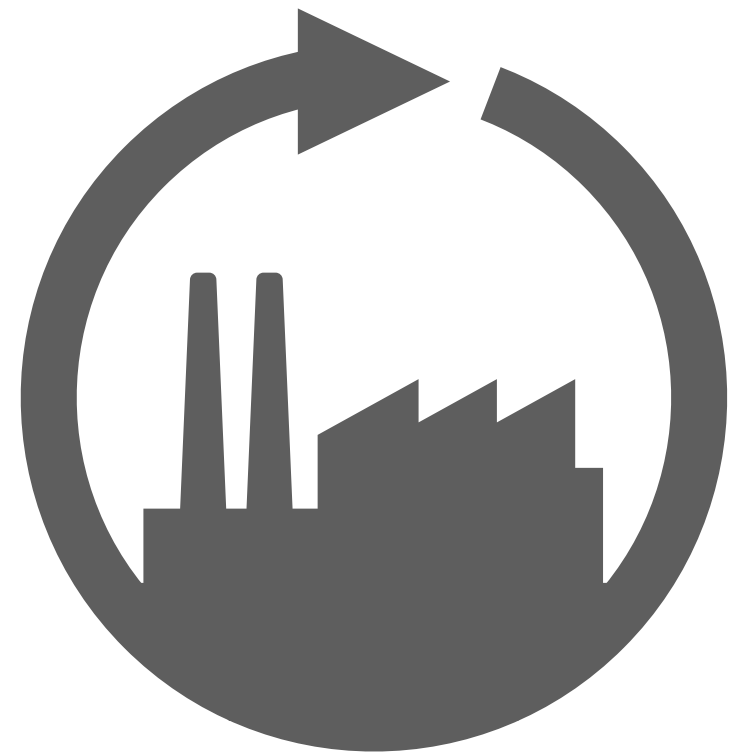
- **Curated, high-quality data**
 - Repeated trials, Independent observations
 - Ability to intervene/experiment
- **Controlled environments**
 - The world is fairly stationary
 - The observer is outside of the system
- **Well-specified/specific tasks**
 - Classification, regression, hypothesis testing
 - Clear goal of high accuracy

Learning ~~in the lab~~ in the wild



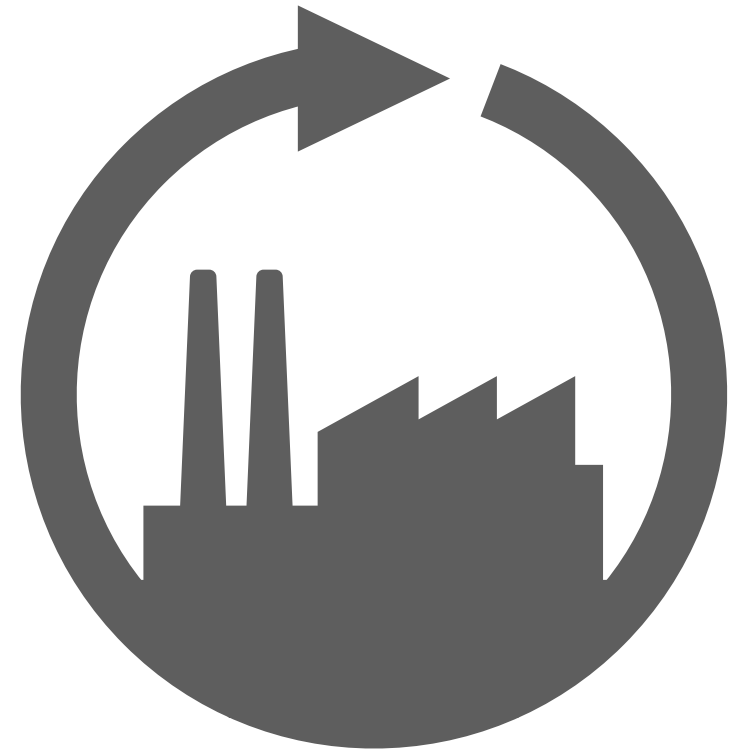
- **Curated, high-quality data**
 - Repeated trials, independent observations
 - Ability to intervene/experiment
- **Controlled environments**
 - The world is fairly stationary
 - The observer is outside of the system
- **Well-specified/specific tasks**
 - Classification, regression, hypothesis testing
 - Clear goal of high accuracy

Learning ~~in the lab~~ in the wild



- **Weak curation, low-quality (cheap?) data**
 - Data is adaptively generated, repeated trials often impossible
 - Costly interventions/experiments in production environments
- **Limited control over environment**
 - The world is non-stationary and strategic
 - The observer is part of the system (feedback loops exist)
- **Well-specified/specific tasks**
 - Classification, regression, hypothesis testing
 - Clear goal of high accuracy

Learning ~~in the lab~~ in the wild



- **Weak curation, low-quality (cheap?) data**
 - Data is adaptively generated, repeated trials often impossible
 - Costly interventions/experiments in production environments
- **Limited control over environment**
 - The world is non-stationary and strategic
 - The observer is part of the system (feedback loops exist)
- **Nebulous/general/complex tasks**
 - Make good decisions, help end-users, improve social welfare
 - End-goals are highly multi-objective

Complex Tasks

SFGATE

CENTRAL COAST

The heavily congested Calif. highway now controlled by AI

'I would like to find out as soon as possible if AI works'

By Andrew Pridgen, Central California Editor

May 7, 2026

Visitors to the Monterey Peninsula this summer may be unwittingly participating in an AI experiment to determine how fast or slow they arrive at their destination.

It is a big bet that county transportation officials, in conjunction with Caltrans, are making: an AI signal system pilot program that could eliminate the need for roundabouts along the same corridor.

Complex Tasks

SFGATE

CENTRAL COAST

The heavily congested Calif. highway now controlled by AI

'I would like to find out as soon as possible if AI works'

By **Andrew Pridgen**, Central California Editor
May 7, 2026

f

X

🦋

✉

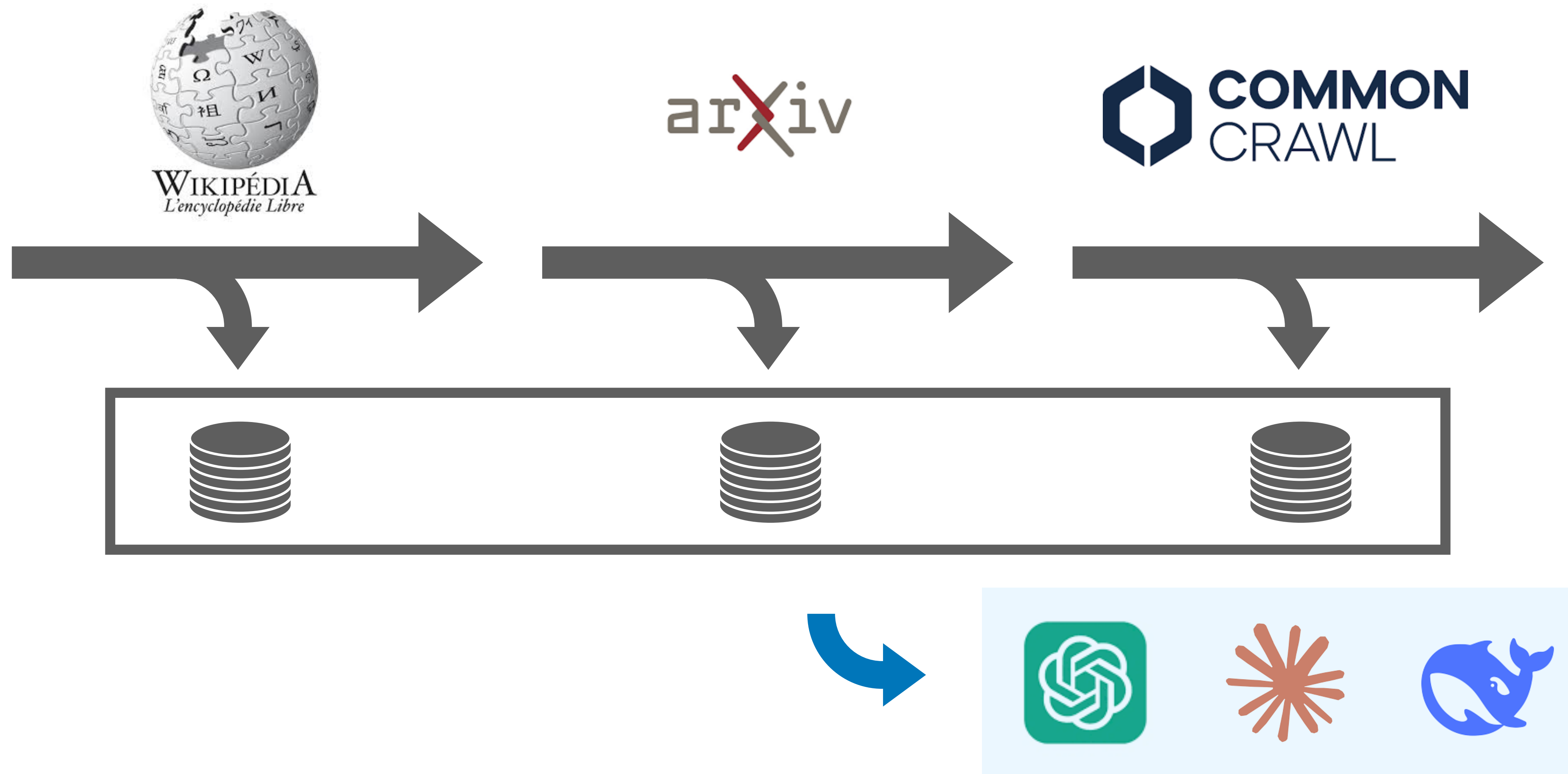
Visitors to the Monterey Peninsula this summer may be unwittingly participating in an AI experiment to determine how fast or slow they arrive at their destination.

It is a big bet that county transportation officials, in conjunction with Caltrans, are making: an AI signal system pilot program that could eliminate the need for roundabouts along the same corridor.

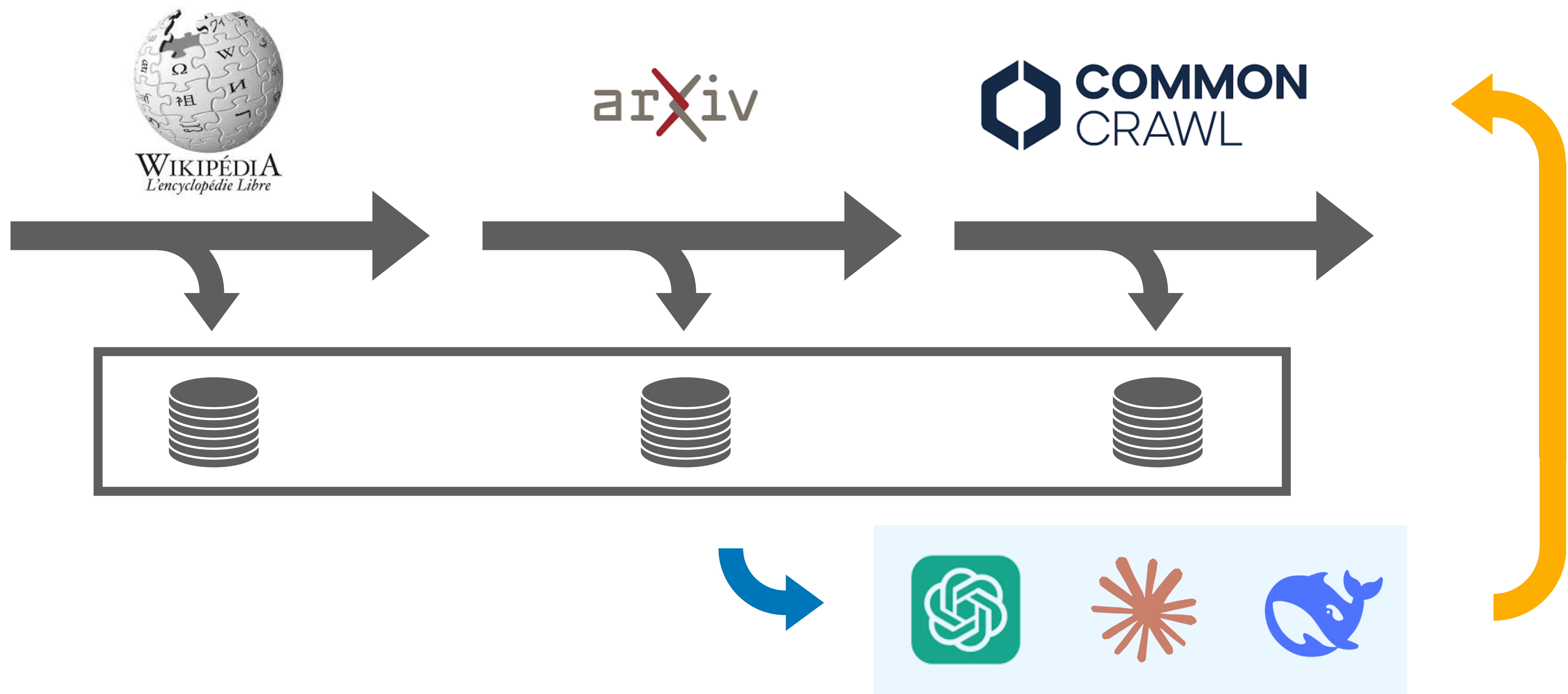
VS.

deer	bird
	
ship	frog
	
deer	truck
	

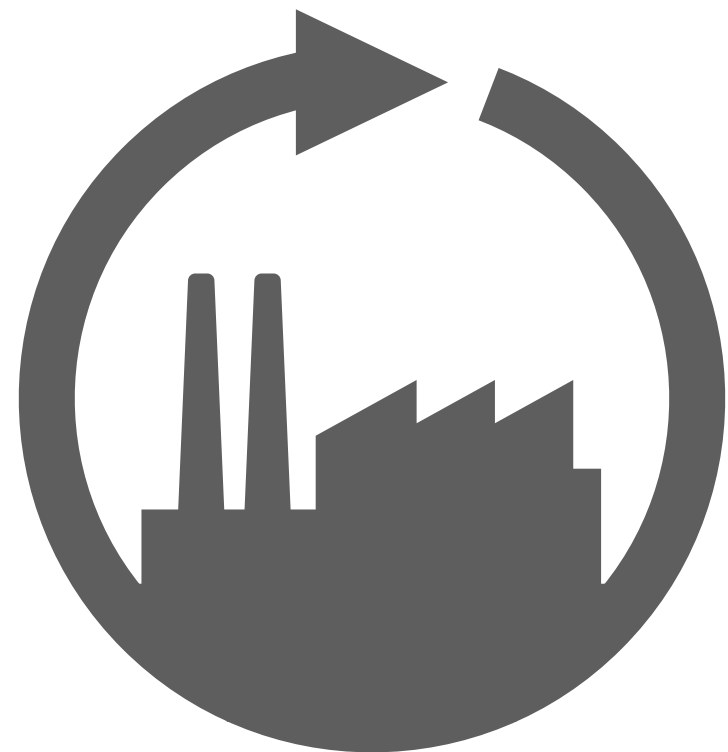
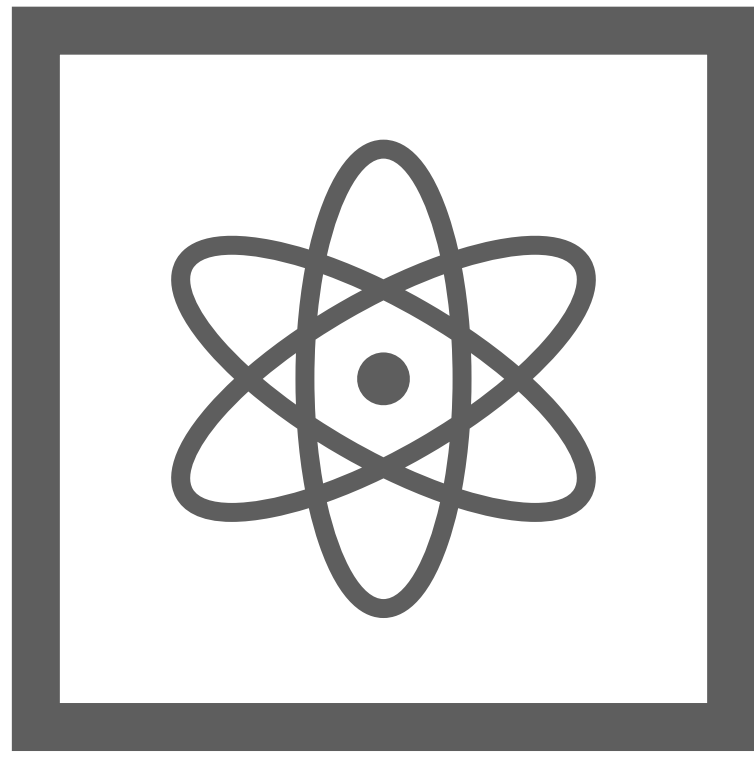
Complex Tasks



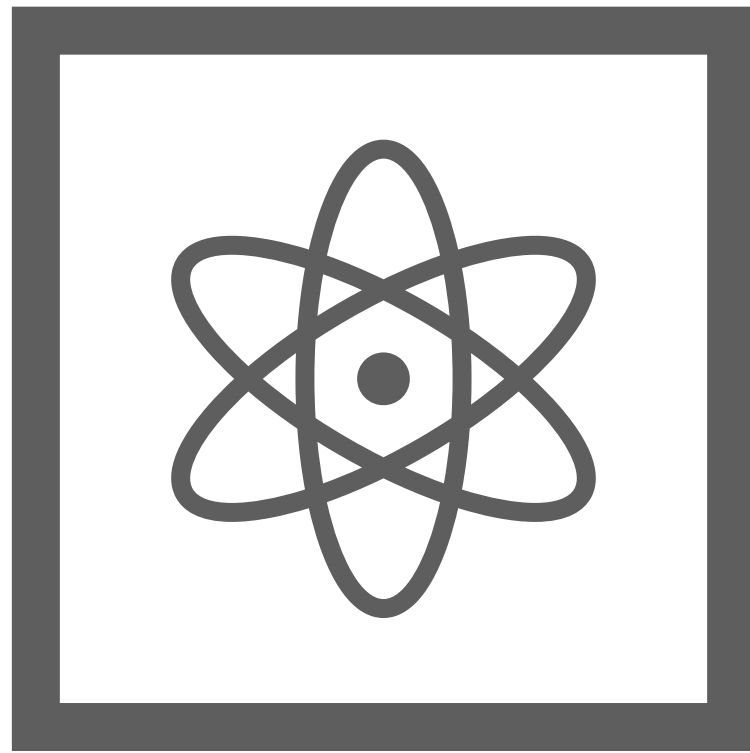
Complex Tasks



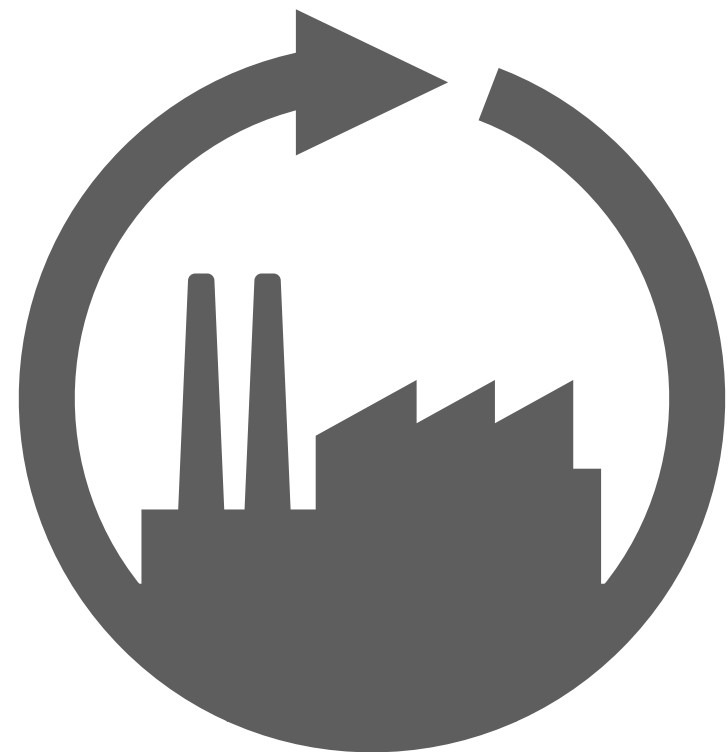
Caricatures of learning



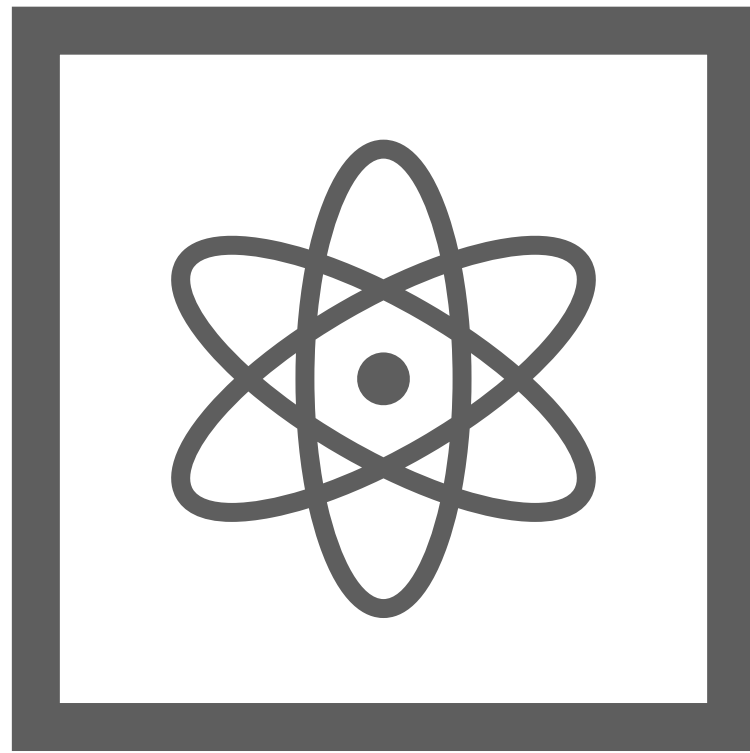
Caricatures of learning



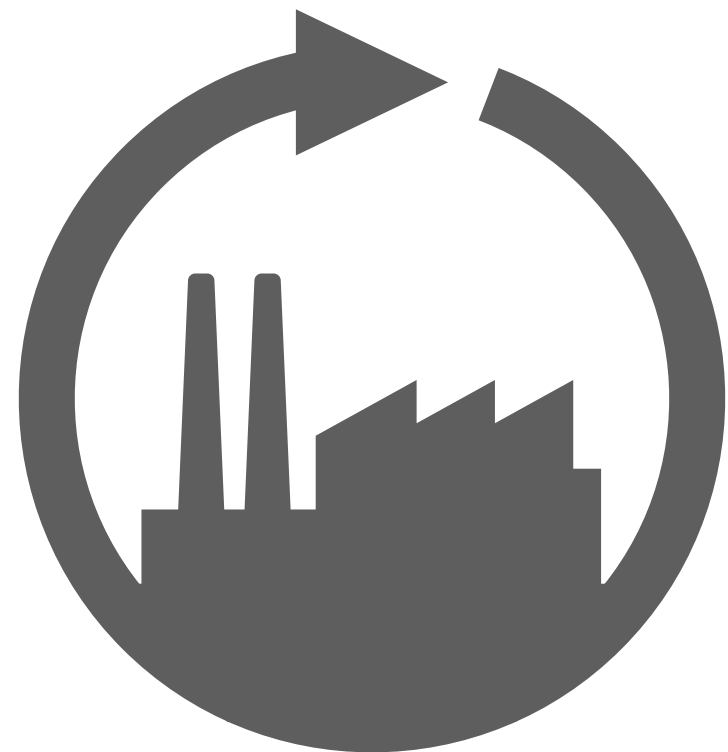
The scientific question comes first. Systems are engineered to generate targeted data that can yield accurate answers.



Caricatures of learning

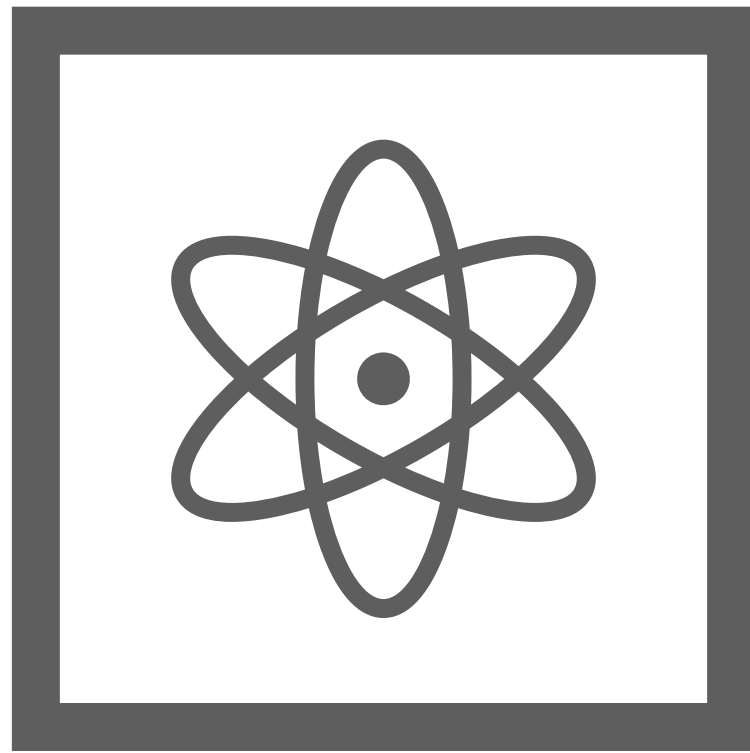


The scientific question comes first. Systems are engineered to generate targeted data that can yield accurate answers.



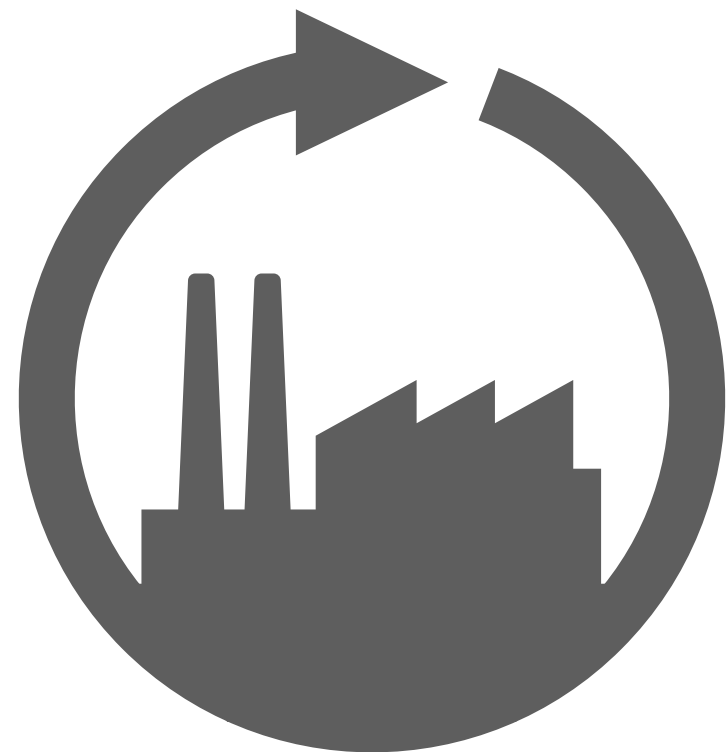
Data is a byproduct of complex systems, and not the goal.

Caricatures of learning



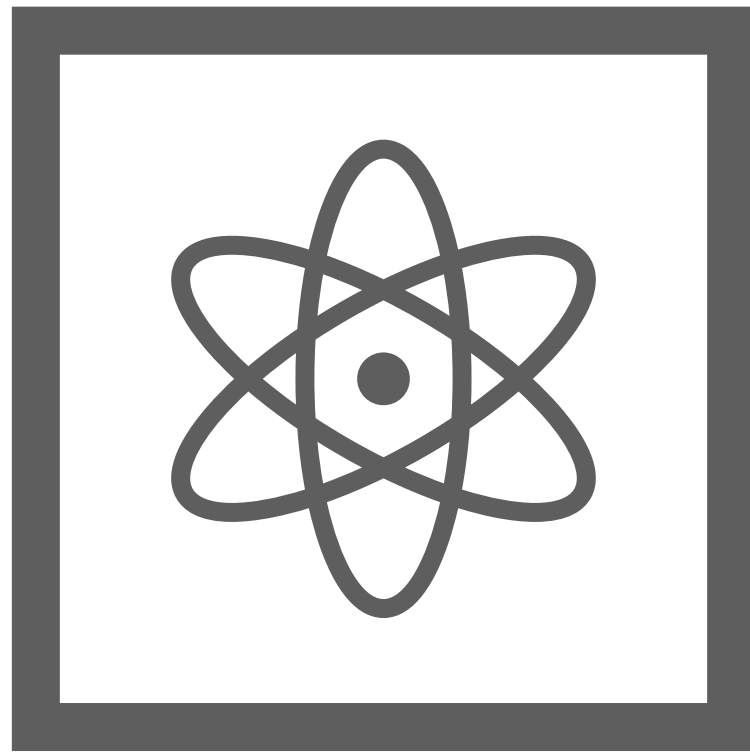
where theory is strong

The scientific question comes first. Systems are engineered to generate targeted data that can yield accurate answers.



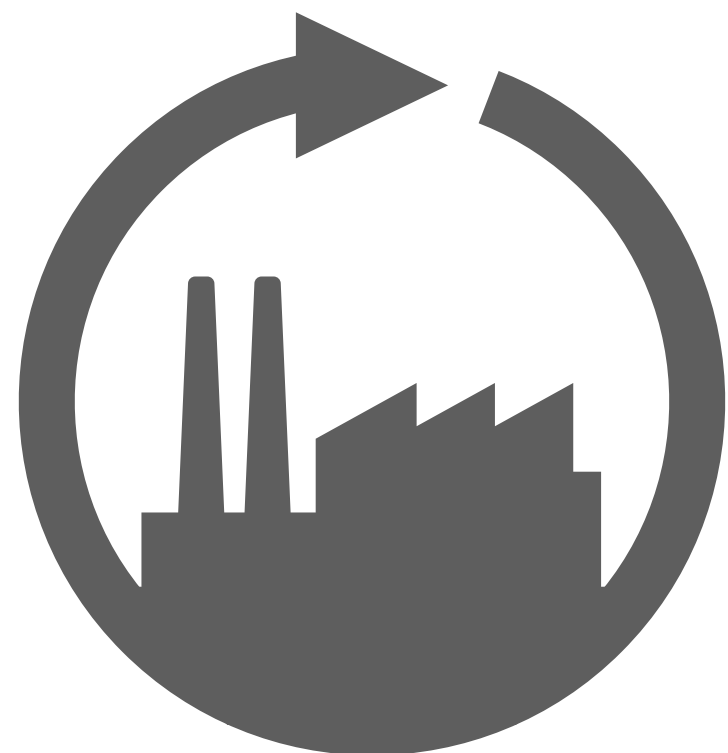
Data is a byproduct of complex systems, and not the goal.

Caricatures of learning



where theory is strong

The scientific question comes first. Systems are engineered to generate targeted data that can yield accurate answers.



where practice has taken off

Data is a byproduct of complex systems, and not the goal.

Challenges for modern machine learning



Challenges for modern machine learning



- Data-generating processes are black-boxes

Challenges for modern machine learning



- Data-generating processes are black-boxes
- The environment is strategic, and we have limited control

Challenges for modern machine learning



- Data-generating processes are black-boxes
 - The environment is strategic, and we have limited control
-
- Tasks are hard to define, multi-objective, involve human values

Challenges for modern machine learning



- Data-generating processes are black-boxes
- The environment is strategic, and we have limited control
- Tasks are hard to define, multi-objective, involve human values
- Solving complex tasks requires a division of labor

Challenges for modern machine learning



- Data-generating processes are black-boxes
- The environment is strategic, and we have limited control
- Tasks are hard to define, multi-objective, involve human values
- Solving complex tasks requires a division of labor
 - The “learner” is a (coordinated?) network of adapting agents

Broad orientation of my research



- **ML for settings with “limited control”** is driven by practice...

Broad orientation of my research



- **ML for settings with “limited control”** is driven by practice, generating a lot of hard but interesting **basic science questions**.

Broad orientation of my research



- **ML for settings with “limited control”** is driven by practice, generating a lot of hard but interesting **basic science questions**.
- What are the **right abstractions** to get at the interdependent relationship between **machine and human/social processes**?

My Research



My Research



- **Online consistency of the nearest neighbor rule**
with Sanjoy Dasgupta (NeurIPS 2024)
- **Consistency of the k_n -nearest neighbor rule under adaptive sampling**
with Robi Bhattacharjee and Sanjoy Dasgupta (NeurIPS 2025)
- **Actively learning halfspaces with without synthetic data**
with Hadley Black, Kasper Green Larsen, Arya Mazumdar, Barna Saha (COLT 2026)
- **Convergence of online k -means**
with Sanjoy Dasgupta and Gaurav Mahajan (AISTATS 2022)

My Research



- How can we learn from adaptive or sequential data?

My Research



- **Learnable mixed Nash equilibria are collectively rational**
with Yian Ma (preprint)

My Research



- What are the economic / dynamical behaviors of multi-agent systems?

My Research



- **Preference optimization on the Pareto set**
with Abhishek Roy and Yian Ma (NeurIPS 2025)
- **On the semi-supervised sample complexity of multi-objective learning**
with Tobias Wegel, Junhyung Park, and Fanny Yang (NeurIPS 2025)
- **Hedging on the frontier: learning new tasks with few samples**
with Tobias Wegel, Federico Di Gennaro, and Fanny Yang (ICML 2026)
- **Metric learning from limited pairwise preference comparisons**
with Zhi Wang and Ramya Korlakai Vinayak (UAI 2024)

My Research



- How should we break down a complex learning task into simpler ones?
- How do we incorporate human preferences into machine systems?

Roadmap

A. Motivation

B. Research Overview

i. Data

ii. Learner

iii. Task

RESEARCH OVERVIEW: DATA

A Model of Adaptive and Imprecise Data

Some dimensions of what makes learning hard

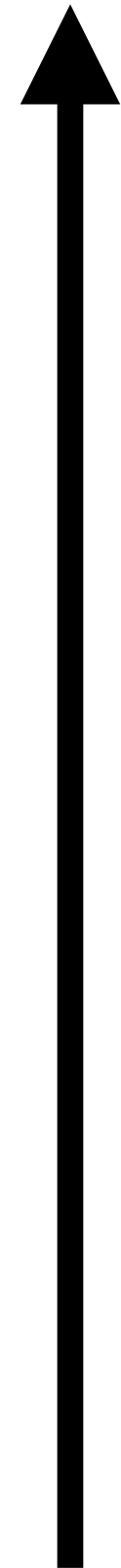
Some dimensions of what makes learning hard

**Model
Complexity**



Some dimensions of what makes learning hard

**Model
Complexity**



How complicated is the thing we want to learn?

Some dimensions of what makes learning hard

**Model
Complexity**



How complicated is the thing we want to learn?

- Bias of a coin (1 parameter)

Some dimensions of what makes learning hard

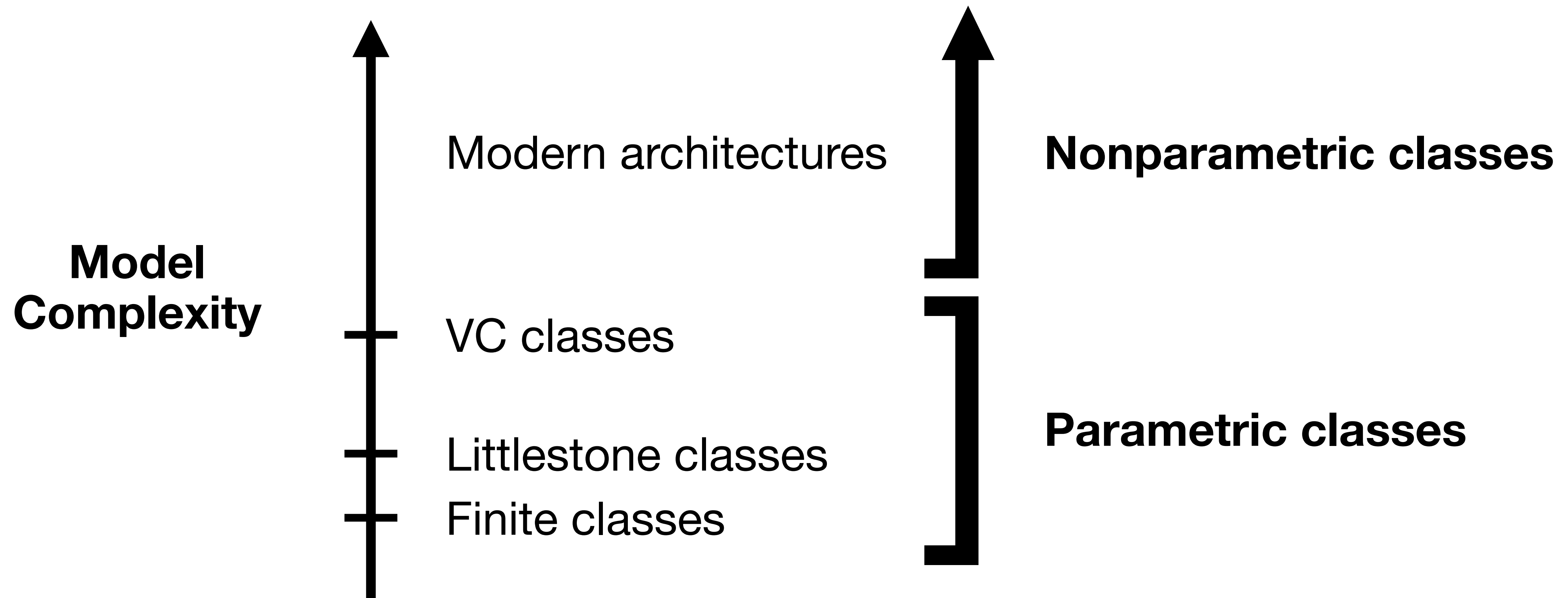
**Model
Complexity**



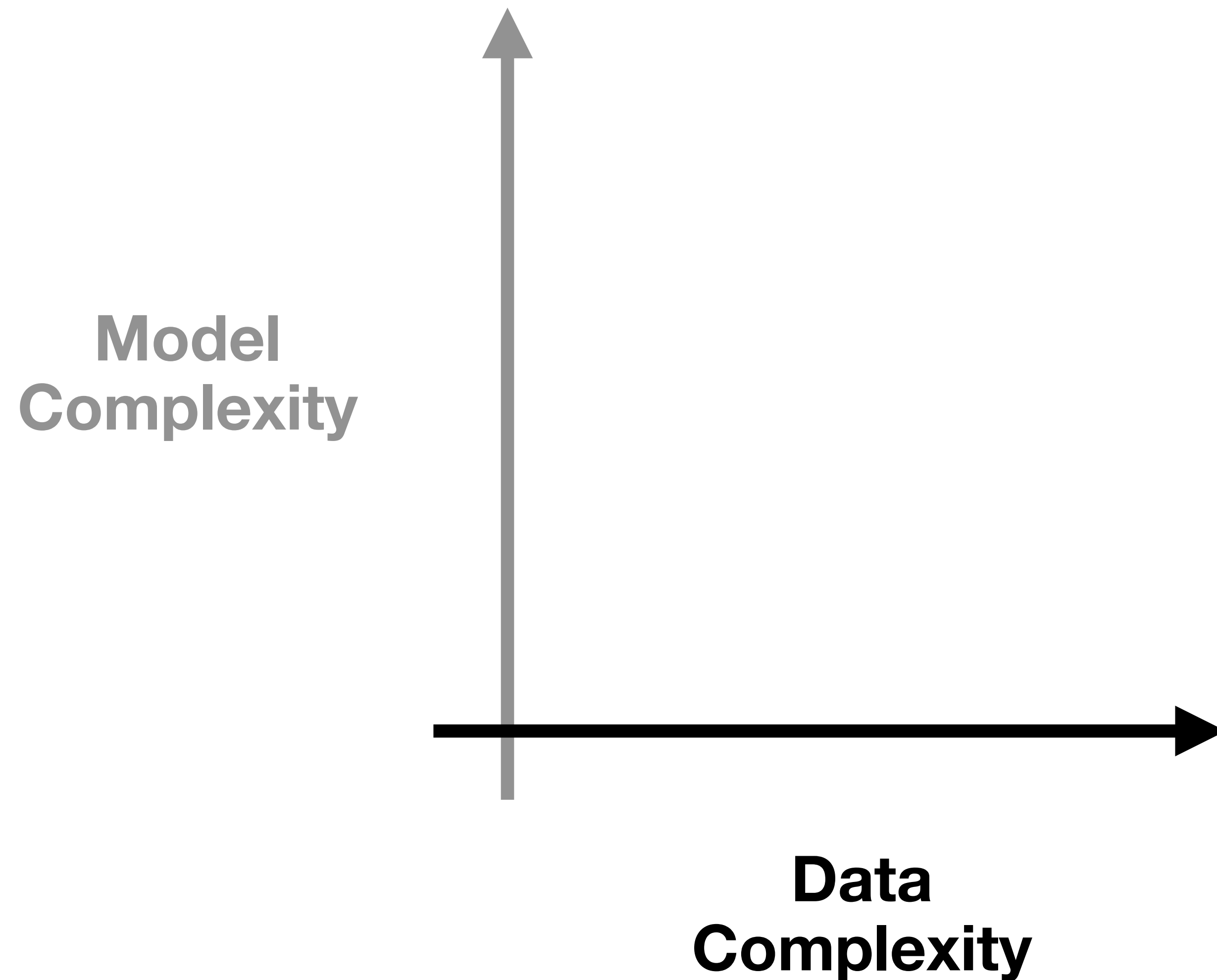
How complicated is the thing we want to learn?

- Bias of a coin (1 parameter)
- A large language model (1e12 parameters)

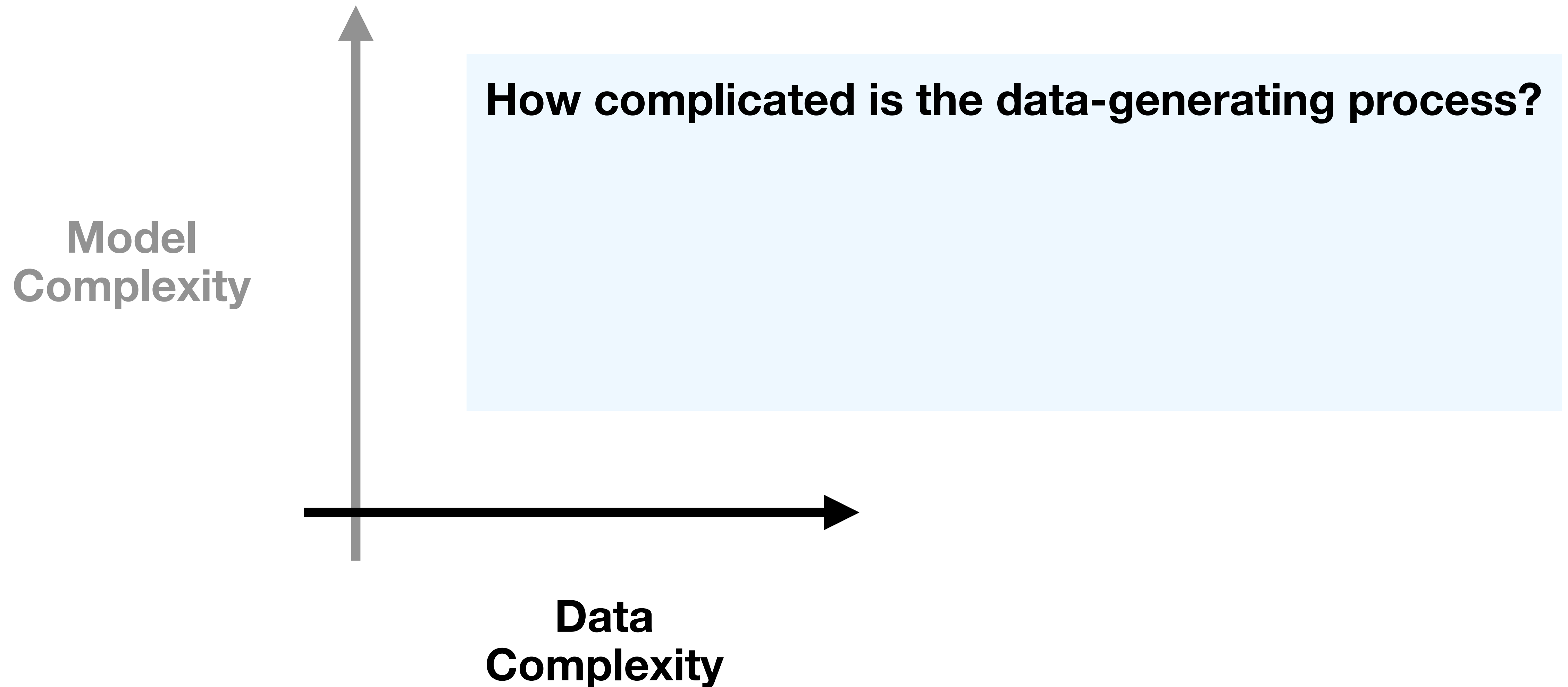
Some dimensions of what makes learning hard



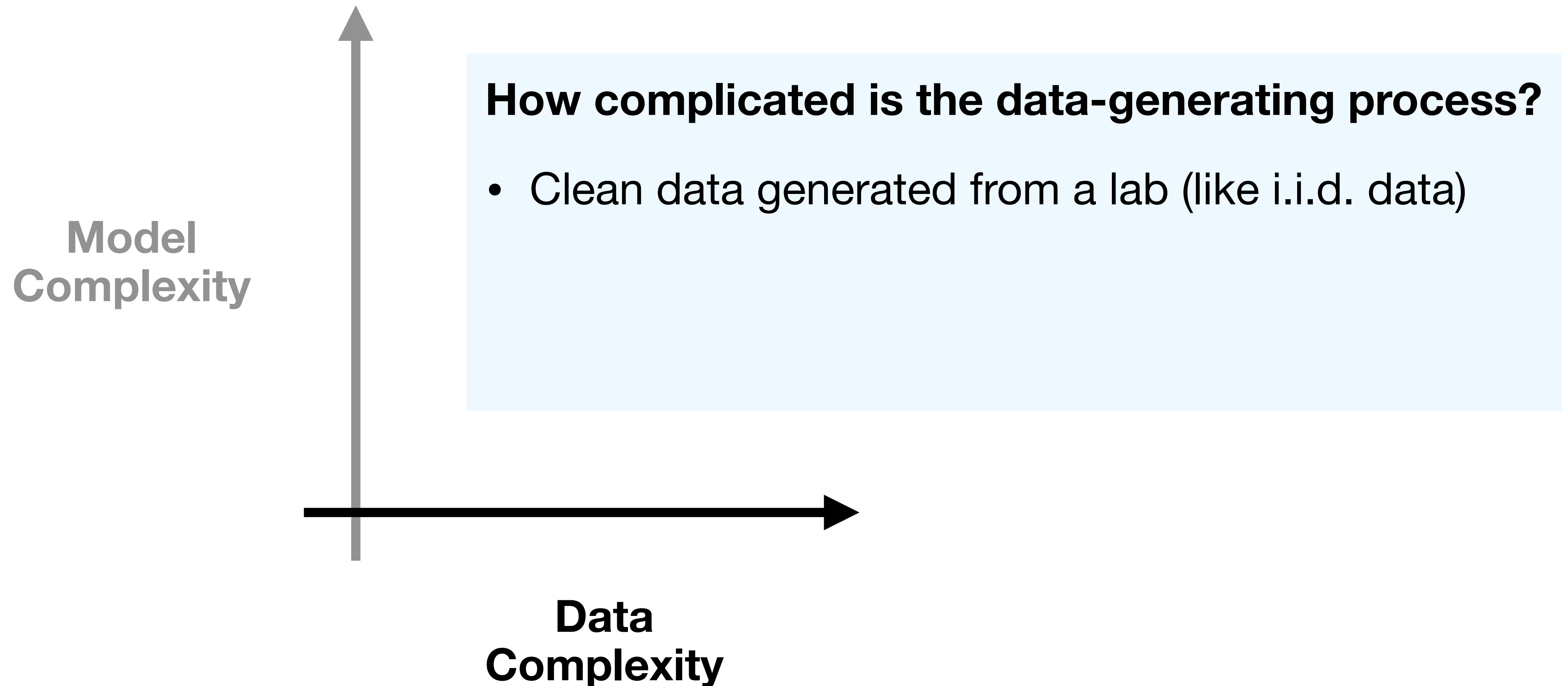
Some dimensions of what makes learning hard



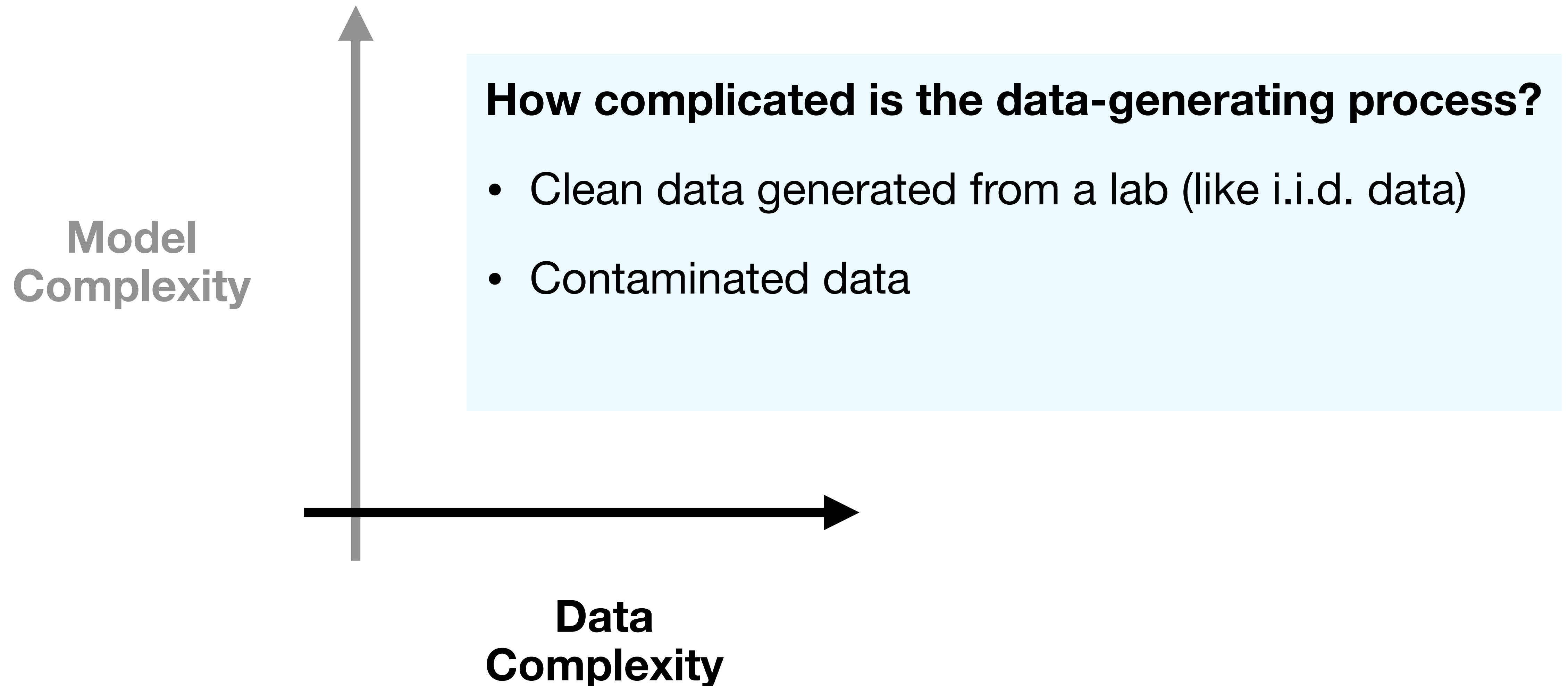
Some dimensions of what makes learning hard



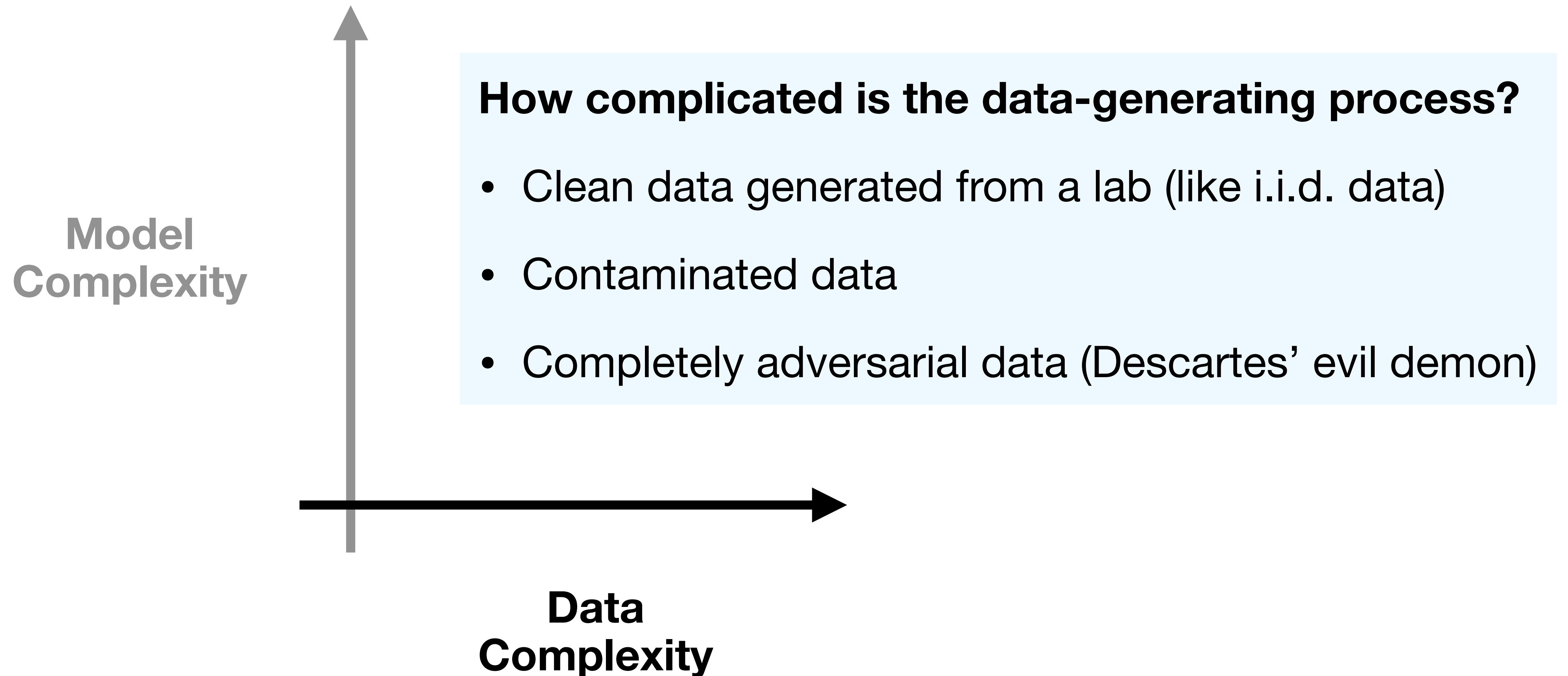
Some dimensions of what makes learning hard



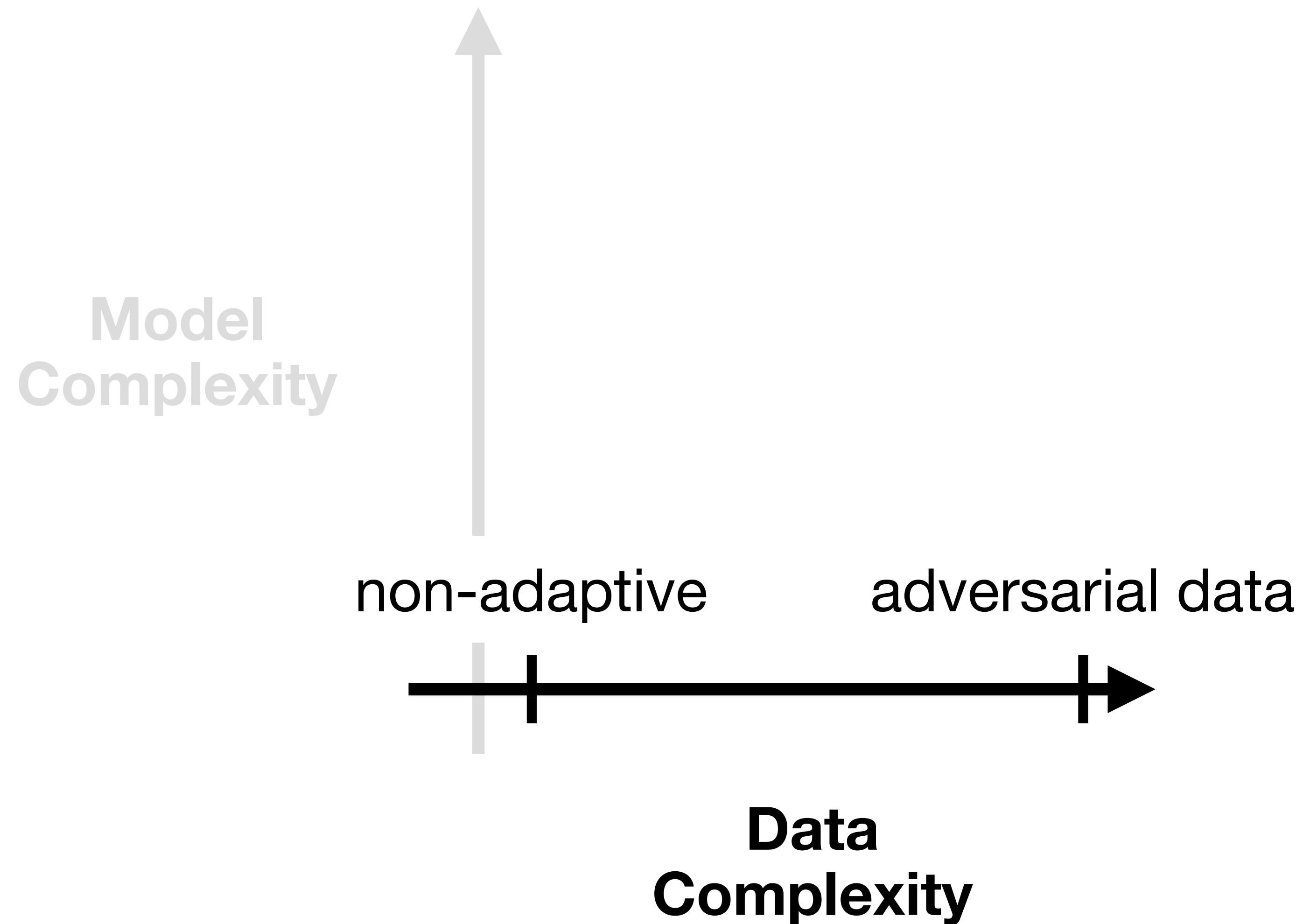
Some dimensions of what makes learning hard



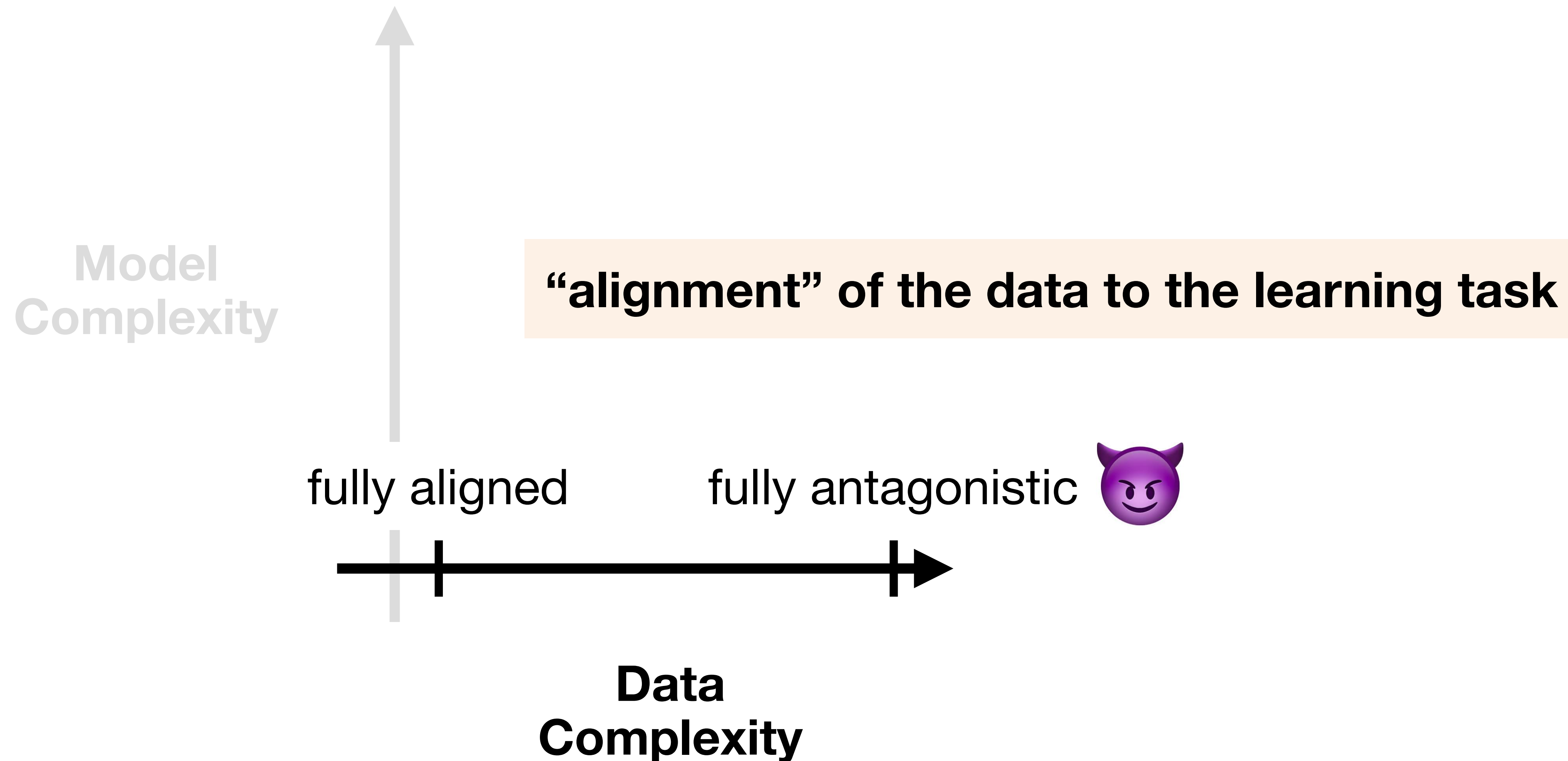
Some dimensions of what makes learning hard



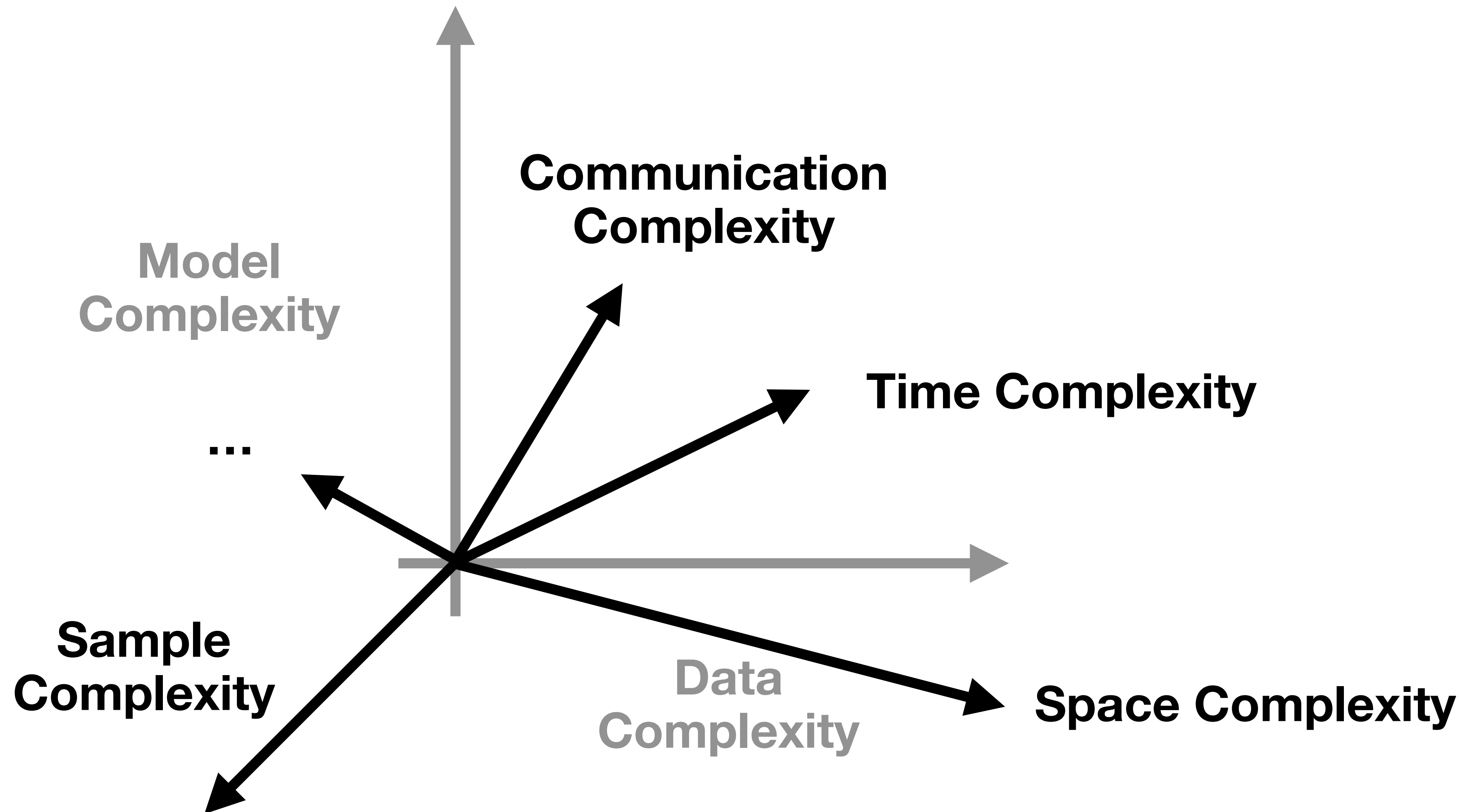
Some dimensions of what makes learning hard



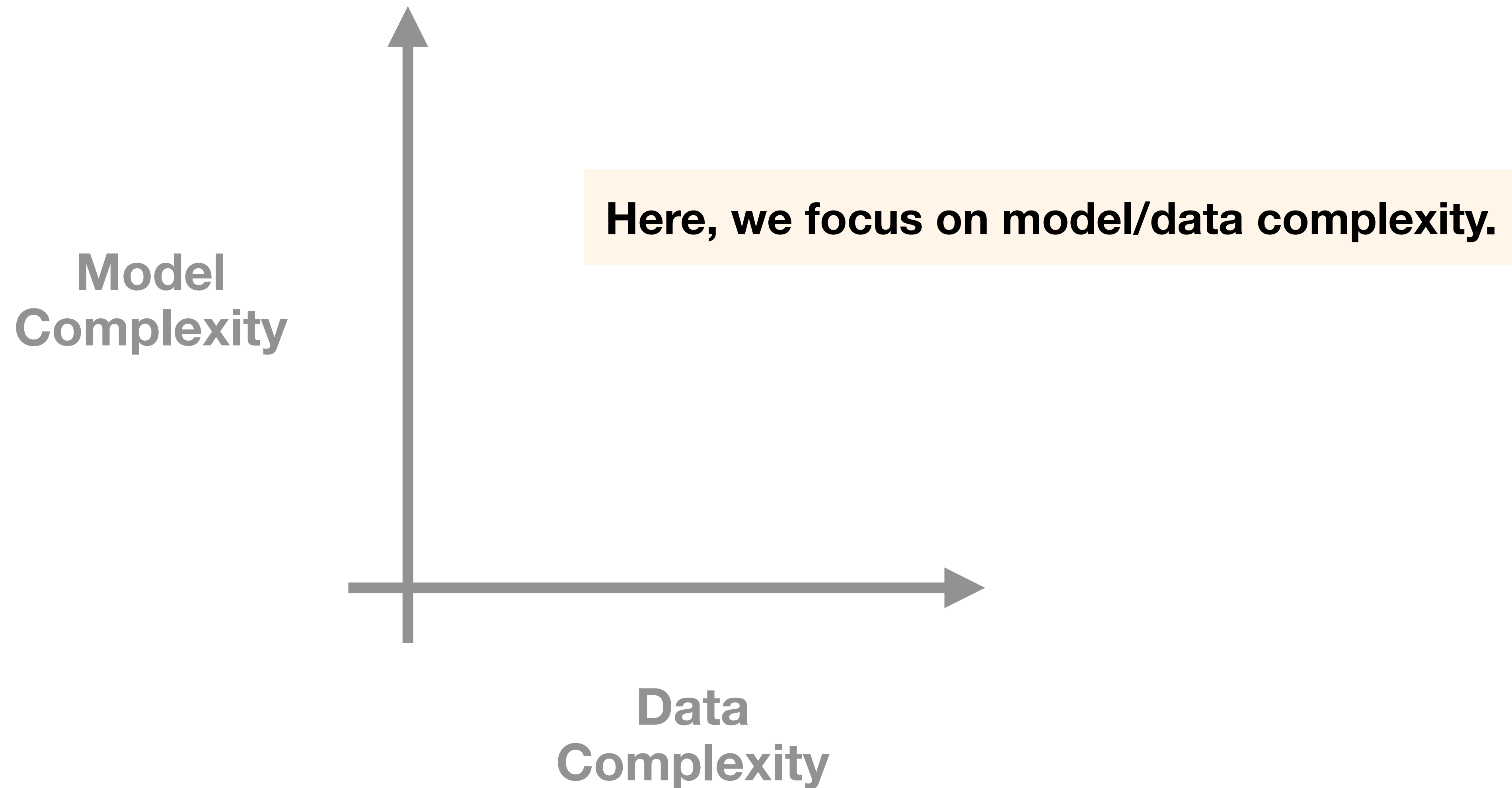
Some dimensions of what makes learning hard



Some dimensions of what makes learning hard



Some dimensions of what makes learning hard



The lay of the land

Model
Complexity



Data
Complexity

The lay of the land

Model
Complexity

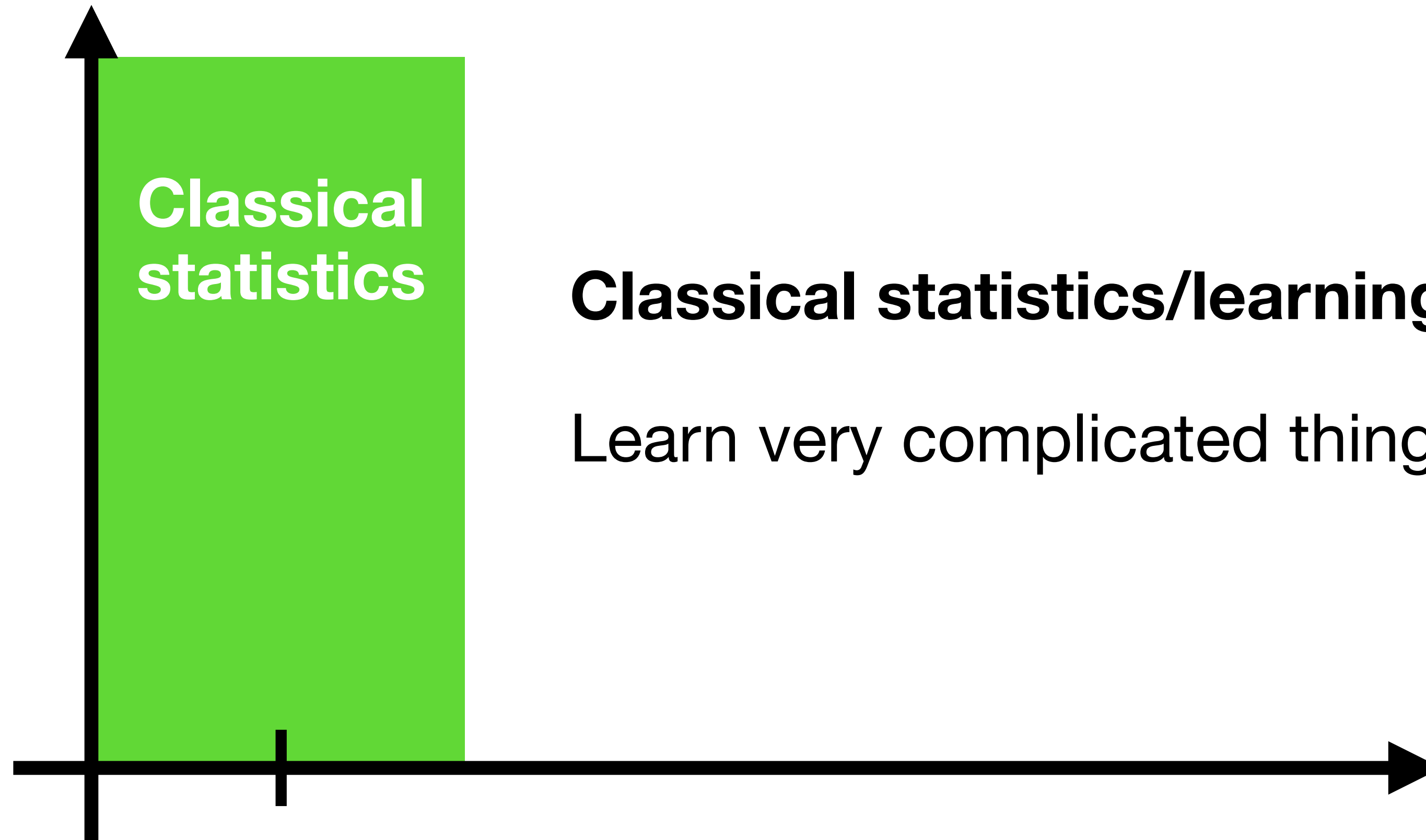
**Classical
statistics**

Classical statistics/learning theory:

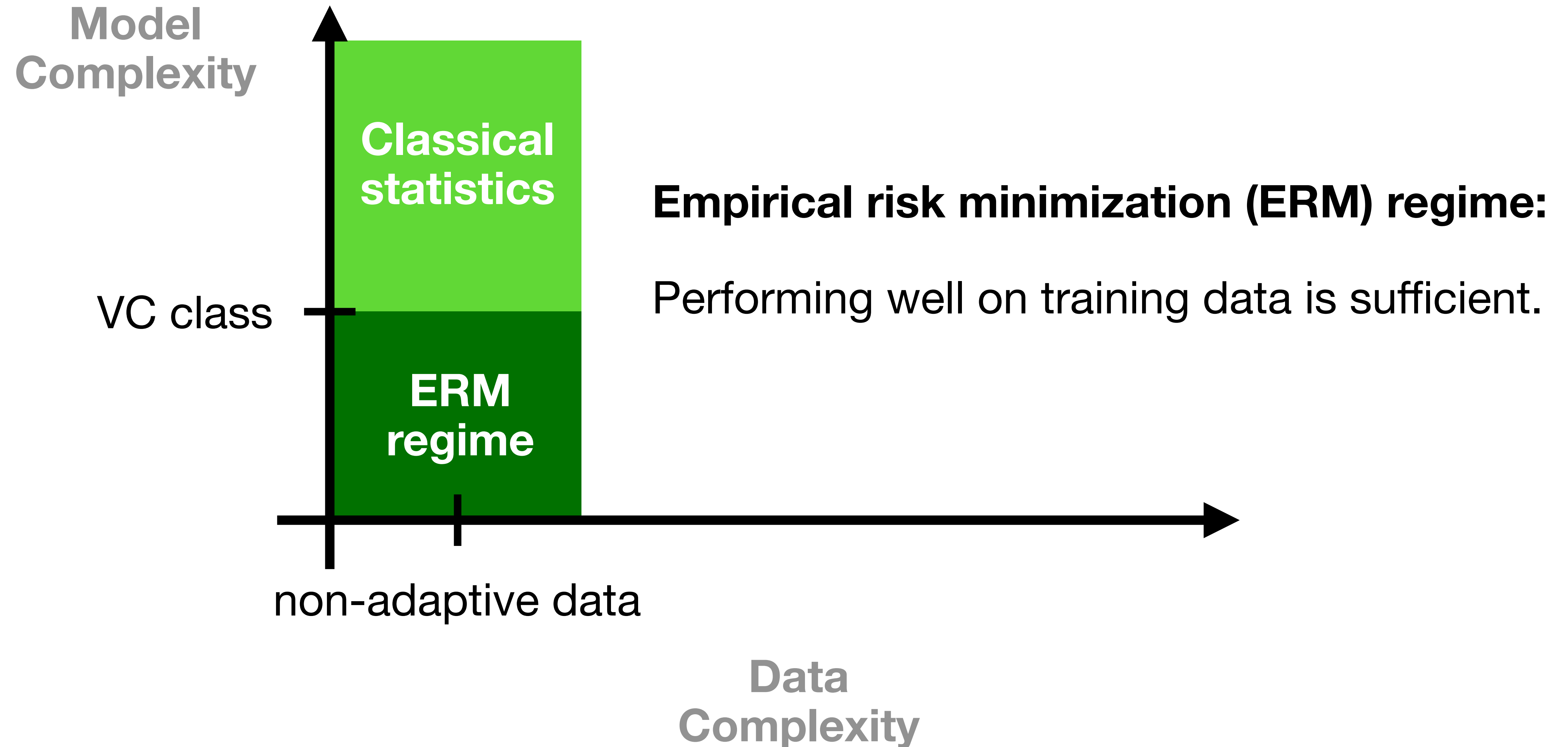
Learn very complicated things from nice data.

non-adaptive data

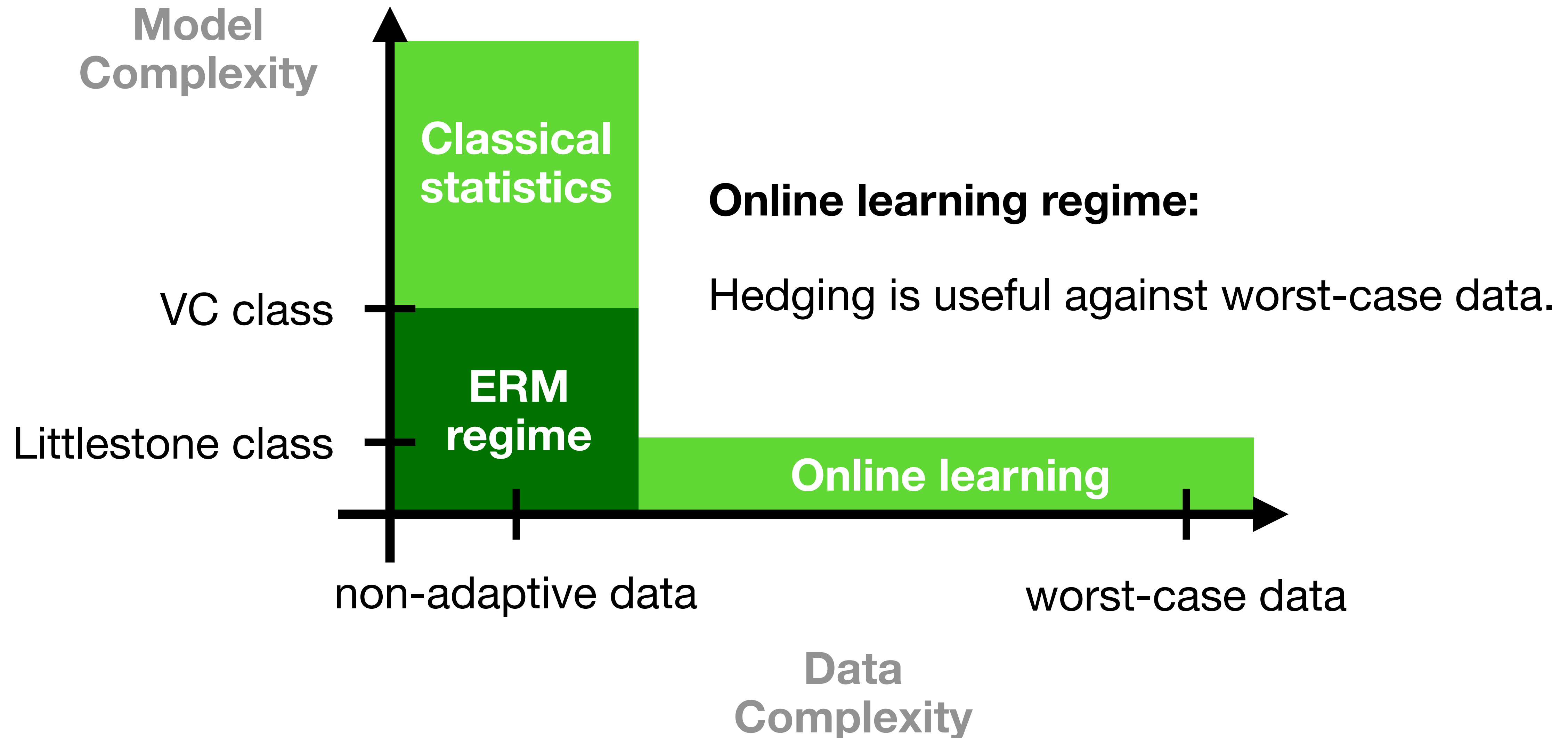
Data
Complexity



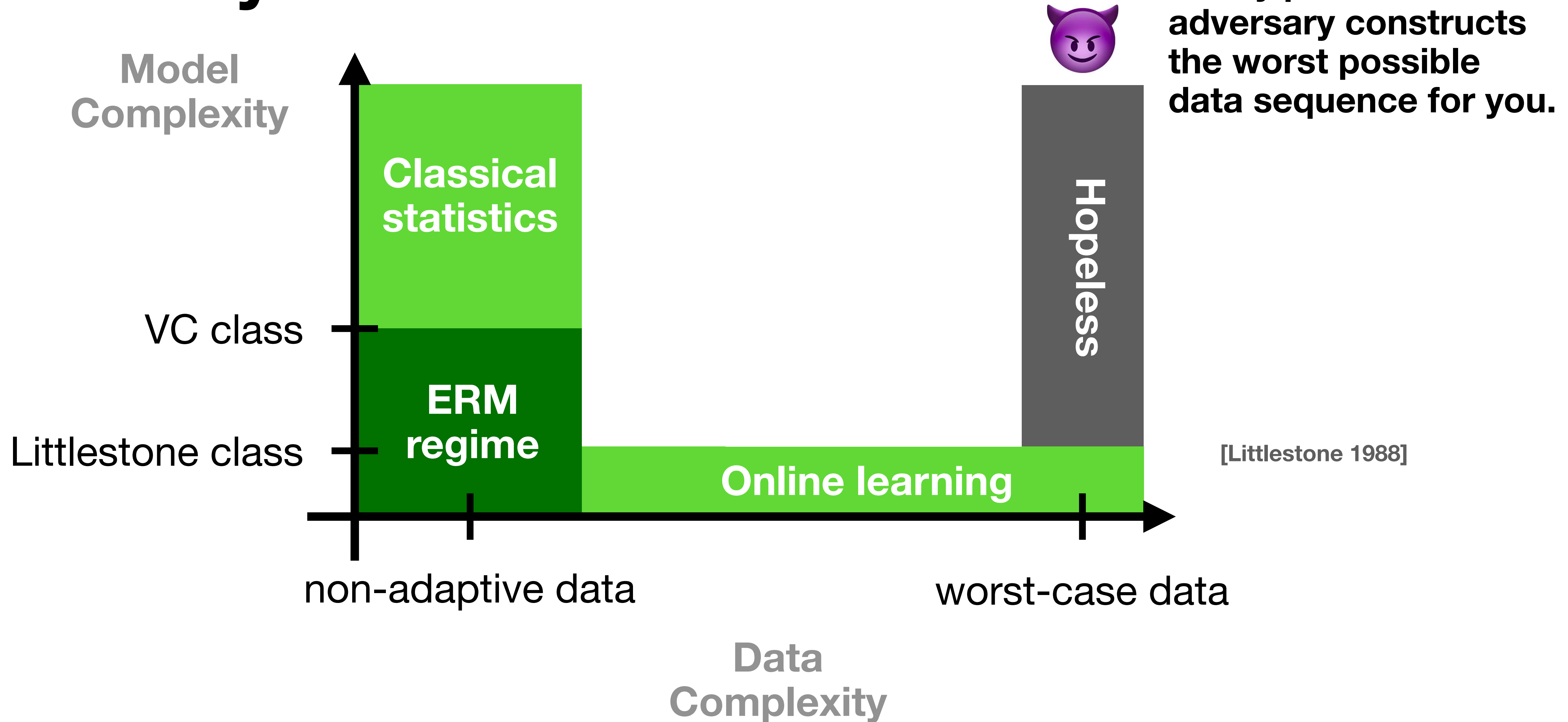
The lay of the land



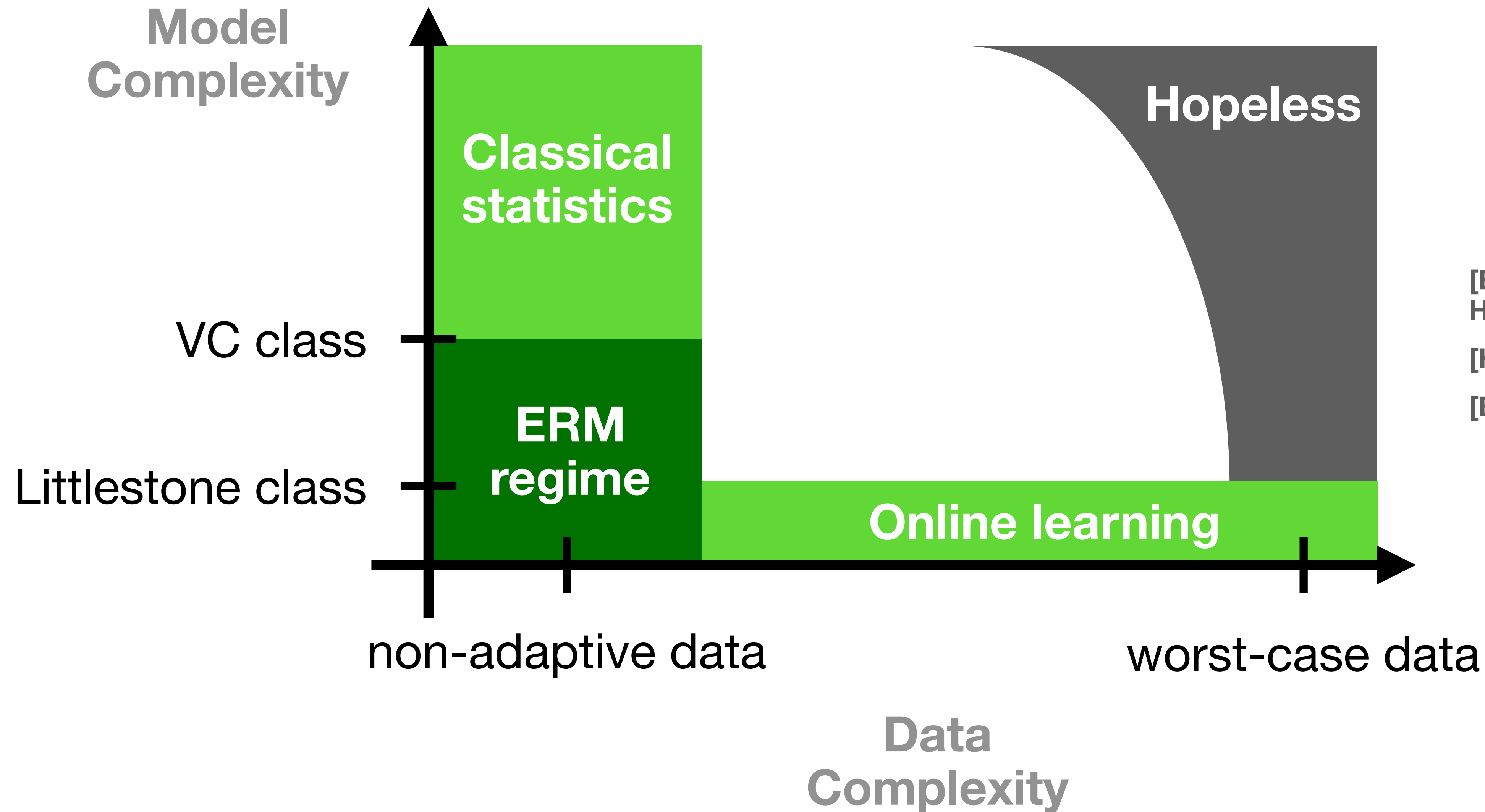
The lay of the land



The lay of the land



The lay of the land

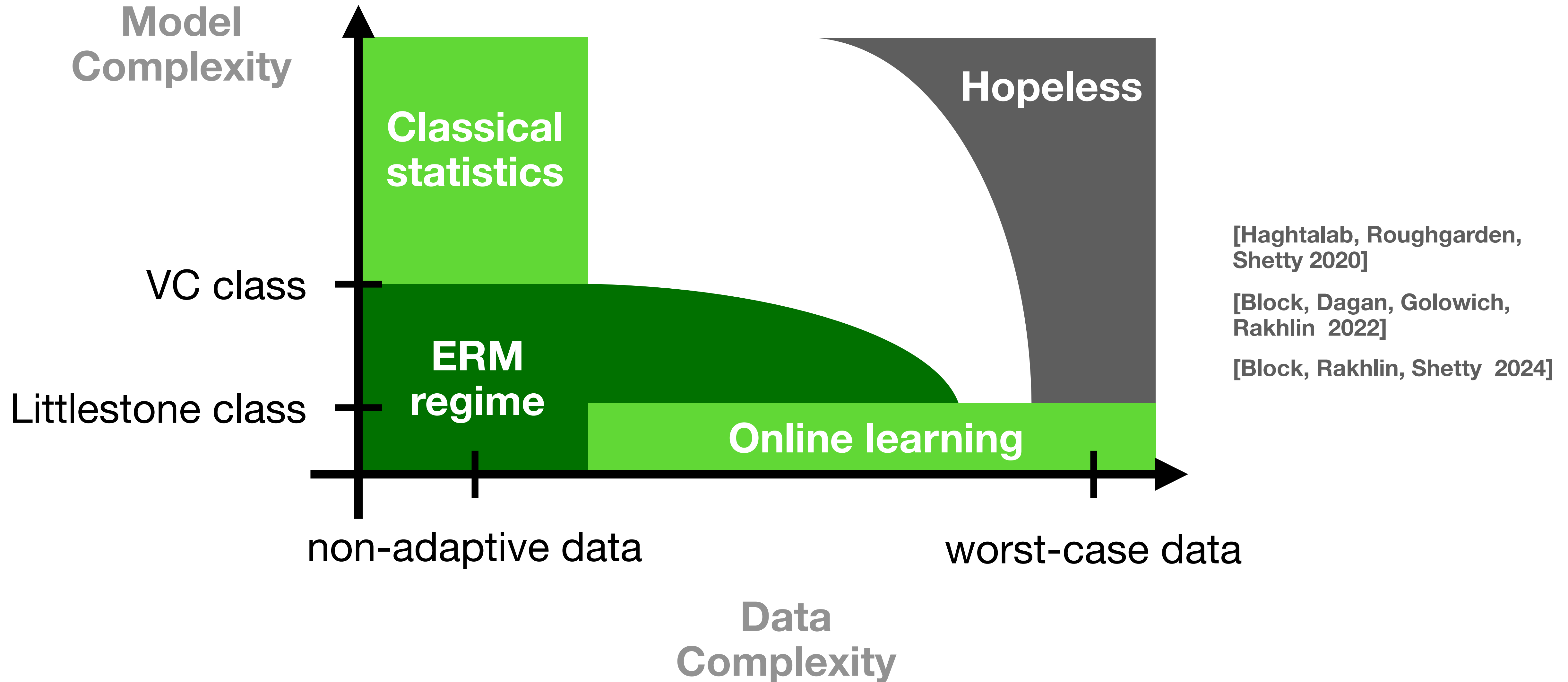


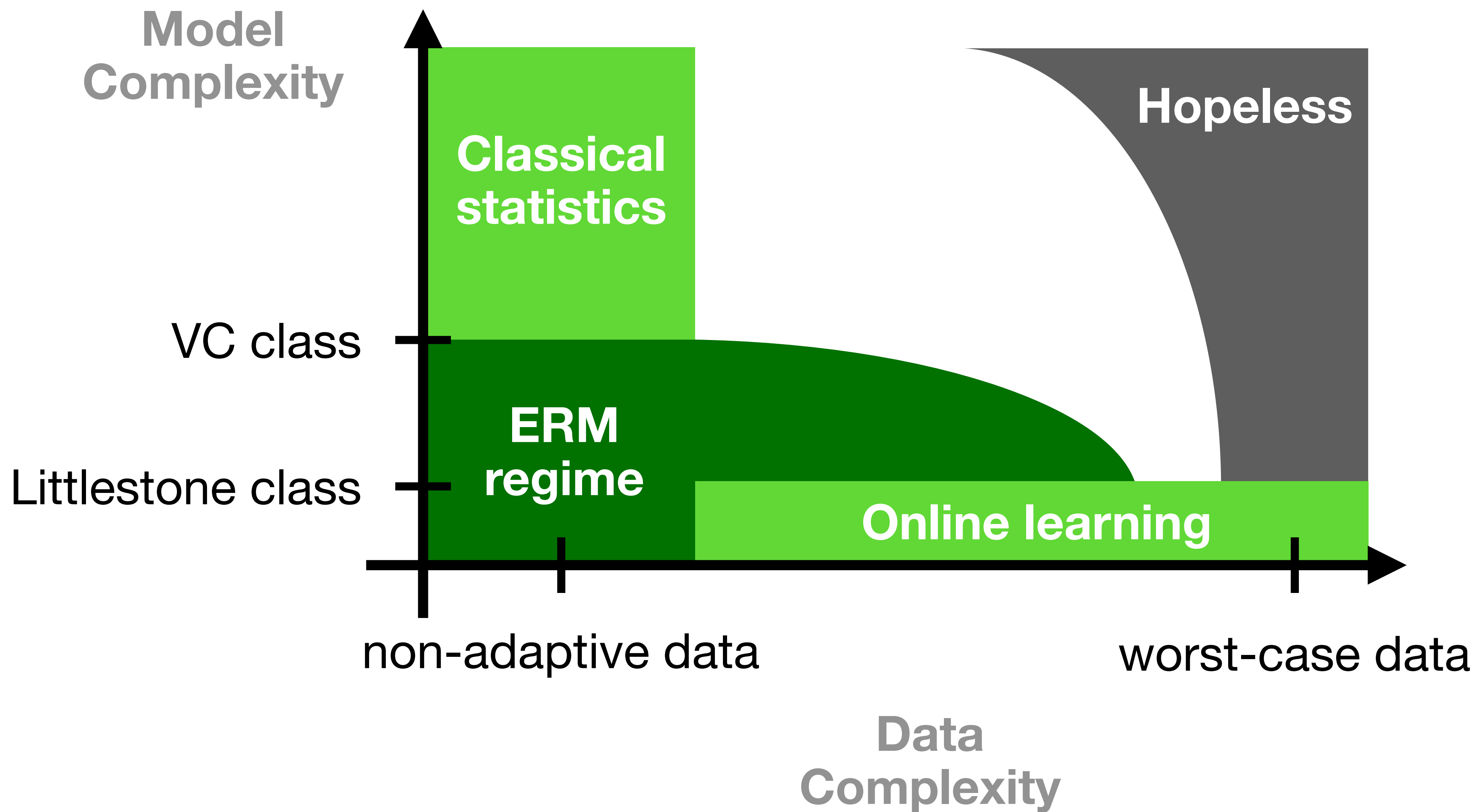
[Bousquet, Hanneke, Moran, Handel, Yehudayoff 2020]

[Hanneke 2021]

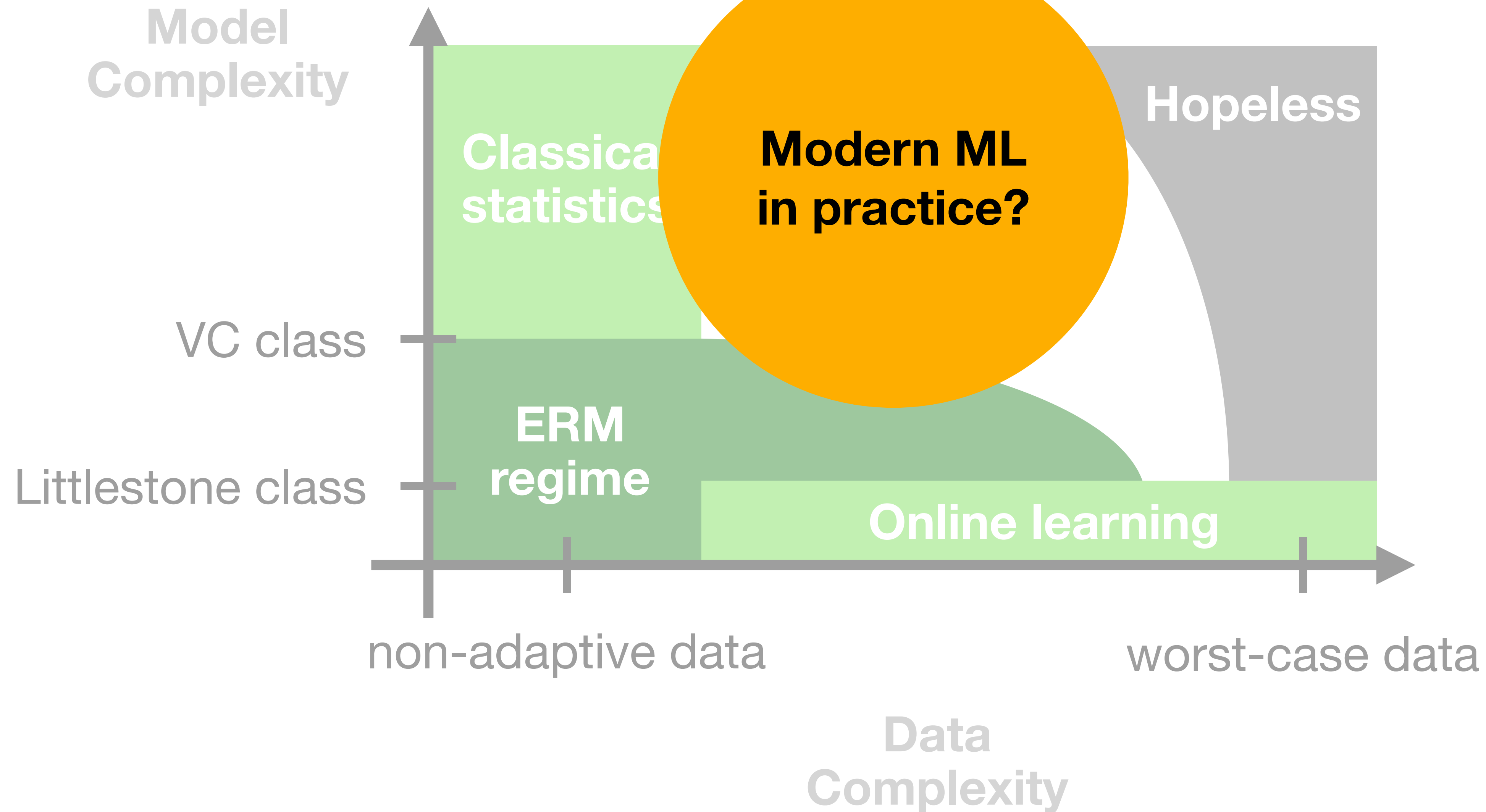
[Blanchard & Jaillet 2023]

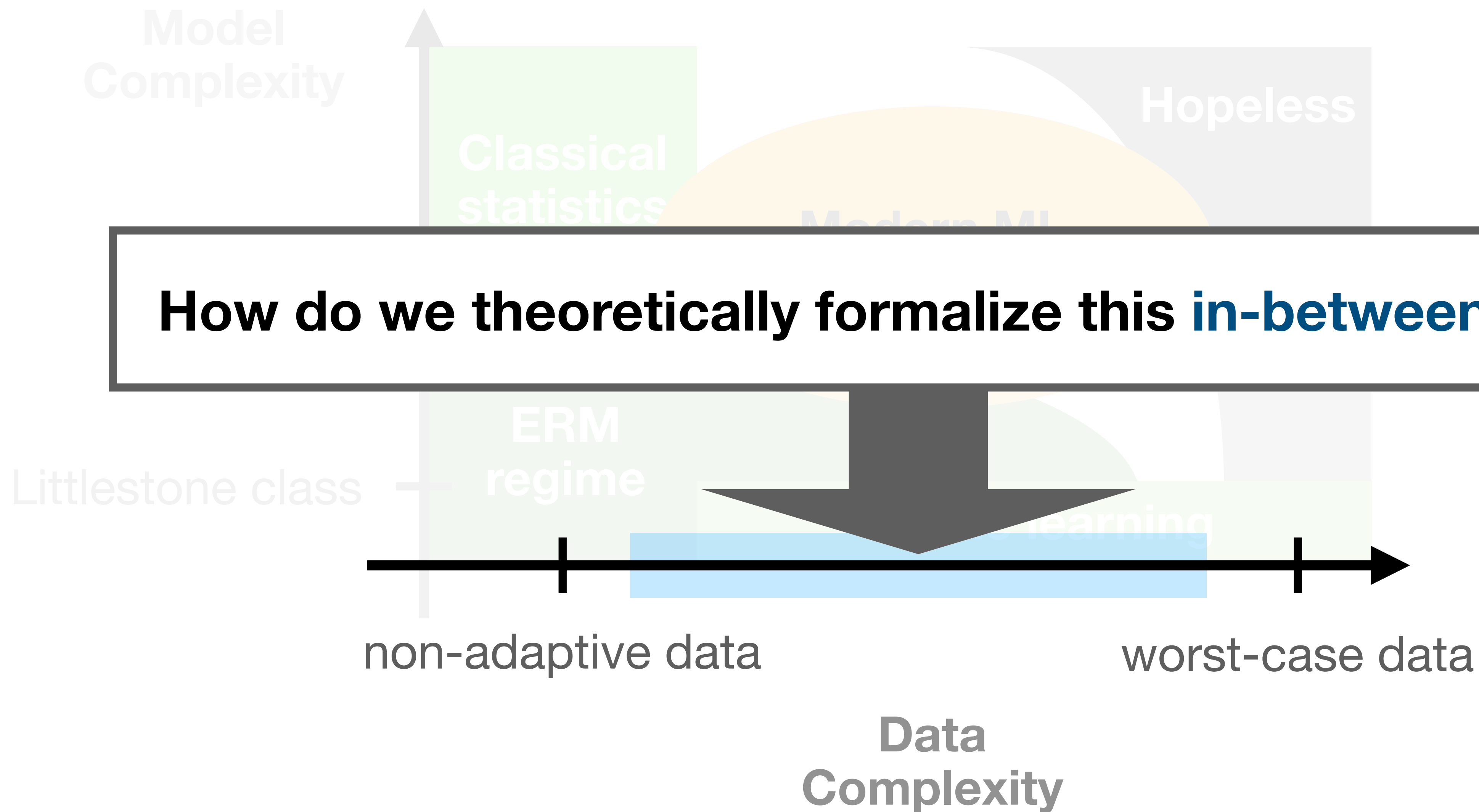
The lay of the land





The lay of the land







1. Online consistency of the nearest neighbor rule

with Sanjoy Dasgupta (NeurIPS 2024)

2. Consistency of the k_n -nearest neighbor rule under adaptive sampling

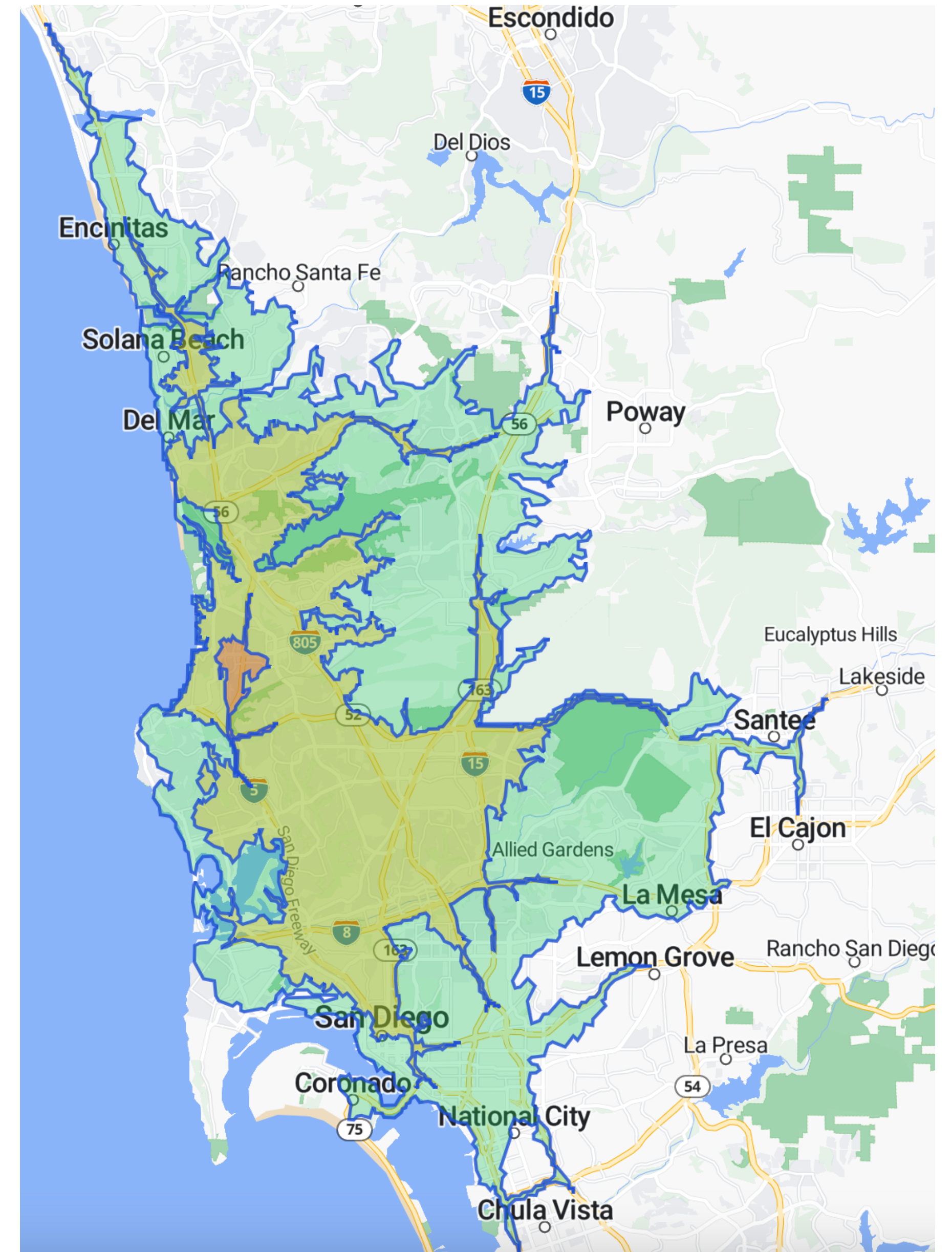
with Robi Bhattacharjee and Sanjoy Dasgupta (NeurIPS 2025)

Danger of adaptive data

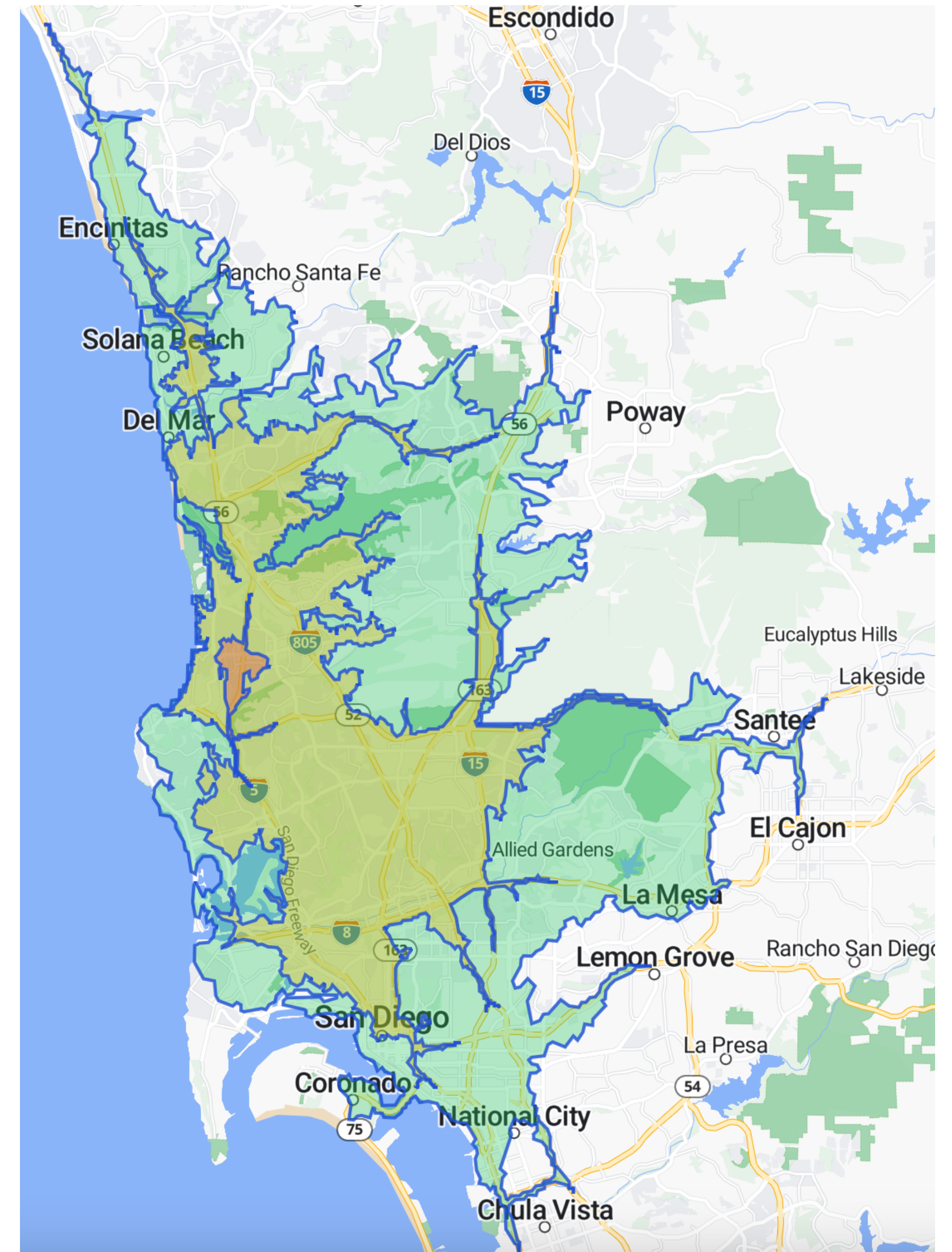
or, what makes adaptive data hard theoretically?

Example problem

Starting from UC San Diego, what can we reach by car within 20 minutes?

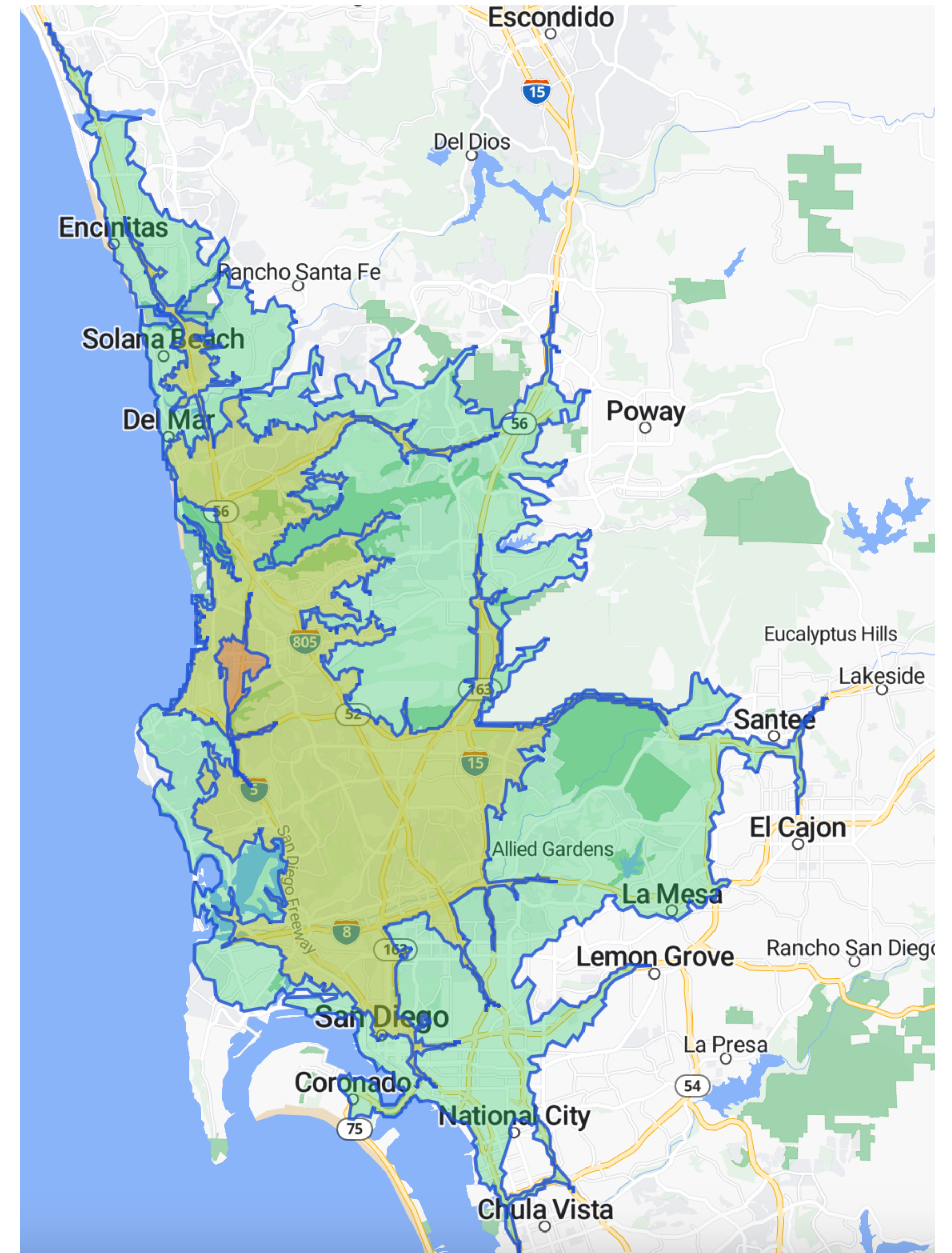


A mathematical model



A mathematical model

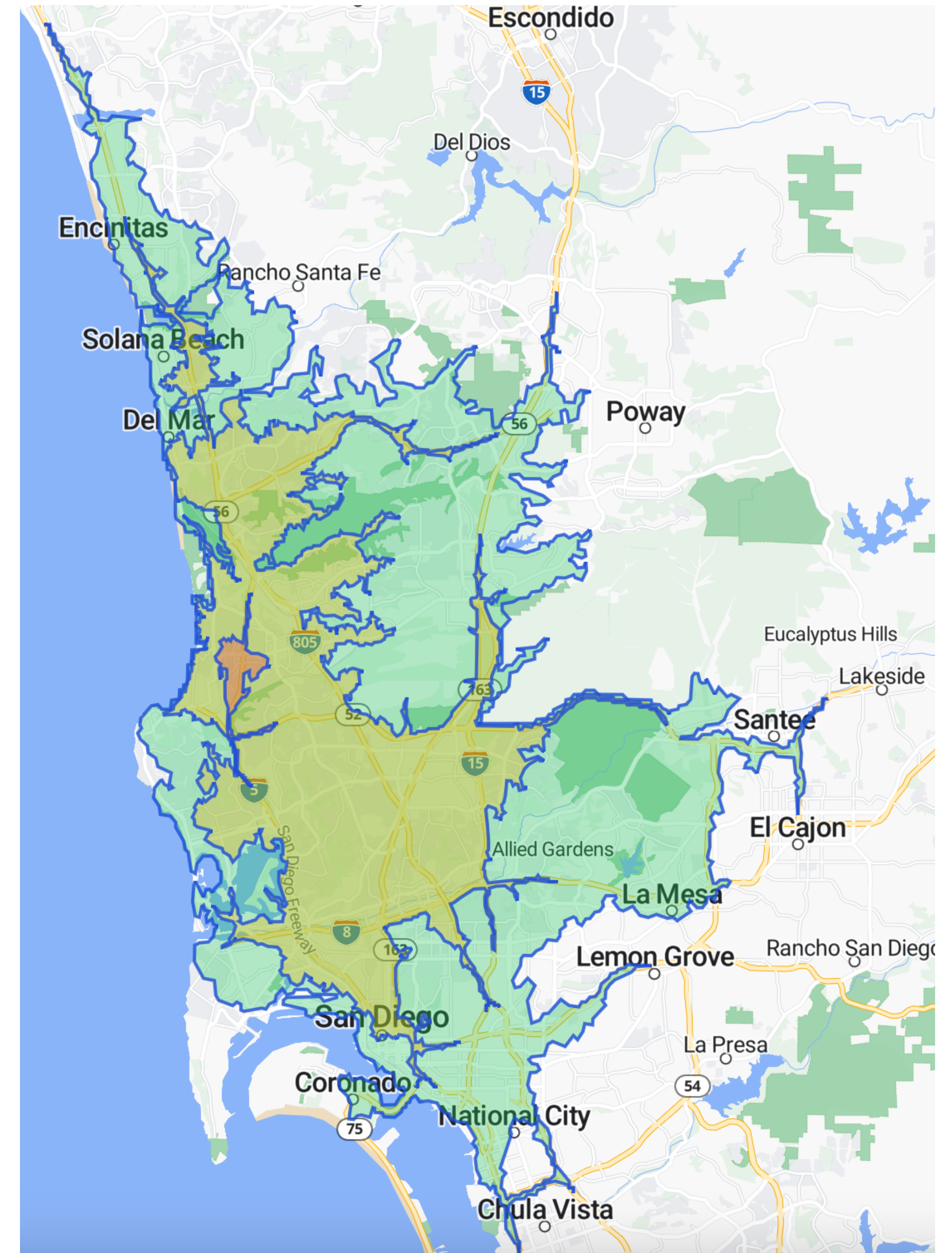
- Call the set of all possible destinations $\mathcal{X}$



A mathematical model

- Call the set of all possible destinations $\mathcal{X}$
- Define the function $\eta : \mathcal{X} \rightarrow [0,1]$

$$\eta(x) = \Pr \left(\text{reach } x \text{ within 20 min} \right)$$

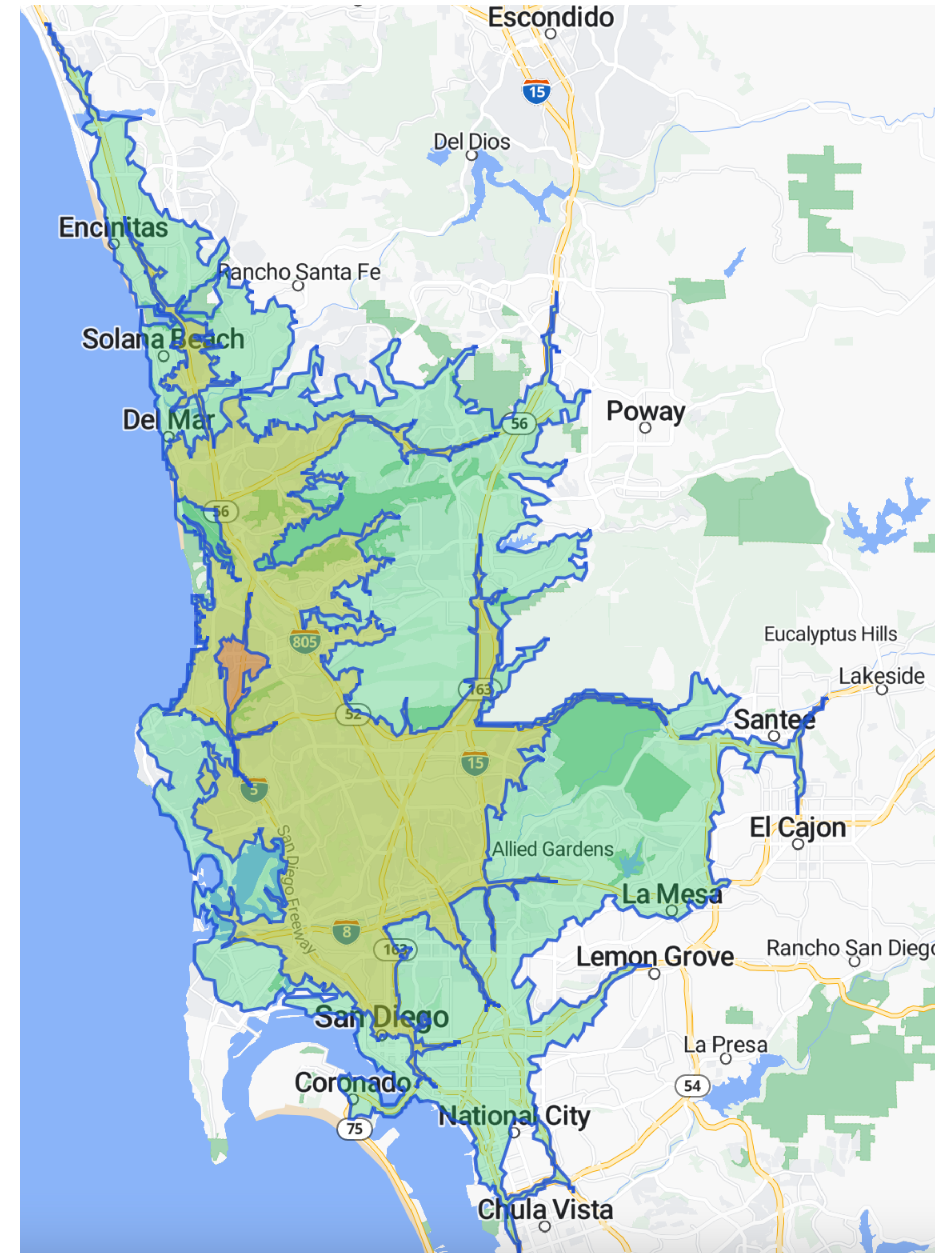


A mathematical model

- Call the set of all possible destinations $\mathcal{X}$
- Define the function $\eta : \mathcal{X} \rightarrow [0,1]$

$$\eta(x) = \Pr \left(\text{reach } x \text{ within 20 min} \right)$$

- We get a bunch of data points (X_n, Y_n)

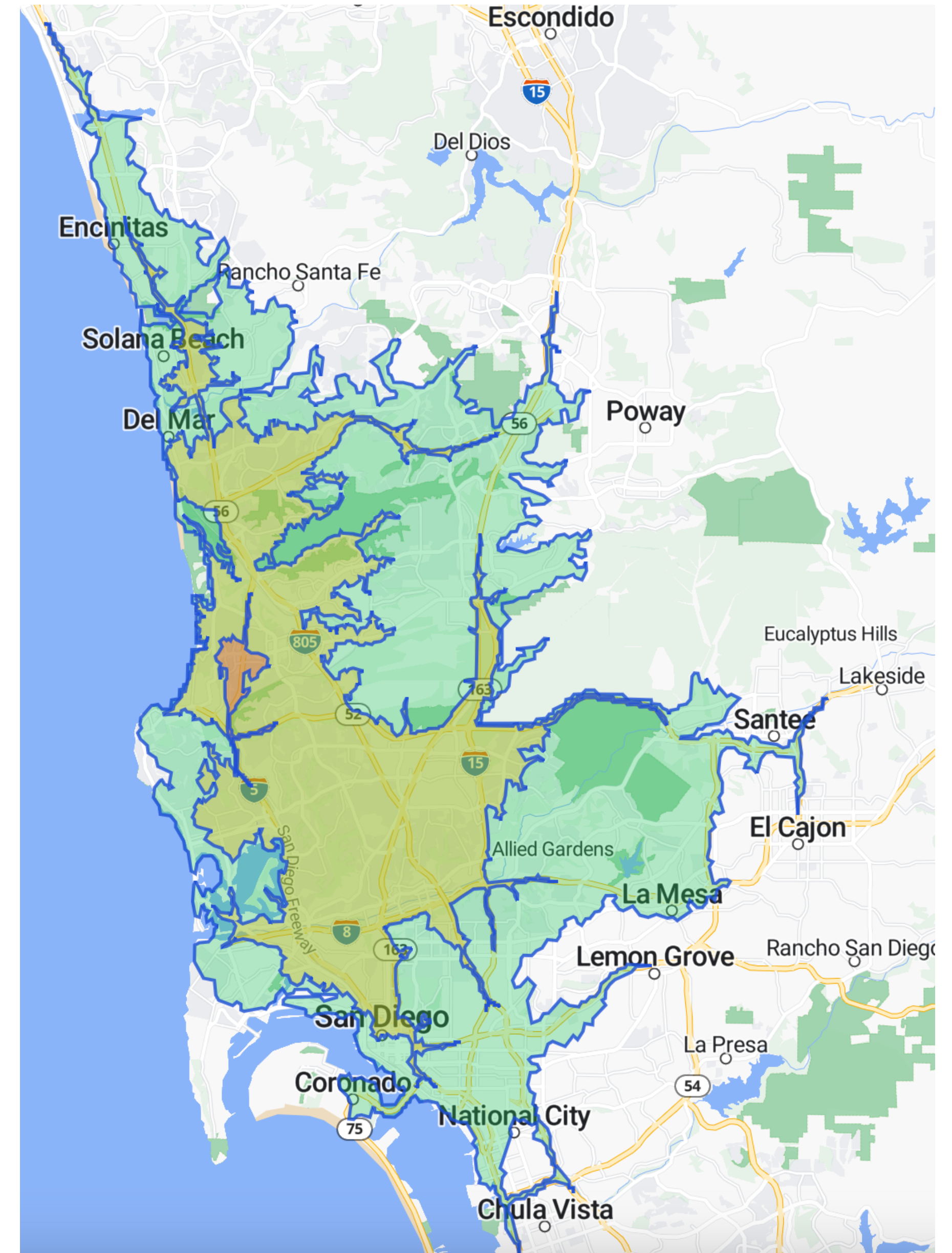


A mathematical model

- Call the set of all possible destinations $\mathcal{X}$
- Define the function $\eta : \mathcal{X} \rightarrow [0,1]$

$$\eta(x) = \Pr \left(\text{reach } x \text{ within 20 min} \right)$$

- We get a bunch of data points (X_n, Y_n)
 - X_n = destination

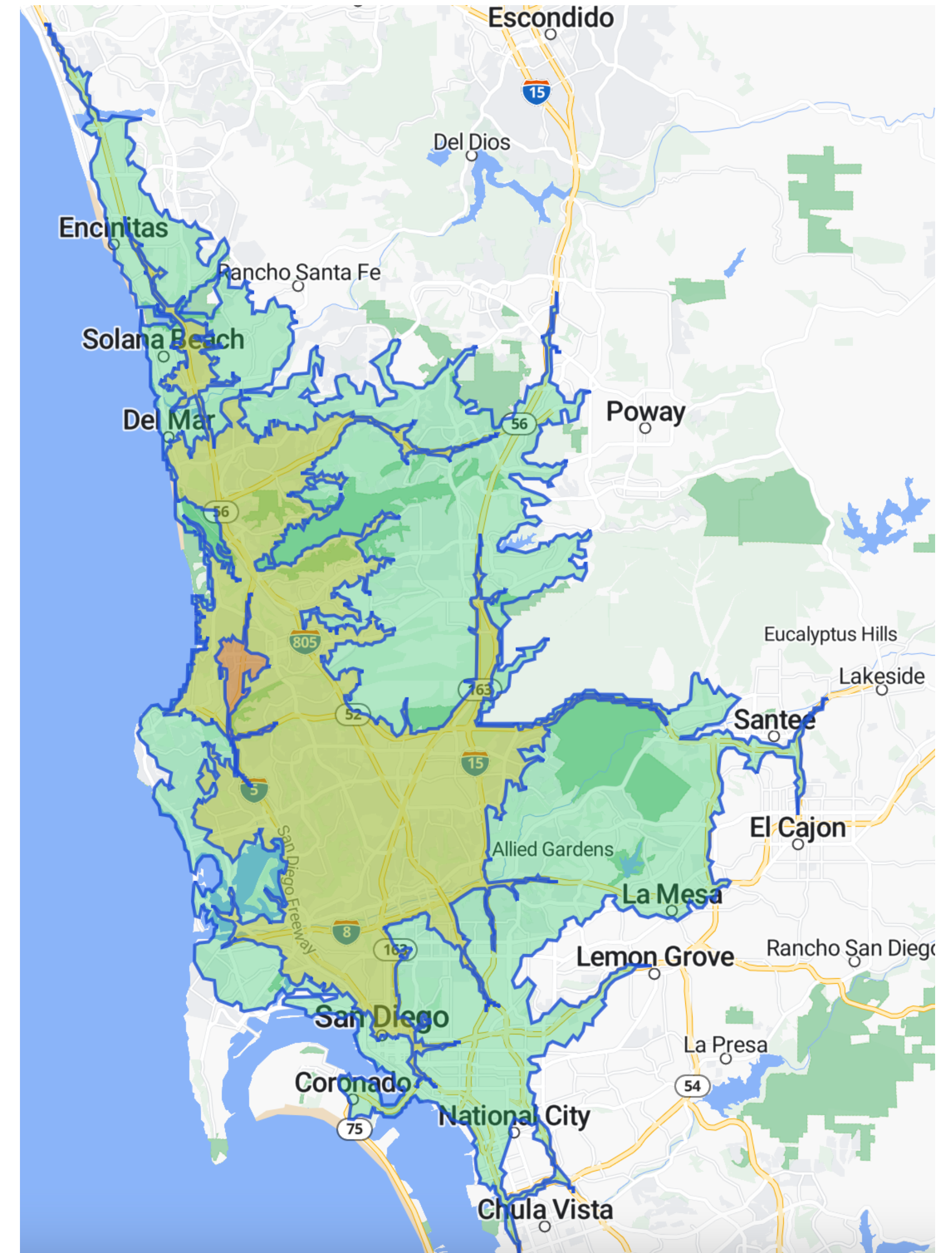


A mathematical model

- Call the set of all possible destinations $\mathcal{X}$
- Define the function $\eta : \mathcal{X} \rightarrow [0,1]$

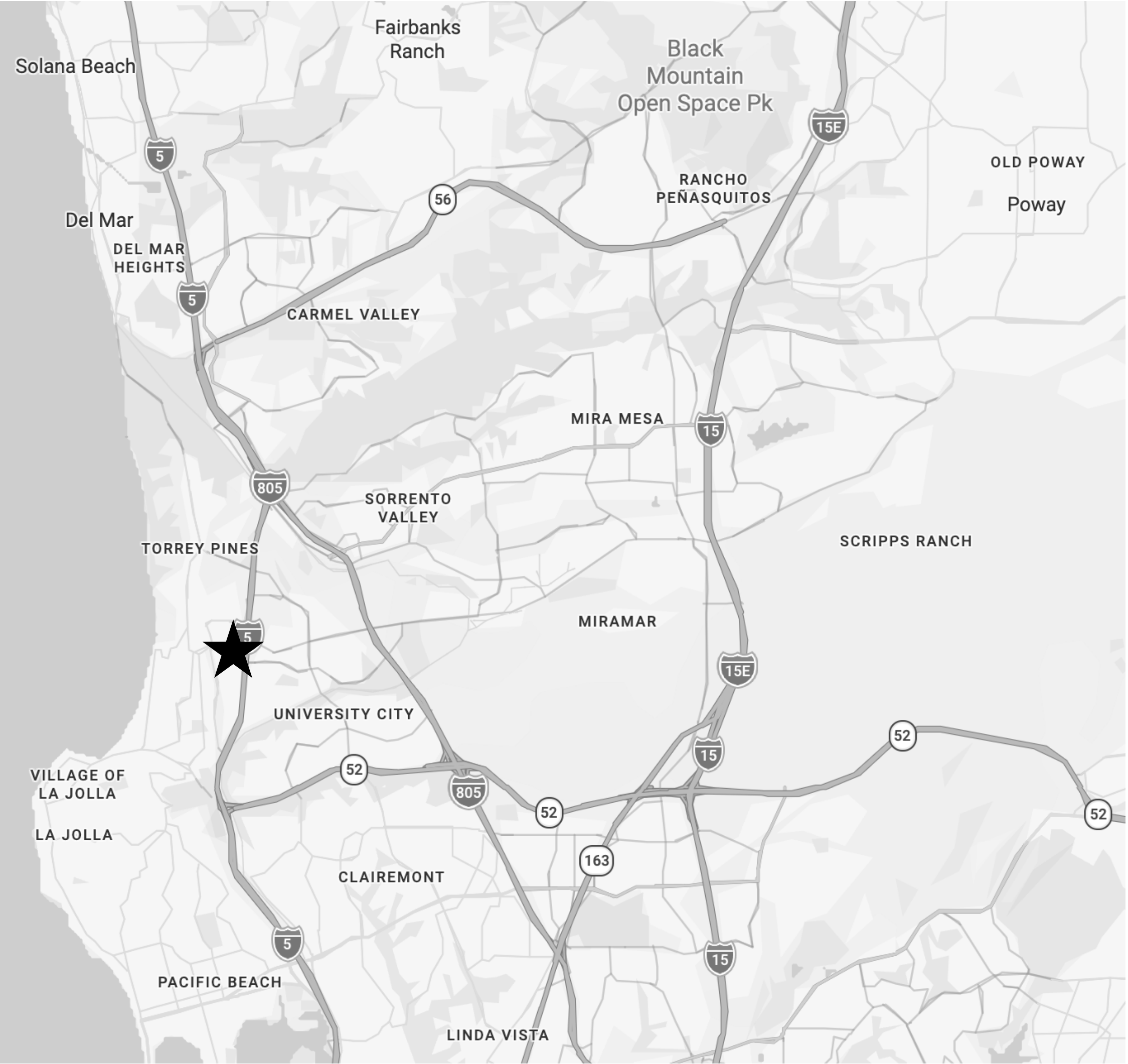
$$\eta(x) = \Pr \left(\text{reach } x \text{ within 20 min} \right)$$

- We get a bunch of data points (X_n, Y_n)
 - X_n = destination
 - $Y_n = \mathbf{1}\{\text{arrived in } < 20 \text{ min}\}$



Given dataset

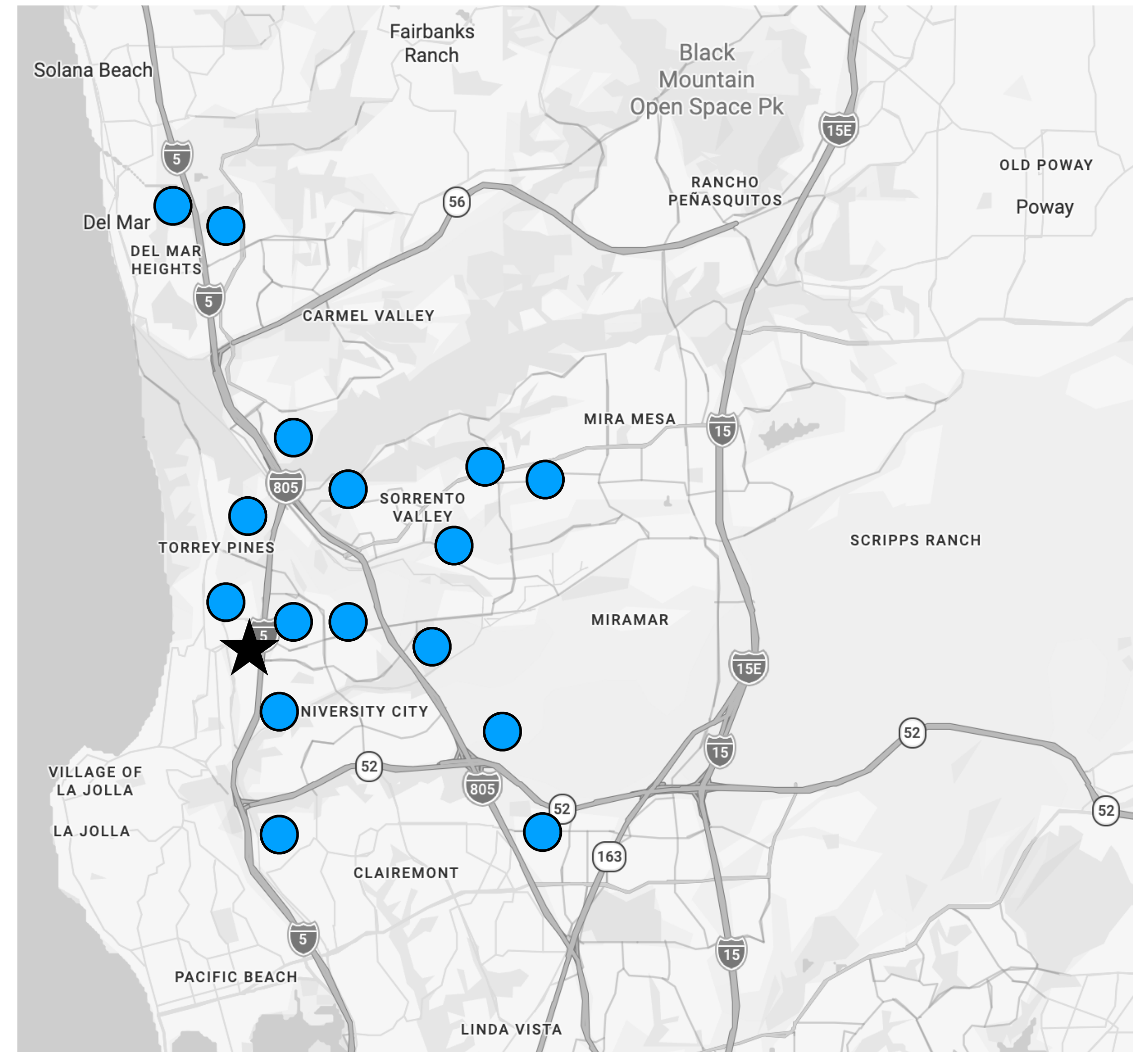
★ = UC San Diego



Given dataset

★ = UC San Diego

● = instance of arriving within 20 min

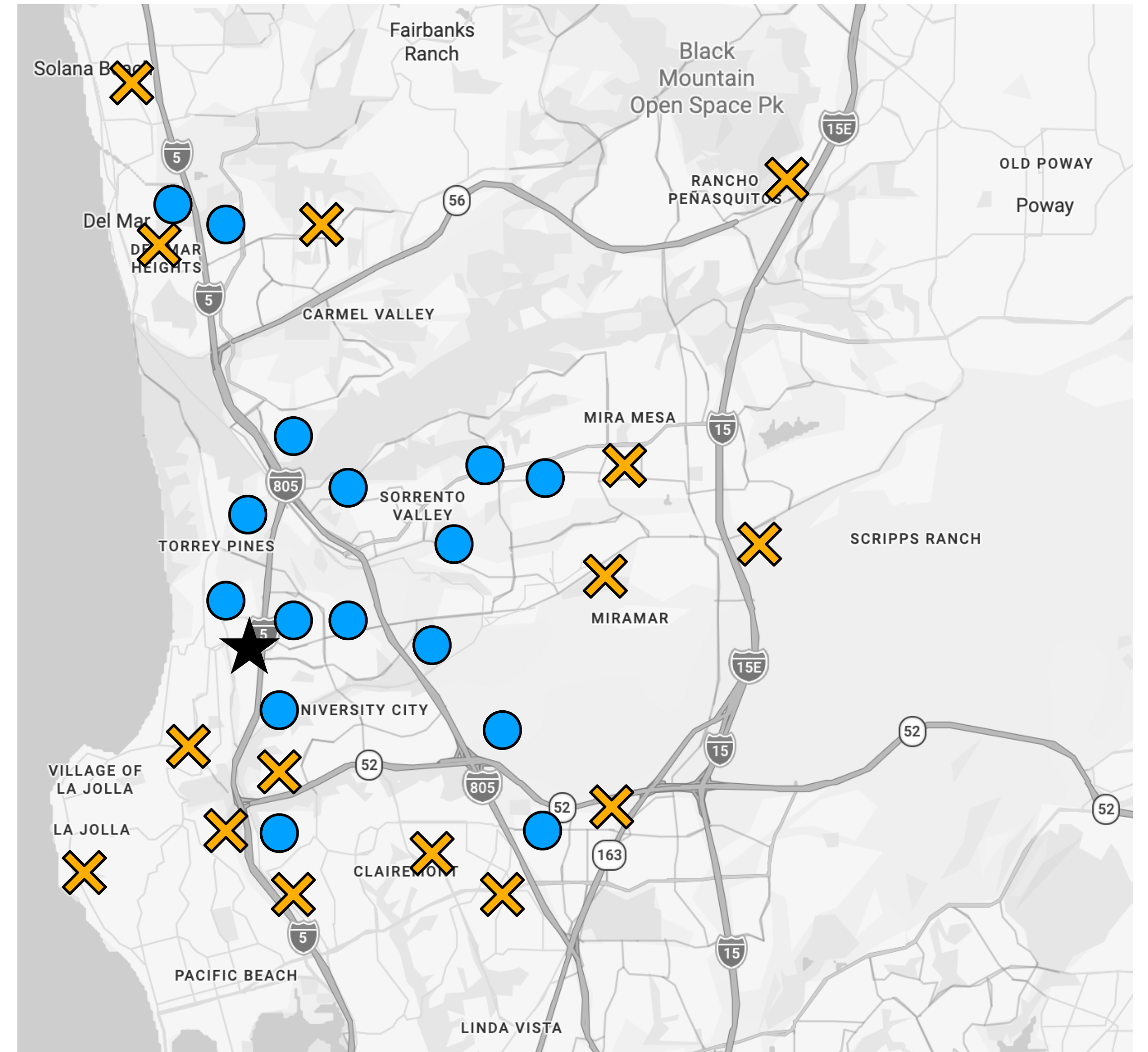


Given dataset

★ = UC San Diego

● = instance of arriving within 20 min

✕ = instance of failing to make it.



How was this data generated?

Adaptive Sampling Procedure

How was this data generated?

Adaptive Sampling Procedure

For $n = 1, 2, \dots$:

How was this data generated?

Adaptive Sampling Procedure

For $n = 1, 2, \dots$:

- The data process chooses X_n

How was this data generated?

Adaptive Sampling Procedure

For $n = 1, 2, \dots$:

- The data process chooses X_n
- A response is revealed: $Y_n \sim \text{Bernoulli}(\eta(X_n))$

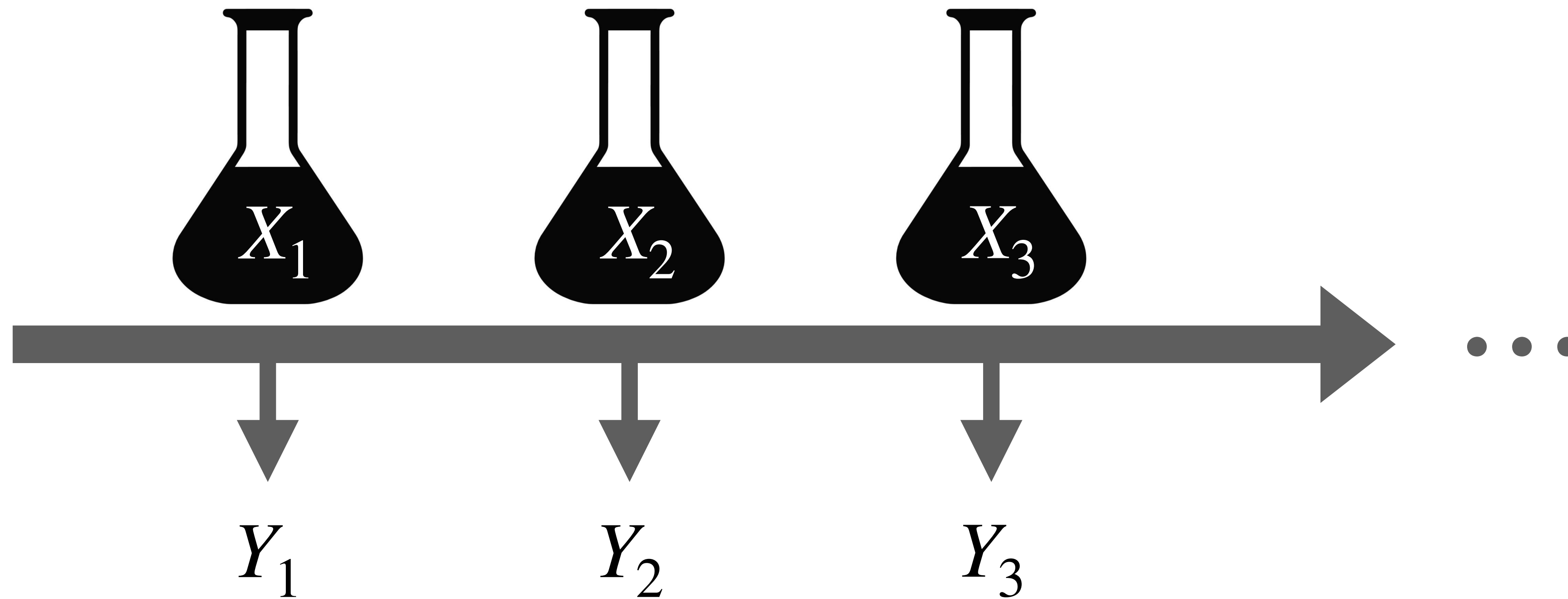
How was this data generated?

Adaptive Sampling Procedure

For $n = 1, 2, \dots$:

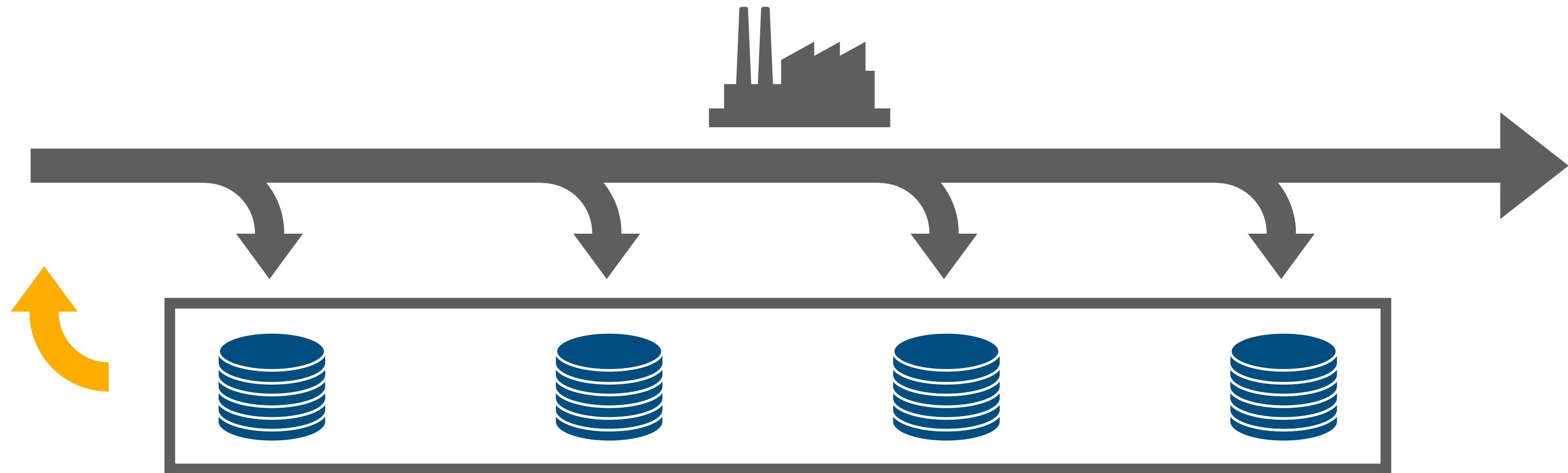
- The data process chooses X_n **[this choice can depend on past data]**
- A response is revealed: $Y_n \sim \text{Bernoulli}(\eta(X_n))$

Example: exploratory data analysis



Often, the next question X_n to ask depends on past findings $Y_{<n}$.

Example: data from adaptive processes



A company's decisions X_n will depend on past outcomes $\mathbb{Y}_{<n}$.

Example: data from a crowd

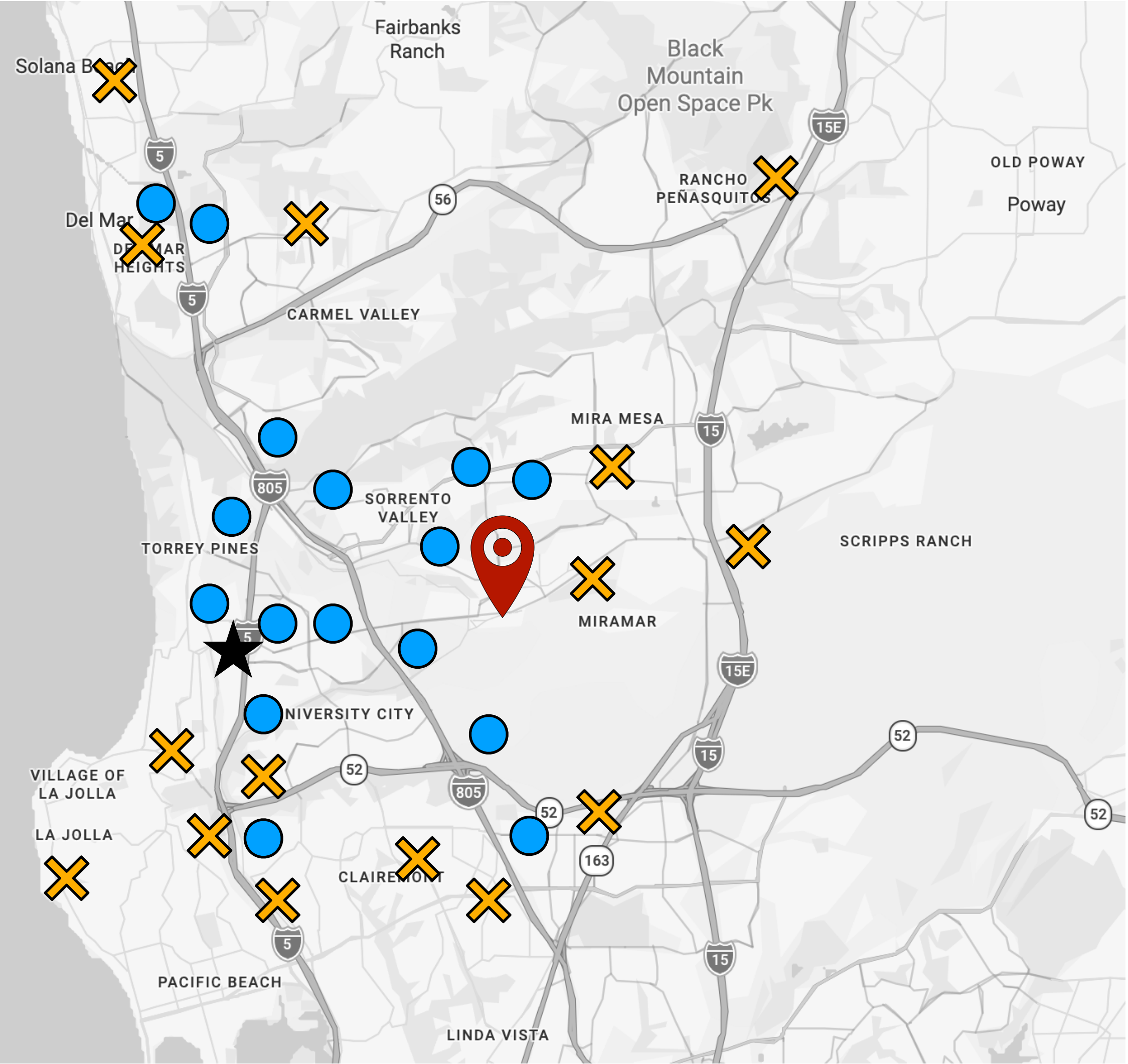


A citizen science scenario

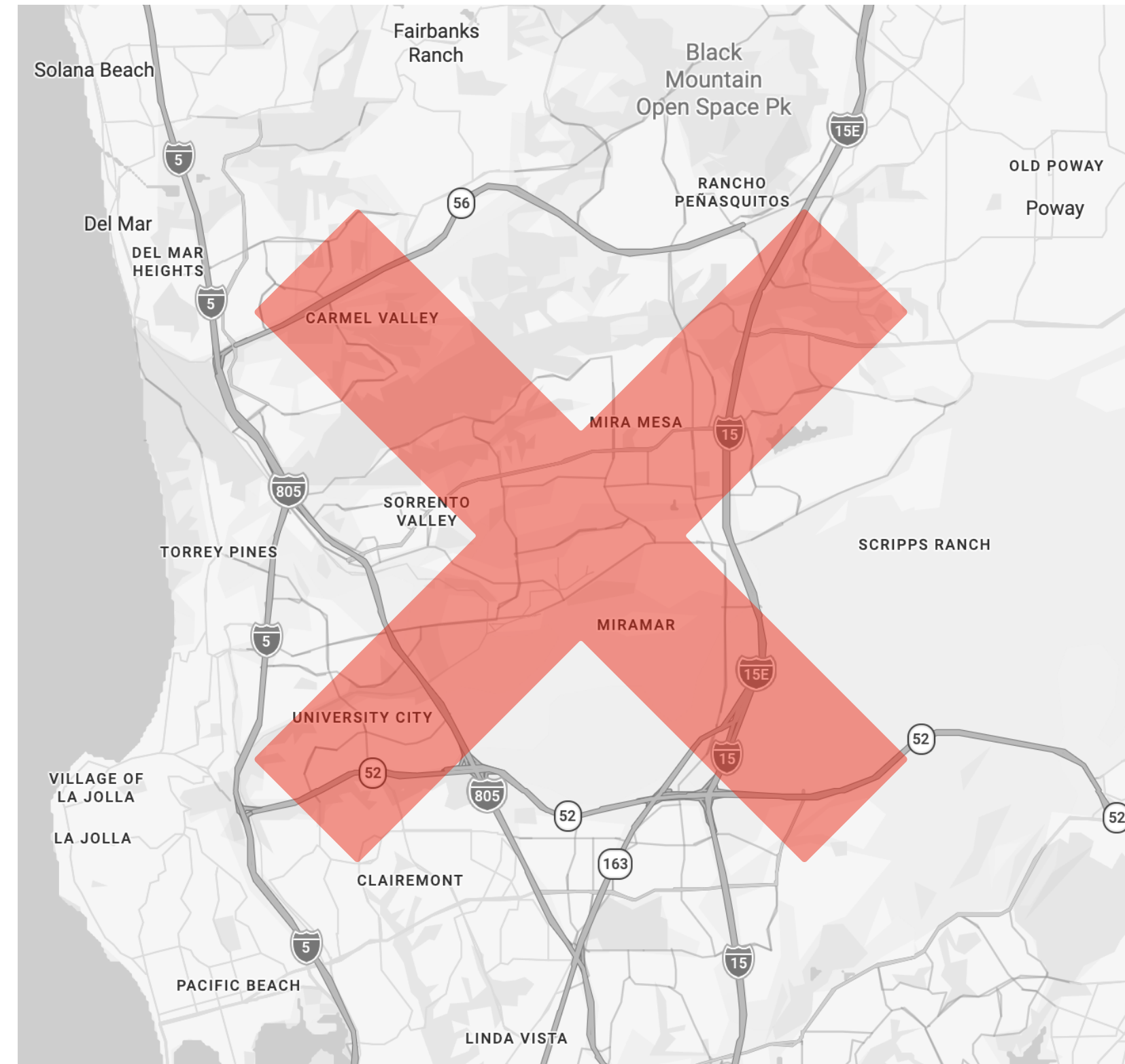
Alice: I saw a chuckwalla in Joshua Tree yesterday 🤯

Bob: I *must* go to JT now!

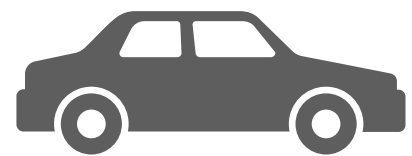
Can we reach  within 20 minutes?



An even simpler world

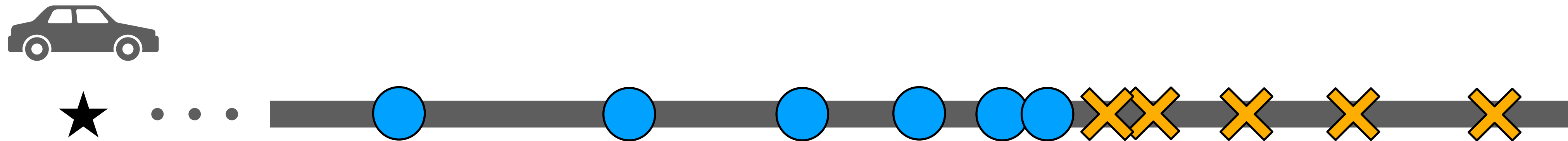


An even simpler world



$$\mathcal{X} = [0, 1]$$

An even simpler world



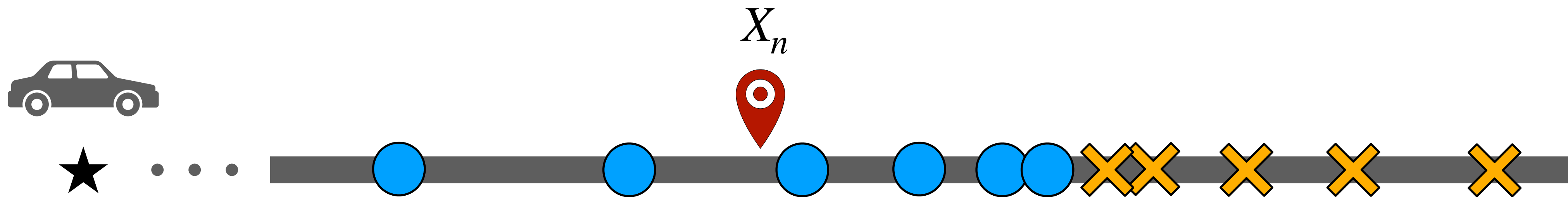
Adaptively sampled dataset

★ = UC San Diego

● = instance of arriving within 20 min

✕ = instance of failing to make it.

What is $\eta(X_n)$?



Adaptively sampled dataset

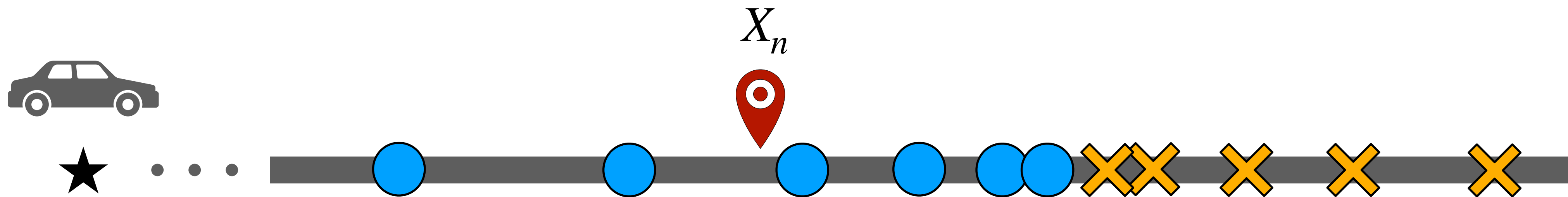
★ = UC San Diego

● = instance of arriving within 20 min

✕ = instance of failing to make it.

What is $\eta(X_n)$?

$$\eta(X_n) = \Pr \left(\text{reaching } X_n \text{ within 20 min} \right)$$



Adaptively sampled dataset

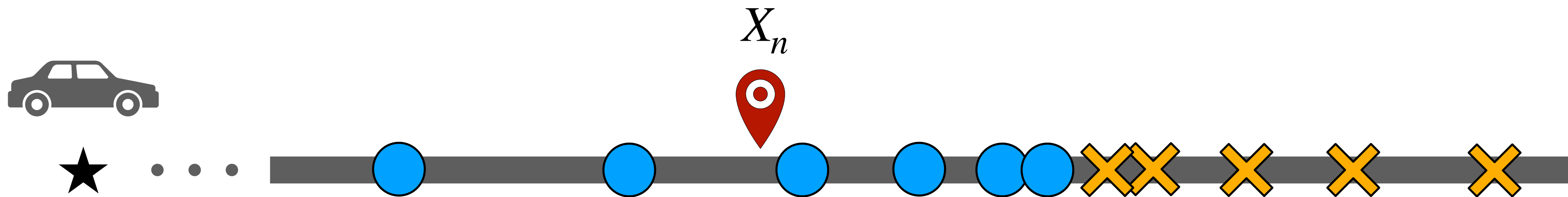
★ = UC San Diego

● = instance of arriving within 20 min

✕ = instance of failing to make it.

What is $\eta(X_n)$?

Seemingly reasonable guess: $\hat{\eta}(X_n) \approx 1$.



Adaptively sampled dataset

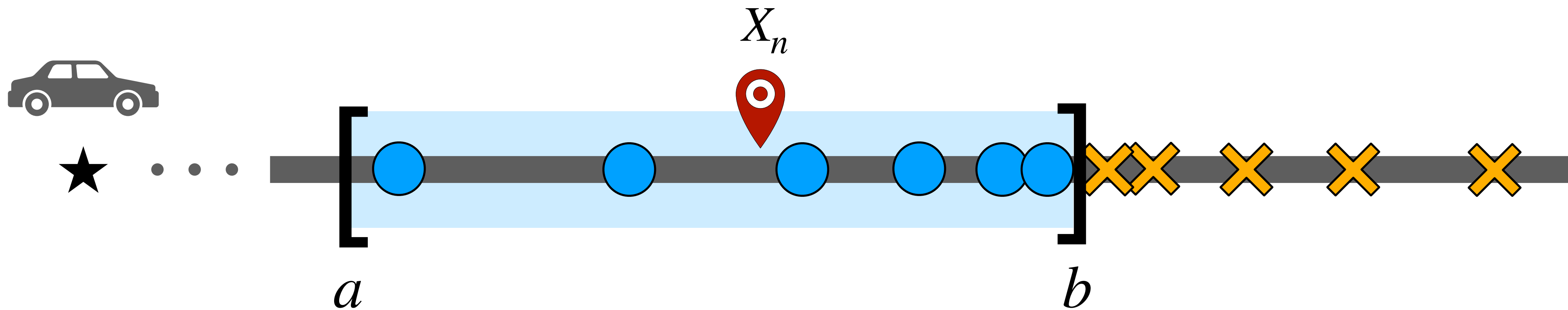
★ = UC San Diego

● = instance of arriving within 20 min

✕ = instance of failing to make it.

What is $\eta(X_n)$?

Seemingly reasonable guess: $\hat{\eta}(X_n) \approx 1$.



There are many instances of success around X_n , so it seems likely that we can make it within 20 minutes.

Let's see how the world actually works and how the data was collected...

The actual world: Schrödinger's traffic

This road either has very heavy traffic...



$$\mathcal{X} = [0,1]$$

The actual world: Schrödinger's traffic

This road either has very heavy traffic, or no traffic at all.



$$\mathcal{X} = [0,1]$$

The actual world: Schrödinger's traffic

This road either has very heavy traffic, or no traffic at all.

Making it anywhere in 20 min is **always a coin toss**.



$$\mathcal{X} = [0,1]$$

The actual world: Schrödinger's traffic

This road either has very heavy traffic, or no traffic at all.

Pr(arrive < 20 min)



$\mathcal{X} = [0, 1]$

The learner's model

The learner's model

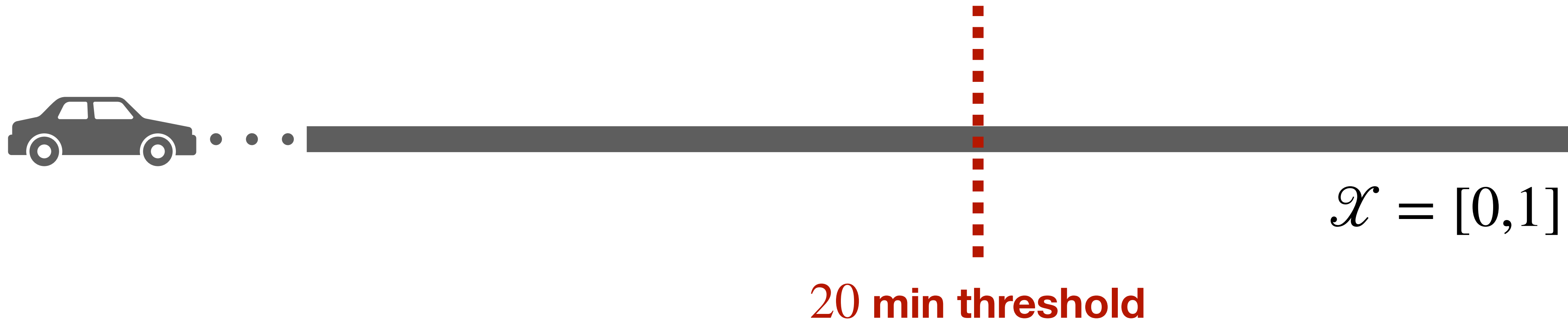
Learner: “I think the **further** away something is...

The learner's model

Learner: “I think the **further** away something is, the **longer** it takes to arrive.”

The learner's model

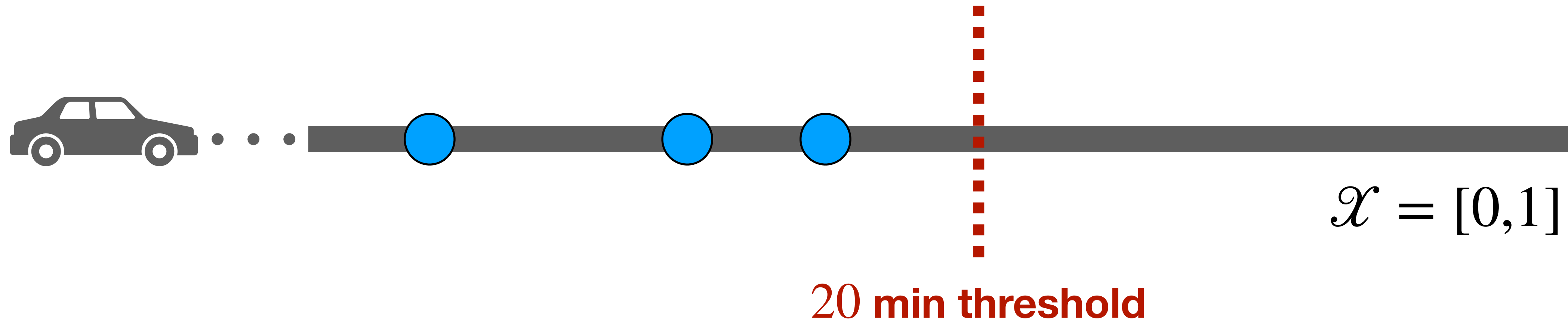
Learner: “I think the **further** away something is, the **longer** it takes to arrive.”



Inductive bias. There is a threshold: everything before takes less than 20 minutes to reach, everything after takes more.

The learner's model

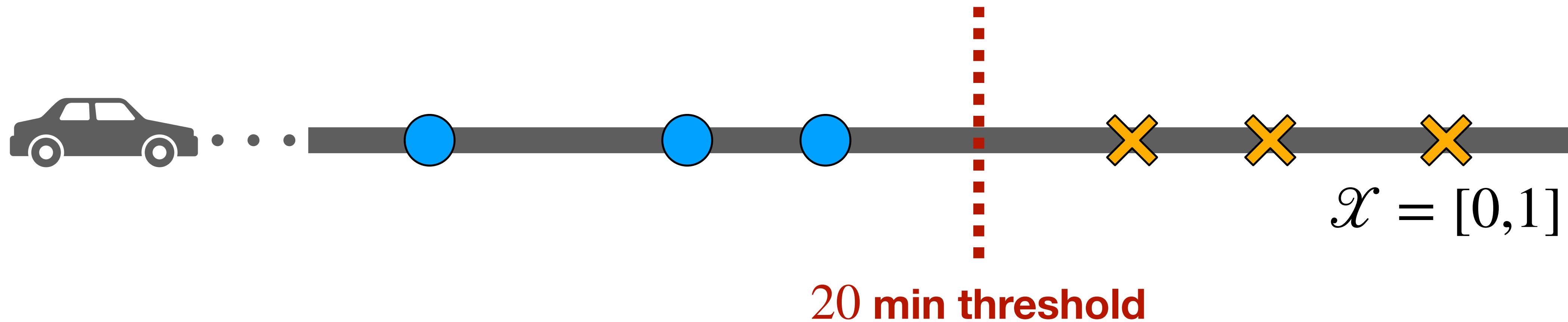
Learner: “I think the **further** away something is, the **longer** it takes to arrive.”



Inductive bias. There is a threshold: **everything before** takes less than 20 minutes to reach, **everything after** takes more.

The learner's model

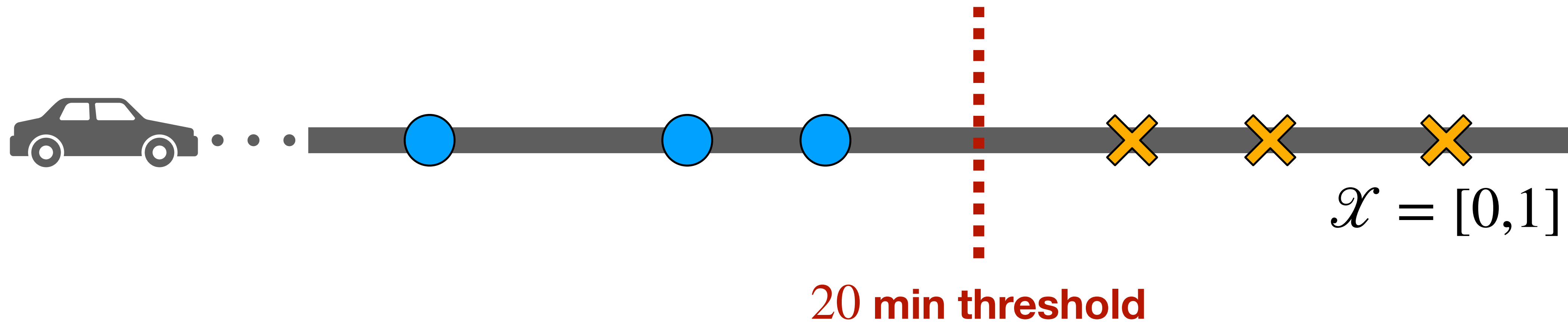
Learner: “I think the **further** away something is, the **longer** it takes to arrive.”



Inductive bias. There is a threshold: everything before takes less than 20 minutes to reach, everything after takes more.

The learner's model (wrong!)

Learner: “I think the **further** away something is, the **longer** it takes to arrive.”



Inductive bias. There is a threshold: everything before takes less than 20 minutes to reach, everything after takes more.

Learning algorithm: binary search



Learning algorithm: binary search



Day 1:

- Learner tries driving to X_1

Learning algorithm: binary search



Day 1:

- Learner tries driving to X_1
- Nature flips a coin: no traffic

Learning algorithm: binary search



Day 1:

- Learner tries driving to X_1
- Nature flips a coin: no traffic
 - Learner arrives < 20 min

Learning algorithm: binary search



Day 1:

- Learner tries driving to X_1
- Nature flips a coin: no traffic
 - Learner arrives < 20 min

Learning algorithm: binary search



Day 2:

- Learner tries driving to X_2

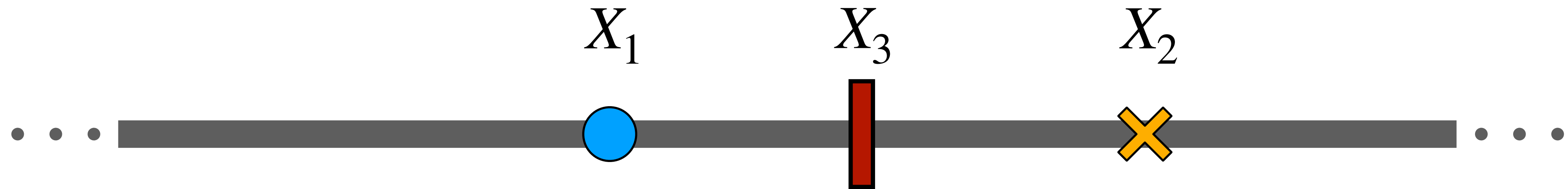
Learning algorithm: binary search



Day 2:

- Learner tries driving to X_2
- Nature flips a coin: traffic
 - Learner arrives > 20 min

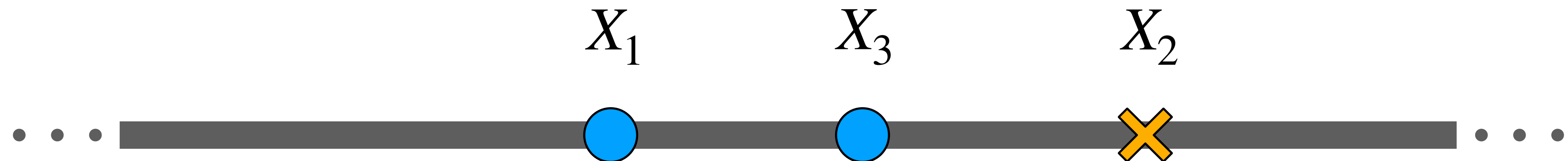
Learning algorithm: binary search



Day 3:

- Learner tries driving to X_3

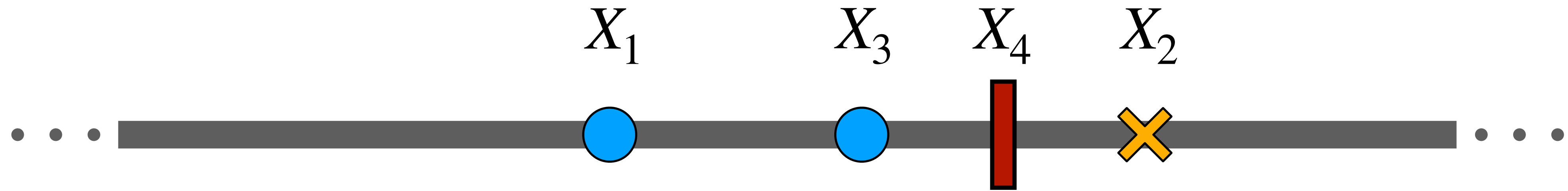
Learning algorithm: binary search



Day 3:

- Learner tries driving to X_3
- Nature flips a coin: no traffic
 - Learner arrives < 20 min

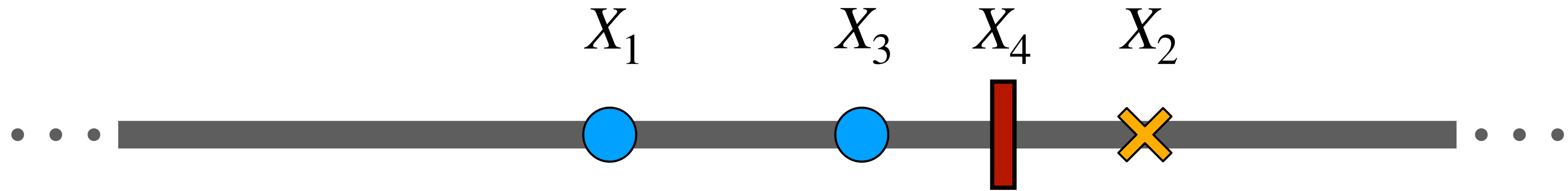
Learning algorithm: binary search



Day 4:

- Learner tries driving to X_4

Learning algorithm: binary search

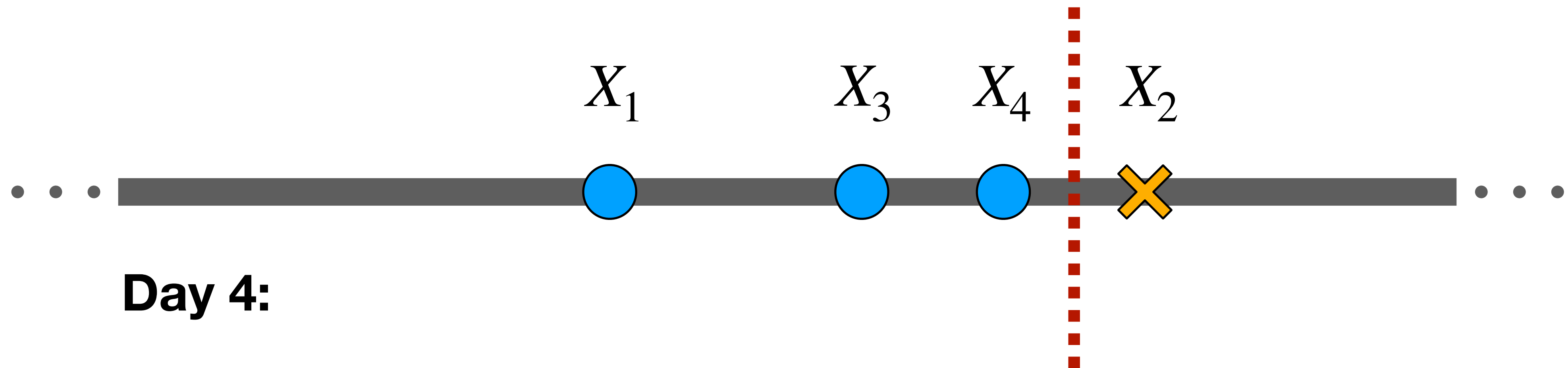


Day 4:

- Learner tries driving to X_4

No matter the outcome, the data continues to look linearly separable!

Learning algorithm: binary search

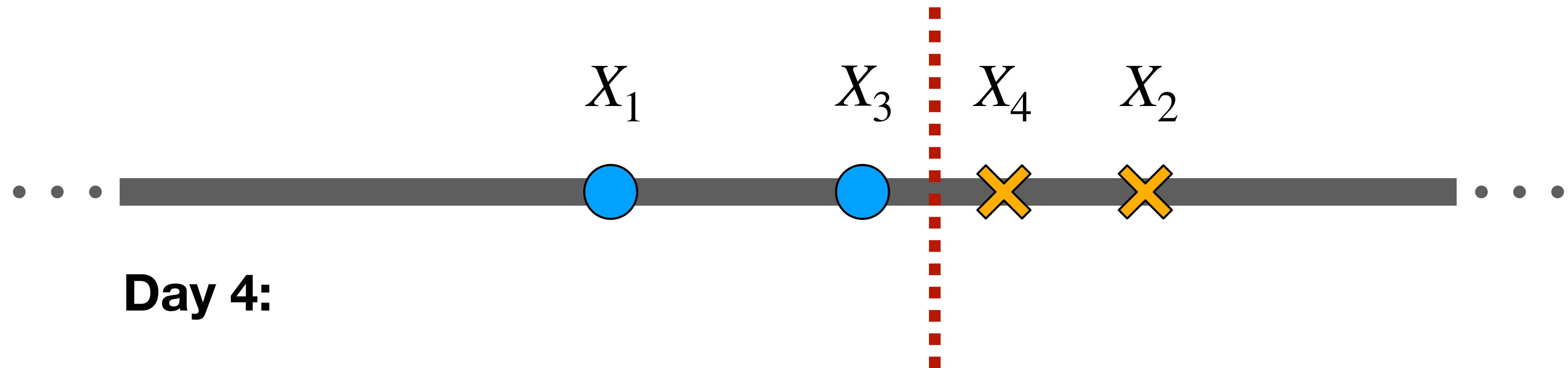


Day 4:

- Learner tries driving to X_4

No matter the outcome, the data continues to look linearly separable!

Learning algorithm: binary search

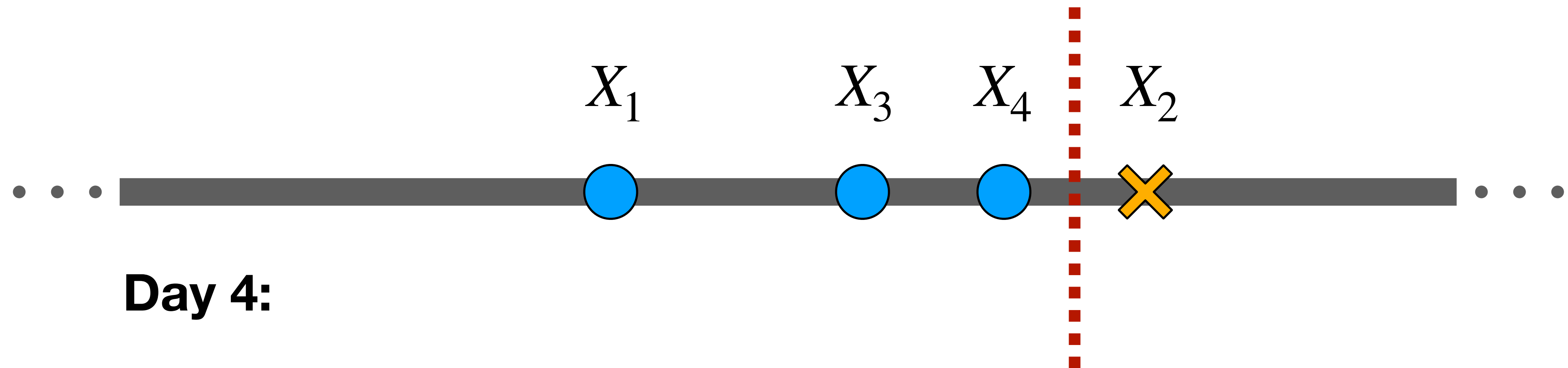


Day 4:

- Learner tries driving to X_4

No matter the outcome, the data continues to look linearly separable!

Learning algorithm: binary search

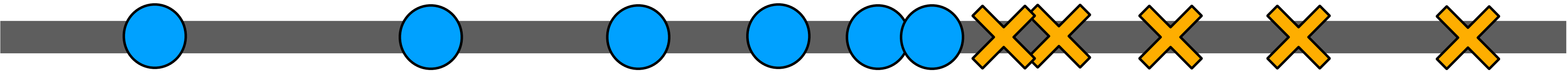


Day 4:

- Learner tries driving to X_4

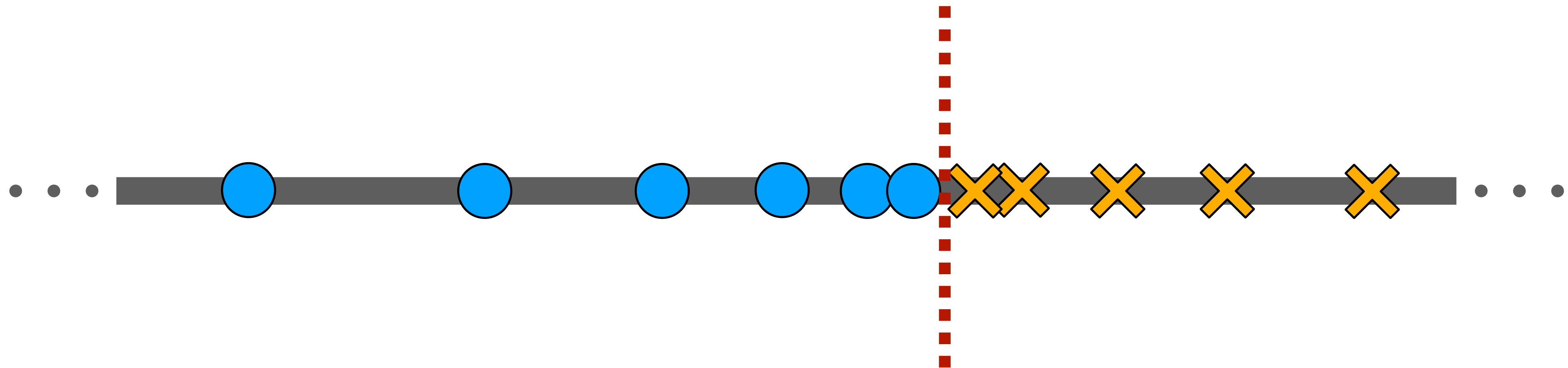
No matter the outcome, the data continues to look linearly separable!

Eventually, the data that we received:



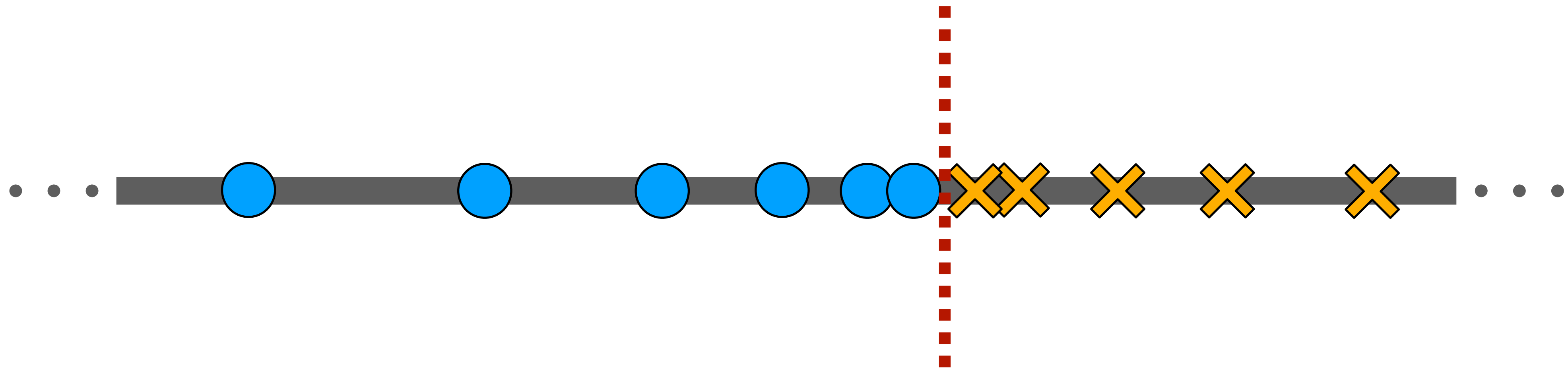
(not to scale)

Spurious patterns can be created from noise



Linear separability here is an artifact of **how the data was sampled**.

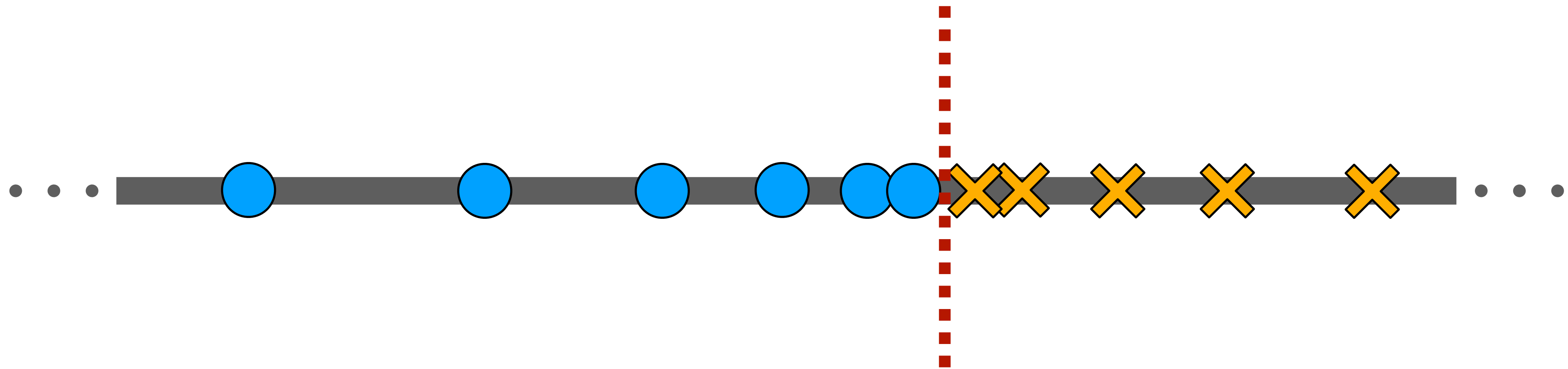
Spurious patterns can be created from noise



Linear separability here is an artifact of **how the data was sampled**.

- It is not reflective of the underlying relationship between X and Y captured by η .

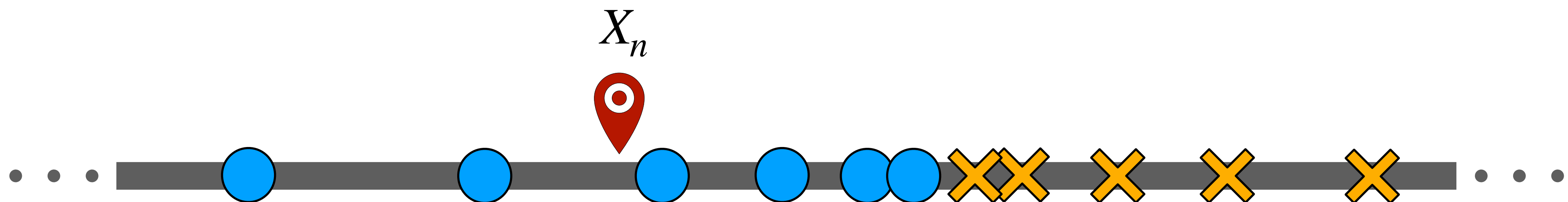
Spurious patterns can be created from noise



Linear separability here is an artifact of **how the data was sampled**.

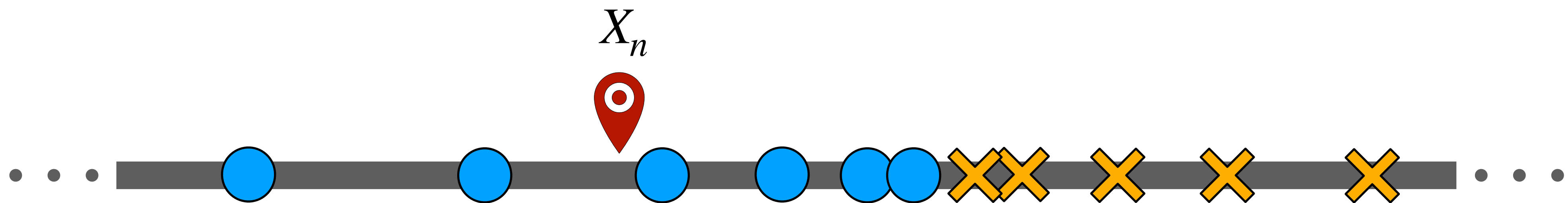
- It is not reflective of the underlying relationship between X and Y captured by η .
- The observed pattern will likely not generalize to data sampled elsewhere.

What is $\eta(X_n)$?

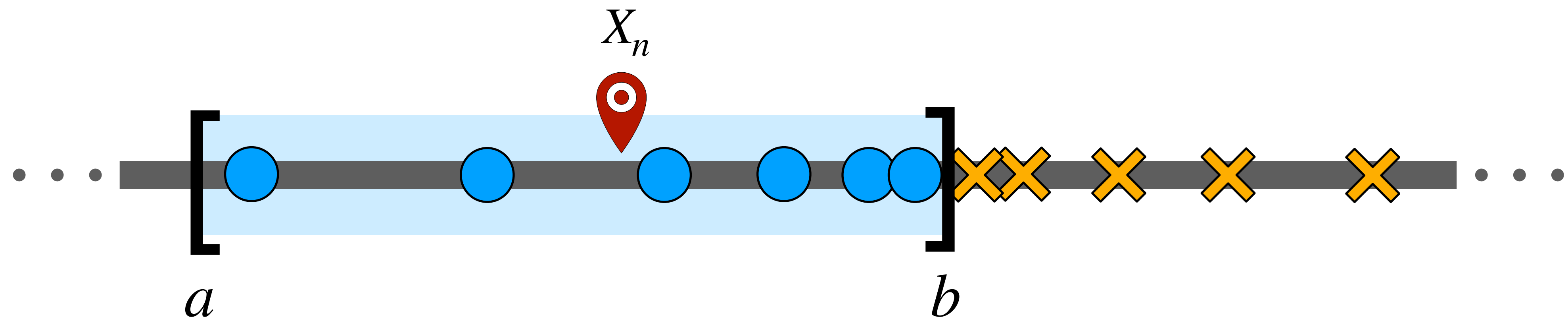


What is $\eta(X_n)$?

The actual answer: $\eta(X_n) = \frac{1}{2}$.

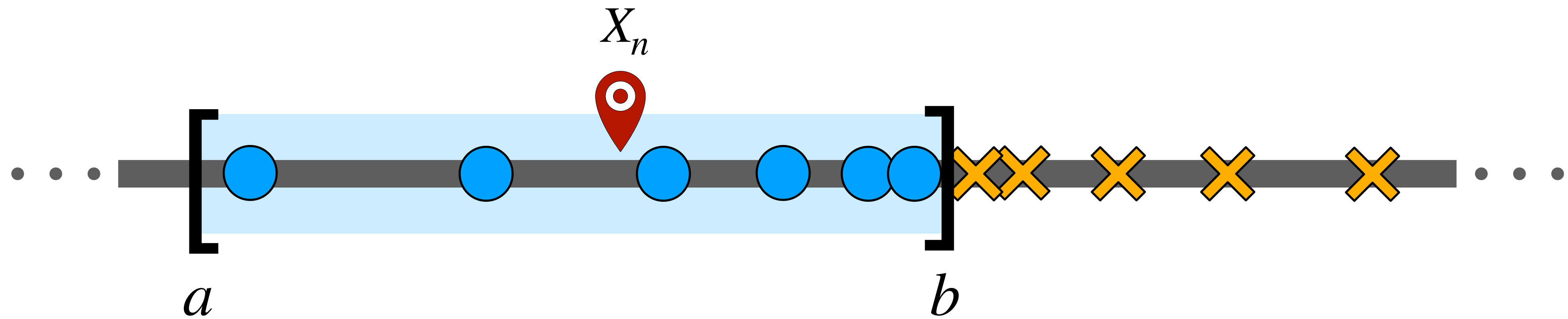


A very imbalanced interval



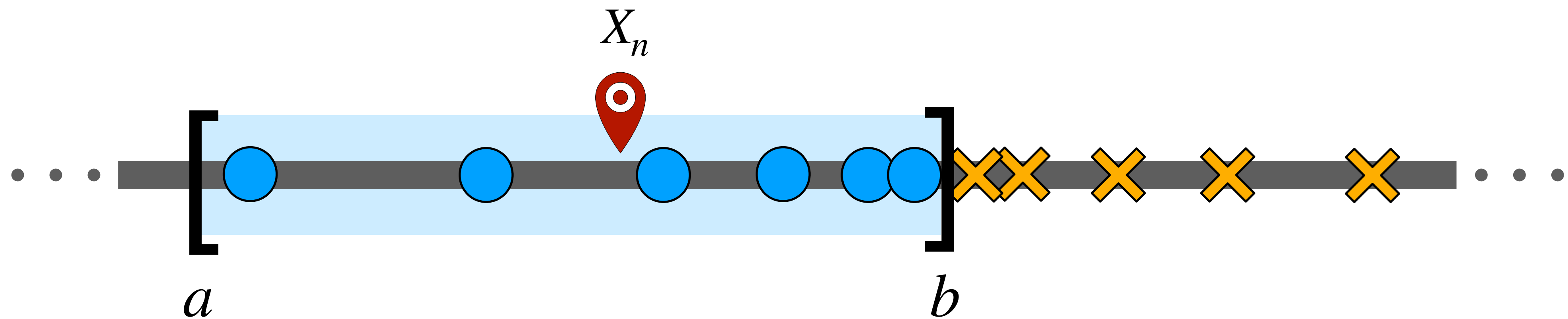
This interval has many points, but its average label is far from the “expected” $\eta = 1/2$.

A very imbalanced interval



In theory, adaptively-generated data can behave very counterintuitively.

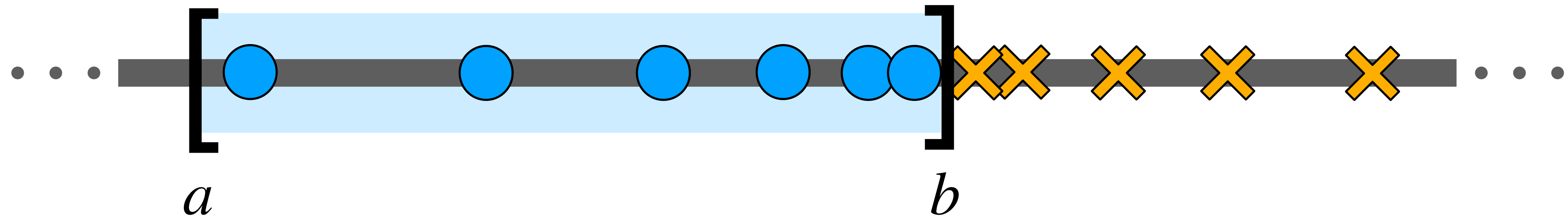
A very imbalanced interval



In theory, adaptively-generated data can behave very counterintuitively.

Uniform convergence over VC classes does not apply to adaptive data.

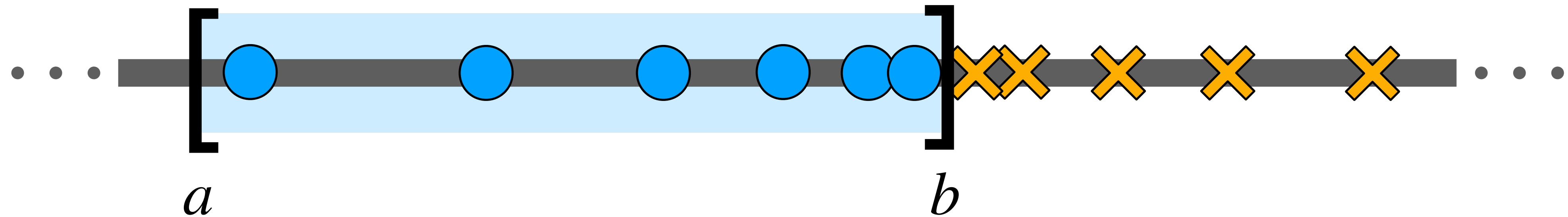
Sequential uniform convergence



A necessary and sufficient condition for uniform convergence to hold in the sequential setting is **finite Littlestone* dimension**.

[Rakhlin, Sridharan, Tewari 2013]

Sequential uniform convergence



A necessary and sufficient condition for uniform convergence to hold in the sequential setting is **finite Littlestone* dimension**.

[Rakhlin, Sridharan, Tewari 2013]

*Technically, the sequential fat-shattering dimension.

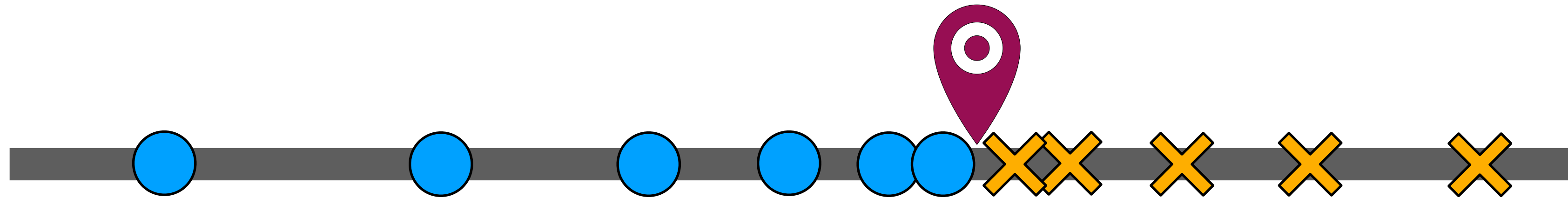
Upshot

There are potential pitfalls with working with data from a black box!

Using adaptive data

or, are things actually that bad in practice?

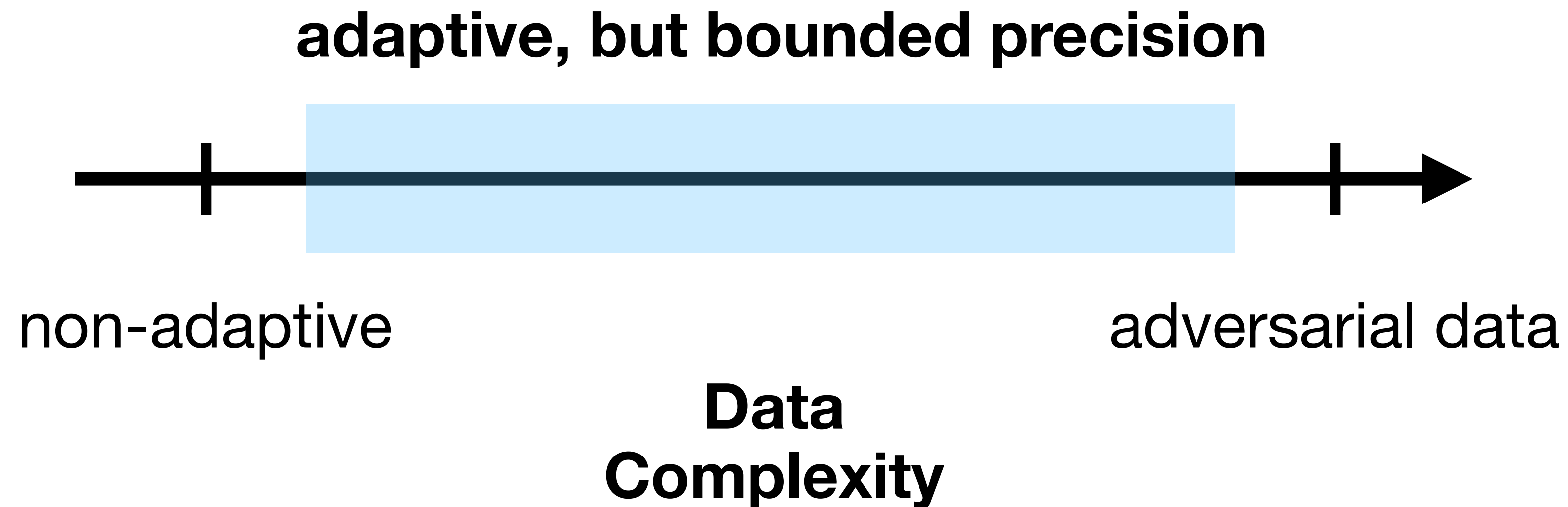
Is this data-process likely to occur?



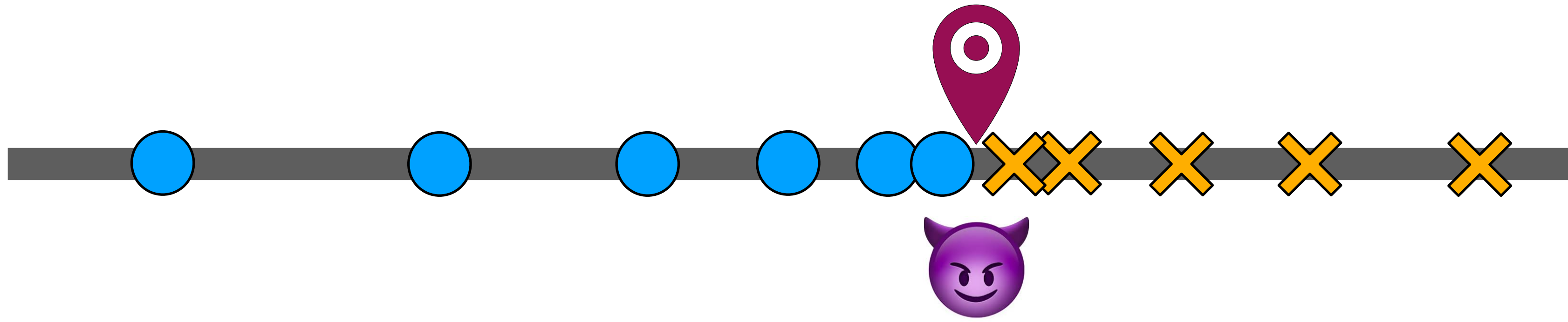
To continue the spurious pattern of linear separability, we need to be able to sample from **bad regions that shrink exponentially quickly.**

Arbitrarily bad data needs infinite precision

Idea: misleading data seems unlikely to persist if it is not generated by a fully adversarial process with **infinite precision**.

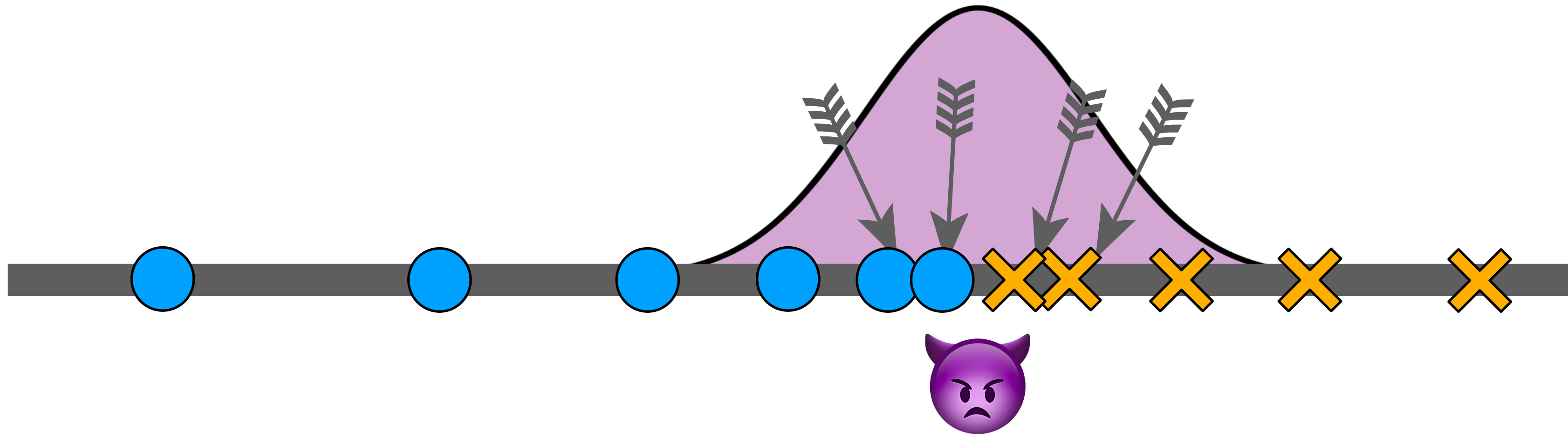


Data processes with bounded precision



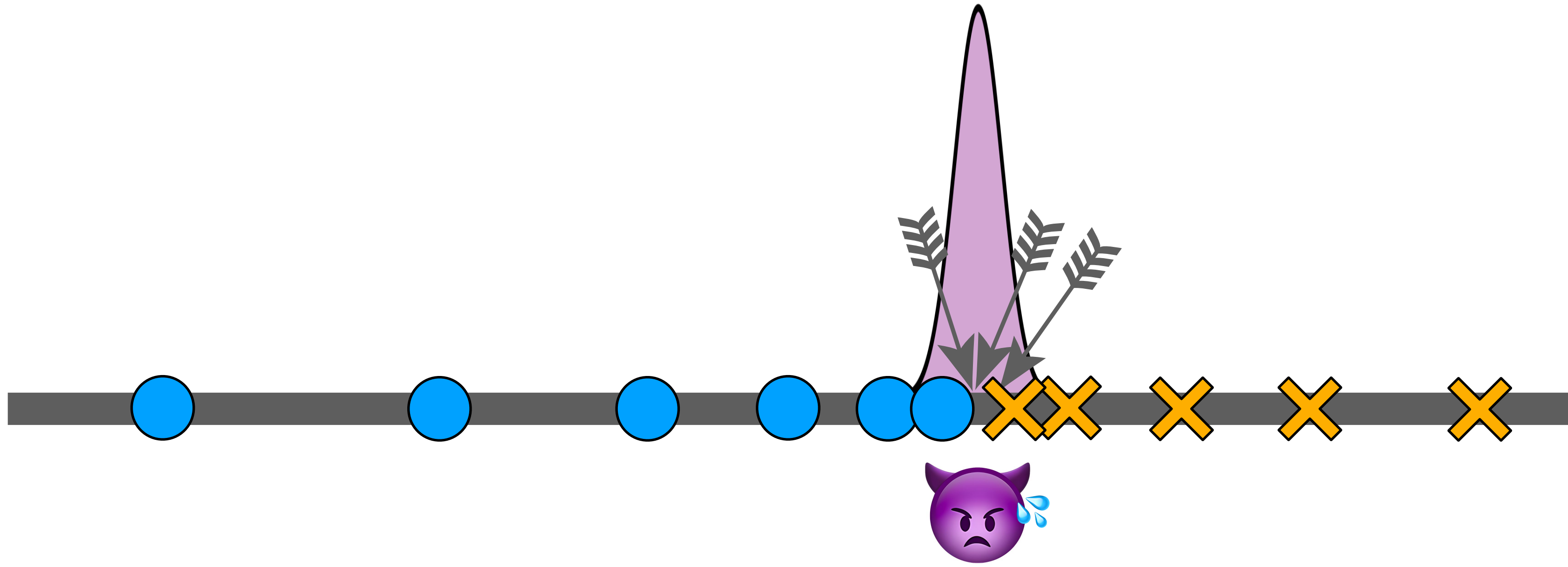
Suppose that the data process “wants” to sample this point.

Data processes with bounded precision



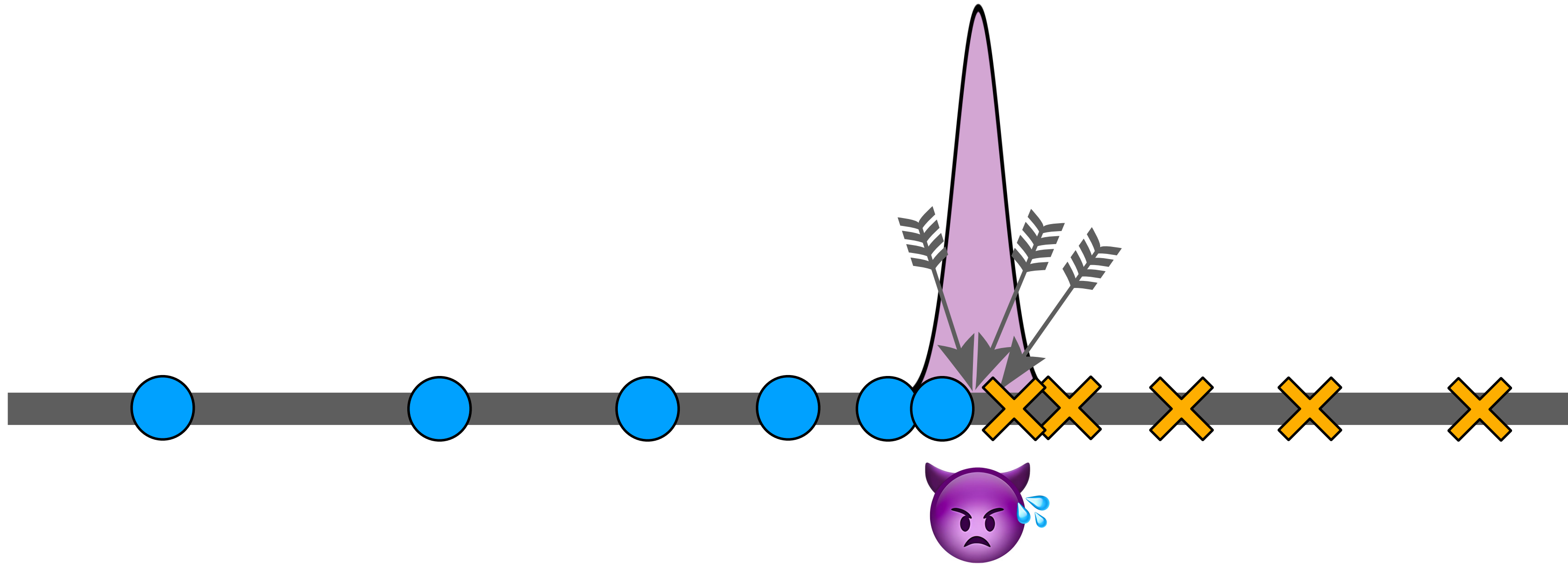
Real data processes are physical mechanisms,
so there is **inherent imprecision**.

Data processes with bounded precision



And **reducing variability** is “costly”.

Data processes with bounded precision



And **reducing variability** is “costly”.

Can we formalize this? ↗

Data processes with bounded precision



Introduce a measure μ on $\mathcal{X}$ capturing a notion of size.

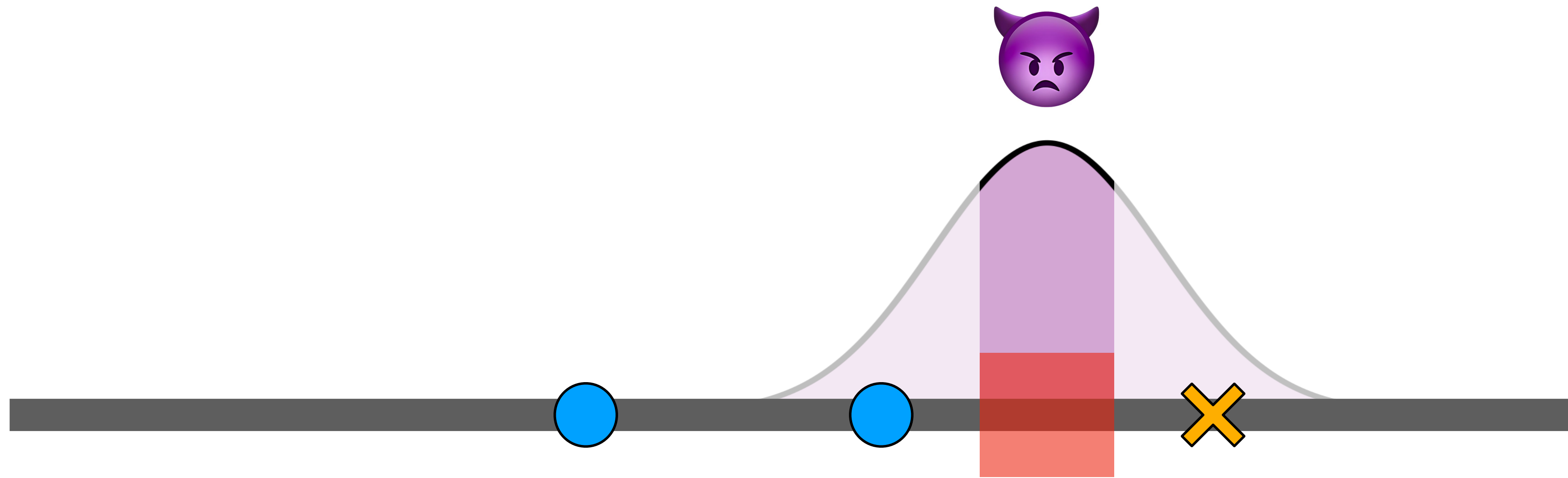
Data processes with bounded precision



Introduce a measure μ on $\mathcal{X}$ capturing a notion of size.

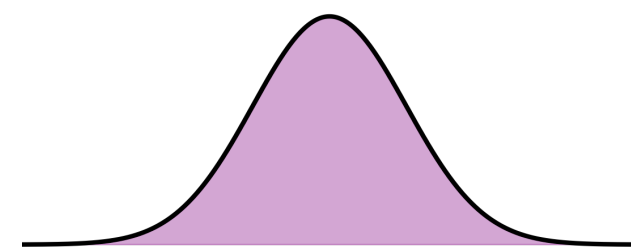
$\mu(\text{red square}) = \text{how large this region of space is}$

Data processes with bounded precision



Introduce a measure μ on $\mathcal{X}$ capturing size

$\mu(\text{red rectangle}) = \text{how large this region of space is}$



it will control how easily an adversary hits that region

Data processes with bounded precision

Definition (Dasgupta and So 2024)

Let $\mathcal{X}$ be an instance space with finite measure μ .

Data processes with bounded precision

Definition (Dasgupta and So 2024)

Let $\mathcal{X}$ be an instance space with finite measure μ . A data process is **uniformly absolutely continuous** if the following holds.

Data processes with bounded precision

Definition (Dasgupta and So 2024)

Let $\mathcal{X}$ be an instance space with finite measure μ . A data process is **uniformly absolutely continuous** if the following holds. For all $\varepsilon > 0$, there is a $\delta > 0$ so:

Data processes with bounded precision

Definition (Dasgupta and So 2024)

Let $\mathcal{X}$ be an instance space with finite measure μ . A data process is **uniformly absolutely continuous** if the following holds. For all $\varepsilon > 0$, there is a $\delta > 0$ so:

$$\forall n \in \mathbb{N}, \quad \mu(A) < \delta \quad \implies \quad \dots$$

whenever $A \subset \mathcal{X}$ is measurable.



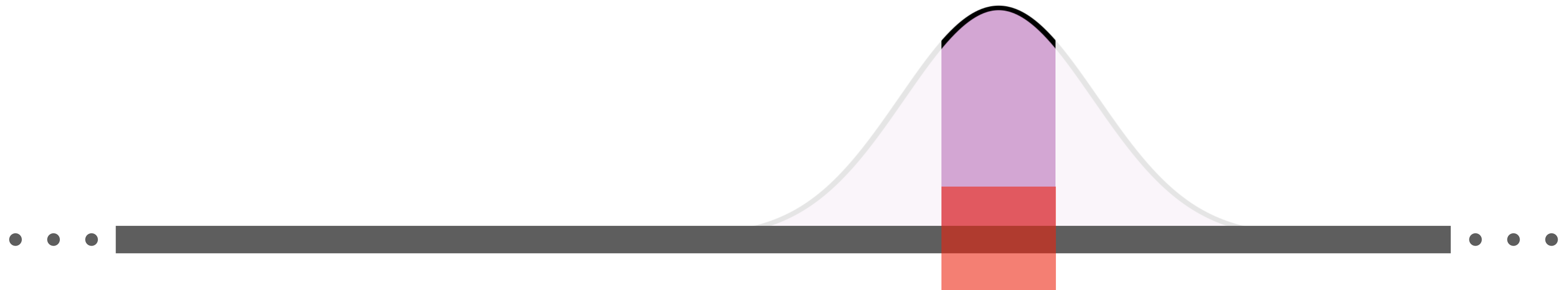
Data processes with bounded precision

Definition (Dasgupta and So 2024)

Let $\mathcal{X}$ be an instance space with finite measure μ . A data process is **uniformly absolutely continuous** if the following holds. For all $\varepsilon > 0$, there is a $\delta > 0$ so:

$$\forall n \in \mathbb{N}, \quad \mu(A) < \delta \quad \implies \quad \Pr(X_n \in A \mid \mathcal{F}_{n-1}) < \varepsilon,$$

whenever $A \subset \mathcal{X}$ is measurable.



Data processes with bounded precision

Definition (Dasgupta and So 2024)

Let $\mathcal{X}$ be an instance space with finite measure μ . A data process is **uniformly absolutely continuous** if the following holds. For all $\varepsilon > 0$, there is a $\delta > 0$ so:

$$\forall n \in \mathbb{N}, \quad \mu(A) < \delta \quad \implies \quad \Pr(X_n \in A \mid \mathcal{F}_{n-1}) < \varepsilon,$$

whenever $A \subset \mathcal{X}$ is measurable.

The adversary choose samples from small regions with small probability.

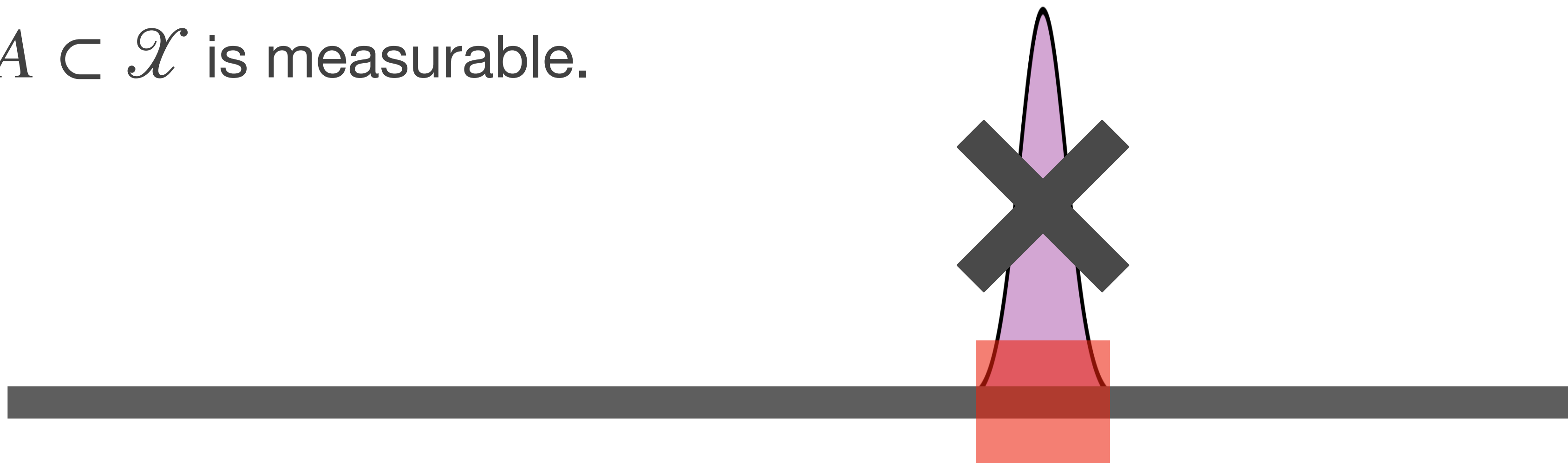
Data processes with bounded precision

Definition (Dasgupta and So 2024)

Let $\mathcal{X}$ be an instance space with finite measure μ . A data process is **uniformly absolutely continuous** if the following holds. For all $\varepsilon > 0$, there is a $\delta > 0$ so:

$$\forall n \in \mathbb{N}, \quad \mu(A) < \delta \quad \implies \quad \Pr(X_n \in A \mid \mathcal{F}_{n-1}) < \varepsilon,$$

whenever $A \subset \mathcal{X}$ is measurable.



Data processes with bounded precision

Definition (Dasgupta and So 2024)

Let $\mathcal{X}$ be an instance space with finite measure μ . A data process is **uniformly absolutely continuous** if the following holds. For all $\varepsilon > 0$, there is a $\delta > 0$ so:

$$\forall n \in \mathbb{N}, \quad \mu(A) < \delta \quad \implies \quad \underbrace{\Pr(X_n \in A \mid \mathcal{F}_{n-1})}_{\nu_m(A)} < \varepsilon,$$

whenever $A \subset \mathcal{X}$ is measurable.

The measure ν_m is absolutely continuous with respect to μ .

Data processes with bounded precision

Definition (Dasgupta and So 2024)

Let $\mathcal{X}$ be an instance space with finite measure μ . A data process is **uniformly absolutely continuous** if the following holds. For all $\varepsilon > 0$, there is a $\delta > 0$ so:

$$\forall n \in \mathbb{N}, \quad \mu(A) < \delta \quad \implies \quad \underbrace{\Pr(X_n \in A \mid \mathcal{F}_{n-1})}_{\nu_m(A)} < \varepsilon,$$

whenever $A \subset \mathcal{X}$ is measurable.

The measure ν_m is absolutely continuous with respect to μ .

- Uniform refers to a time-uniform bound.

Data processes with bounded precision

Definition (Quantitative version, Haghtalab et al 2020)

A data process is **L -Lipschitz** or is **L -smoothed** whenever

$$\forall n \in \mathbb{N}, \quad \mu(A) < \delta \implies \Pr(X_n \in A \mid \mathcal{F}_{n-1}) < L \cdot \delta,$$

whenever $A \subset \mathcal{X}$ is measurable.

Data processes with bounded precision

Definition (Quantitative version, Haghtalab et al 2020)

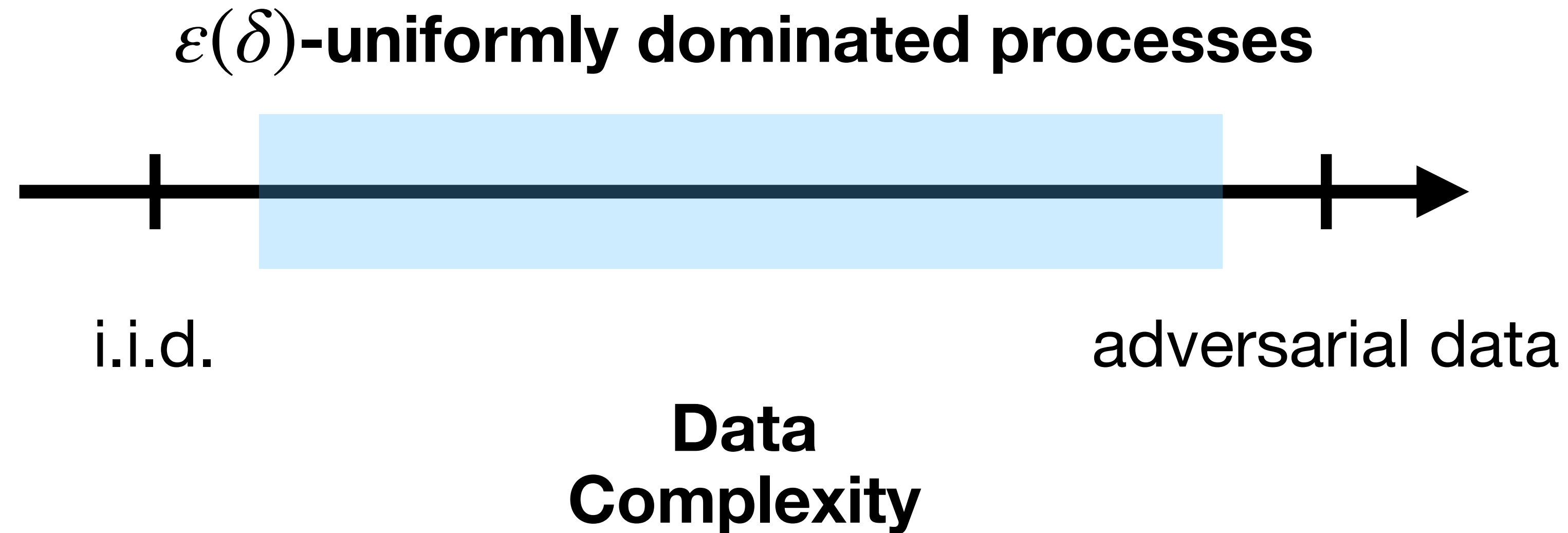
A data process is **L -Lipschitz** or is **L -smoothed** whenever

$$\forall n \in \mathbb{N}, \quad \mu(A) < \delta \implies \Pr(X_n \in A \mid \mathcal{F}_{n-1}) < L \cdot \delta,$$

whenever $A \subset \mathcal{X}$ is measurable.

The Radon-Nikodym derivative is bounded $\frac{d\nu_n}{d\mu} \leq L$.

Interpolation between i.i.d. and worst-case



Uniform convergence for smoothed data

Let $\mathcal{X} = [0,1]$ with the Lebesgue measure.



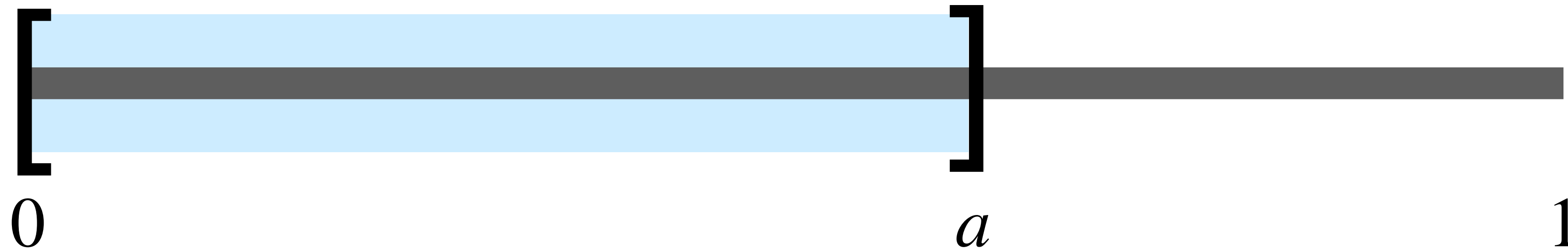
0

1

Uniform convergence for smoothed data

Let $\mathcal{X} = [0,1]$ with the Lebesgue measure. Let $\mathbf{I}$ be the set of thresholds:

$$\mathbf{I} = \{[0,a] : a \in [0,1]\}.$$



Uniform convergence for smoothed data

Let $\mathcal{X} = [0,1]$ with the Lebesgue measure. Let $\mathbf{I}$ be the set of thresholds:

$$\mathbf{I} = \{[0,a] : a \in [0,1]\}.$$

Informal Lemma (Bhattacharjee, Dasgupta, So 2025)

Let $(X_1, Y_1, \dots, X_N, Y_N)$ be adaptively generated by an L -smoothed process.

Uniform convergence for smoothed data

Let $\mathcal{X} = [0,1]$ with the Lebesgue measure. Let $\mathbf{I}$ be the set of thresholds:

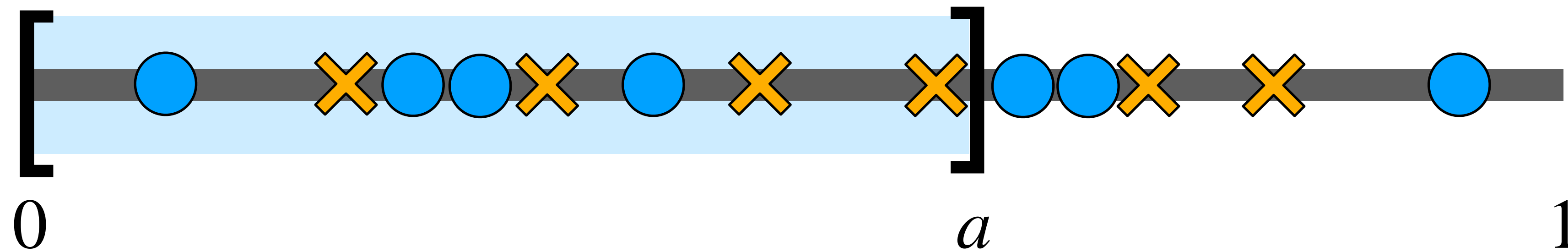
$$\mathbf{I} = \{[0,a] : a \in [0,1]\}.$$

Informal Lemma (Bhattacharjee, Dasgupta, So 2025)

Let $(X_1, Y_1, \dots, X_N, Y_N)$ be adaptively generated by an L -smoothed process. Then, uniform convergence over $\mathbf{I}$ holds with probability at least δ ,

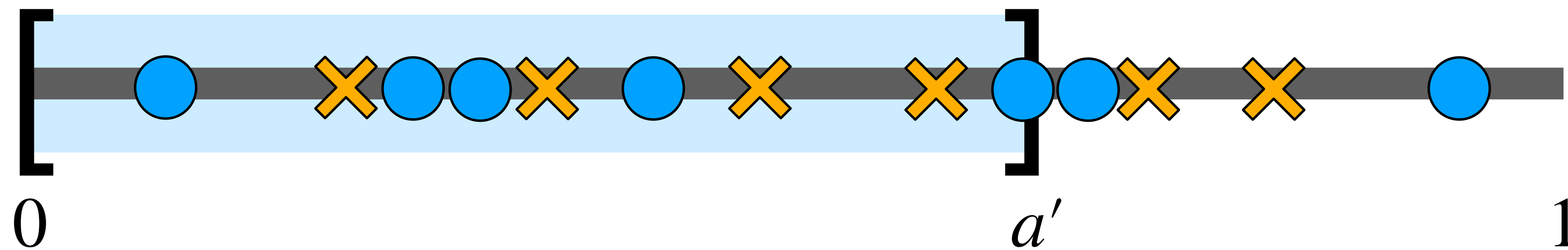
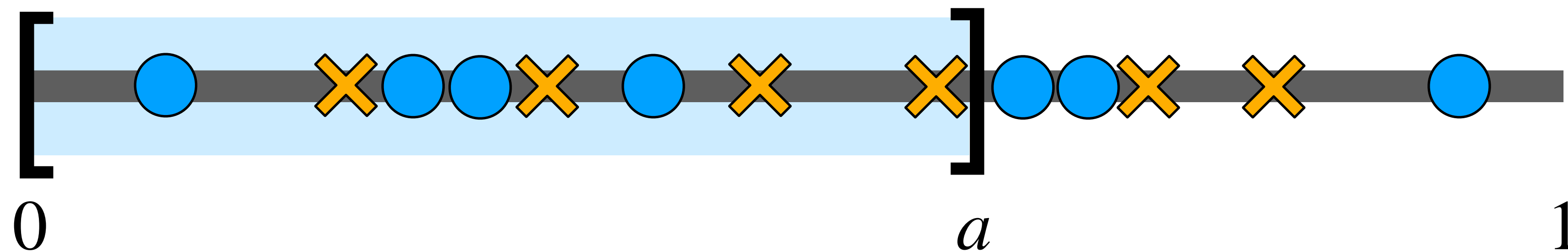
$$\forall I \in \mathbf{I}, \quad \left| \frac{1}{|I|} \sum_{X_i \in I} Y_i - \eta(X_i) \right| \lesssim \sqrt{\frac{1}{|I|} \log \frac{N}{\delta}}$$

Proof idea



The class of intervals we care about.

Proof idea



If two intervals differ by a small region...

Proof idea

...it is unlikely that their symmetric difference Δ contains many data points.



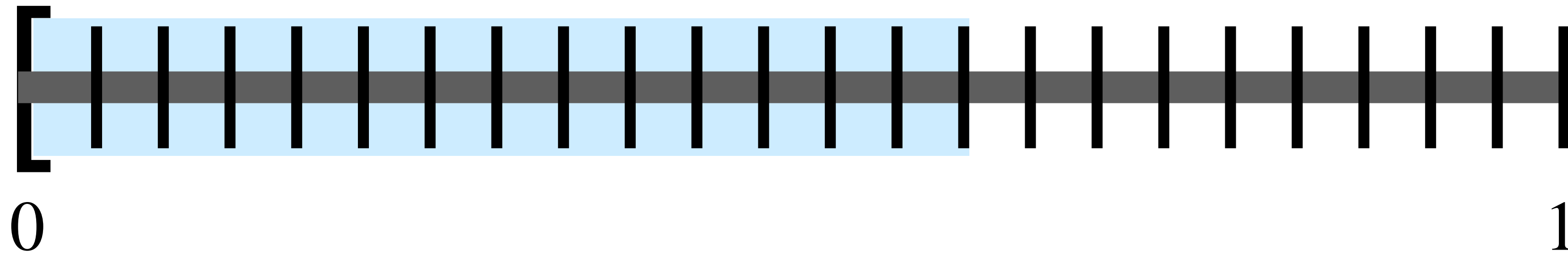
Proof idea

...it is unlikely that their symmetric difference Δ contains many data points.

- ▶ The expected number of data points in Δ is no more than $L \cdot \mu(\Delta)$.



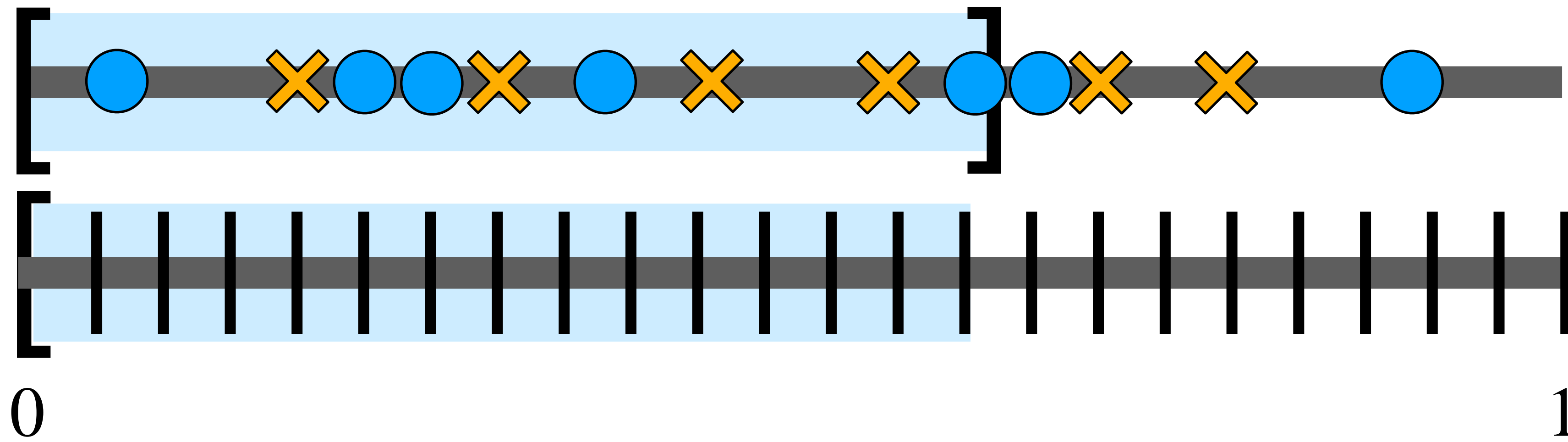
Proof idea



Let $\mathbf{J} \subset \mathbf{I}$ be the subclass:

$$\mathbf{J} = \left\{ \left[0, \frac{i}{LN} \right] : i = 1, 2, \dots, \lceil LN \rceil \right\}.$$

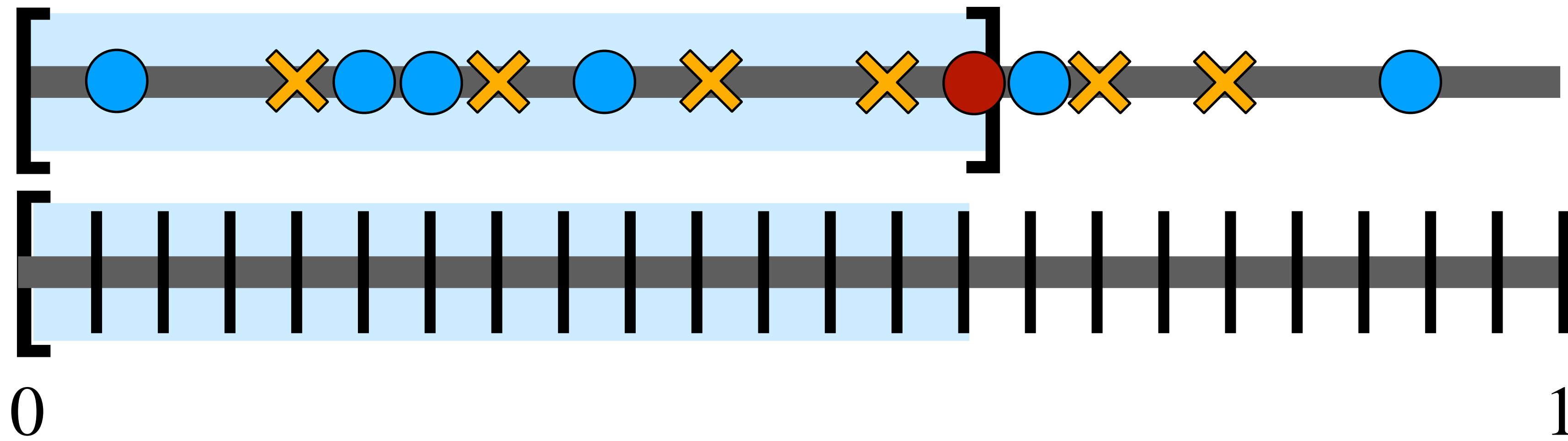
Proof idea



Then, $\mathbf{J}$ is a finite subclass that approximates $\mathbf{I}$:

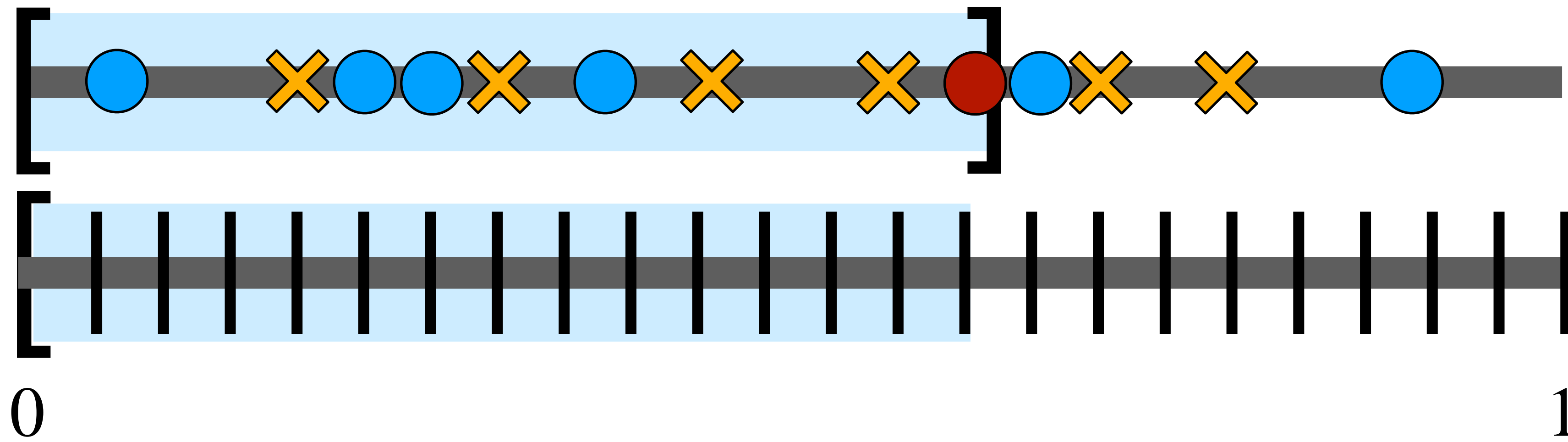
$$\forall I \in \mathbf{I} \quad \min_{J \in \mathbf{J}} \mu(I \Delta J) < \frac{1}{LN}.$$

Proof idea



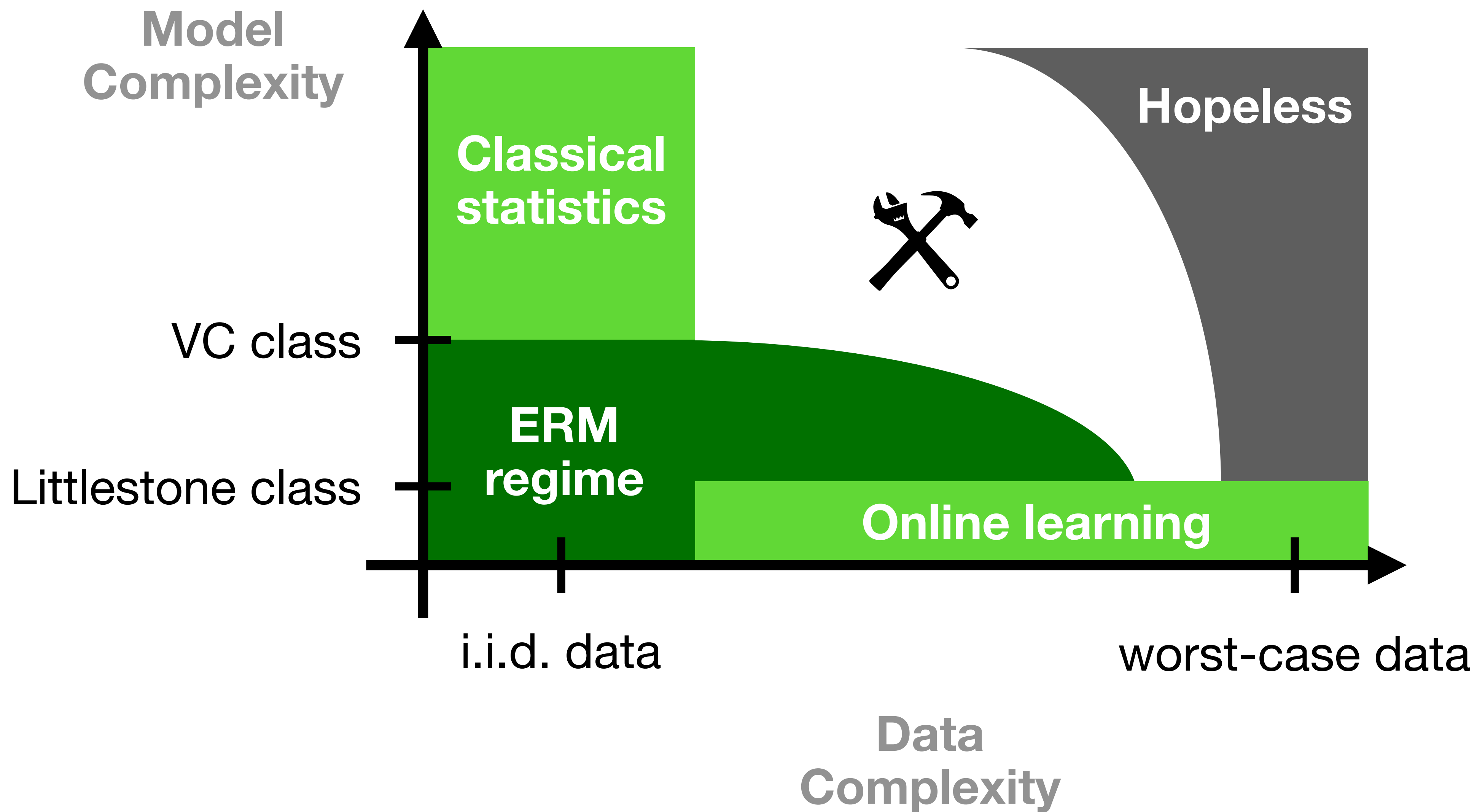
If $\mu(I\Delta J) < (LN)^{-1}$, then by L -smoothness, the expected number of data points that land in $I\Delta J$ is **less than one**.

Proof idea



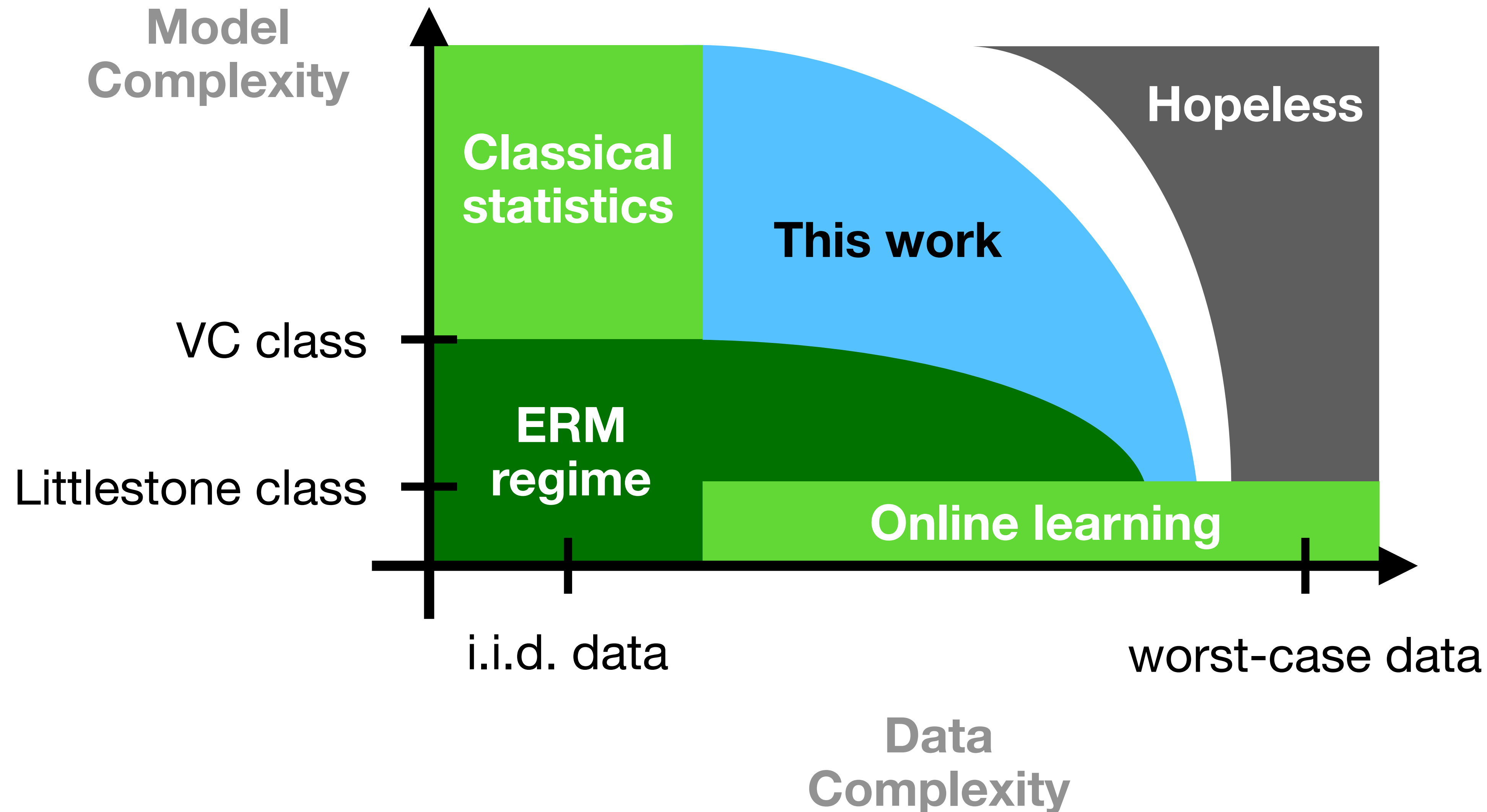
A slightly worse uniform convergence over I can be deduced from uniform convergence over J .





[Nearest neighbor methods]

Simple algorithms can learn complex things



Consistency of the k_n -nearest neighbor rule

Let $\mathcal{X} = \mathbb{R}^d$ be Euclidean space. Use the k_n -nearest neighbor learning.

Theorem (Bhattacharjee, Dasgupta, So 2025).

Let $(k_n)_n$ be a regular* sequence. Let the adaptive data process be uniformly dominated by the Lebesgue measure and $\eta : \mathcal{X} \rightarrow [0,1]$ be measurable. Then:

$$\limsup_{N \rightarrow \infty} \frac{1}{N} \sum_{n=1}^N \mathbf{1}\{\hat{Y}_n \neq Y_n^\star\} = 0 \quad \text{a.s.}$$

Challenges for modern machine learning



- Data-generating processes are black-boxes

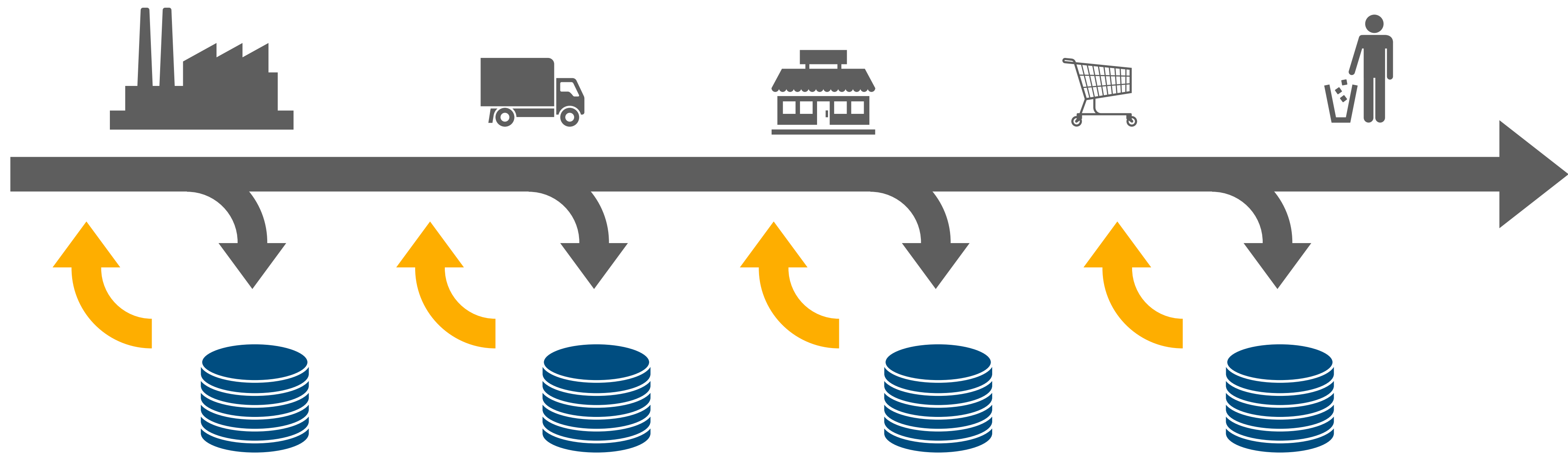
Challenges for modern machine learning

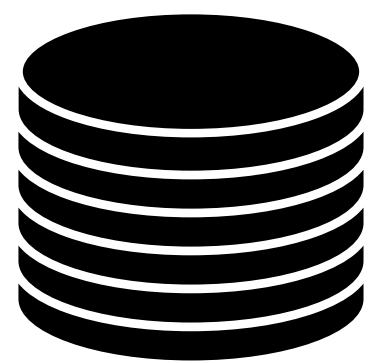


- Solving complex tasks requires a division of labor
 - The “learner” is a (coordinated?) network of adapting agents

RESEARCH OVERVIEW: LEARNER

Economics and Dynamics of Adaptive Agents





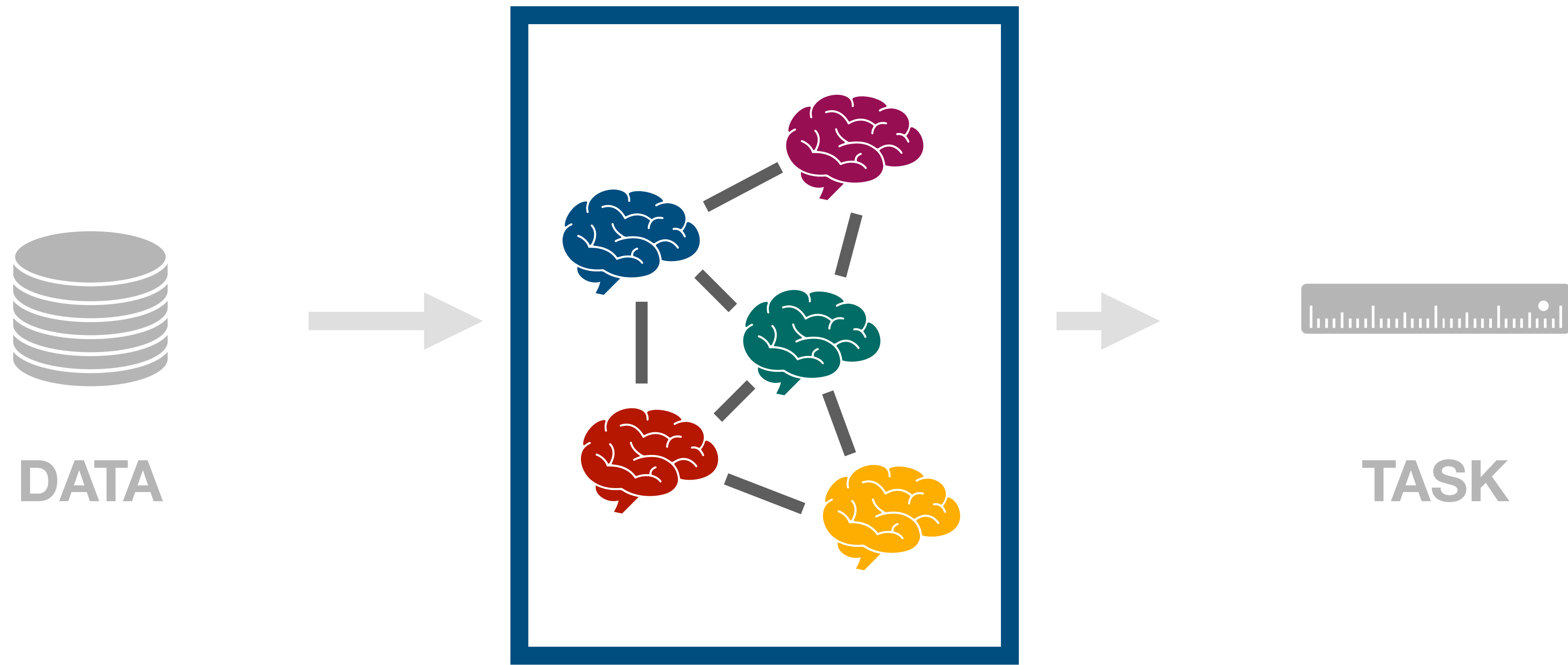
DATA



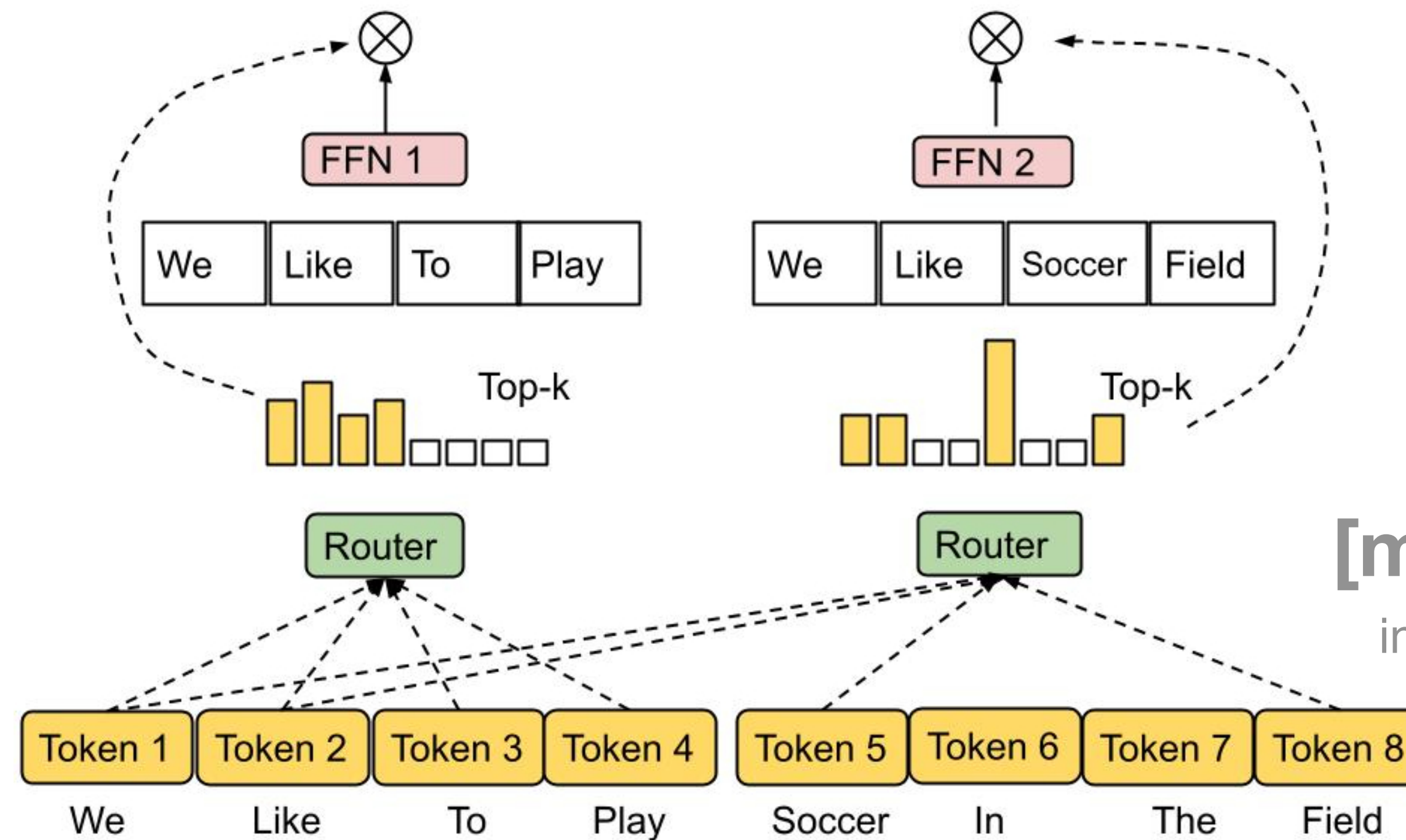
LEARNER



TASK



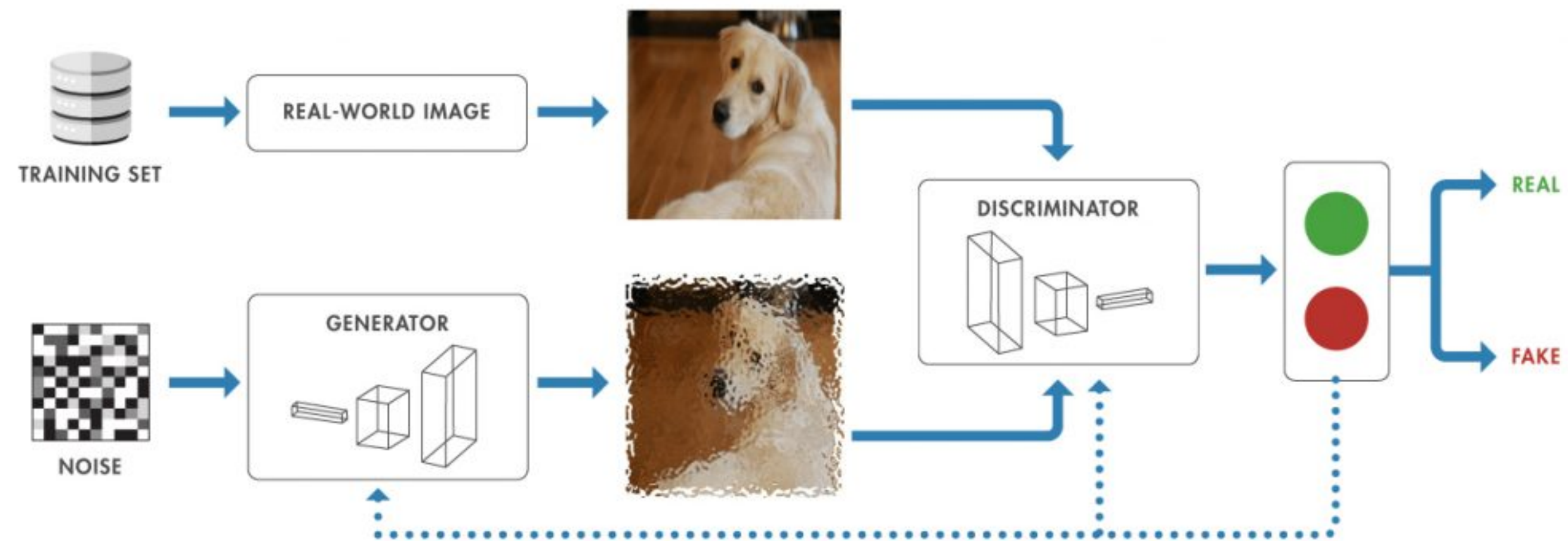
Complex problems often require a division of labor or decentralization.



[mixture of experts]

image from Google Research

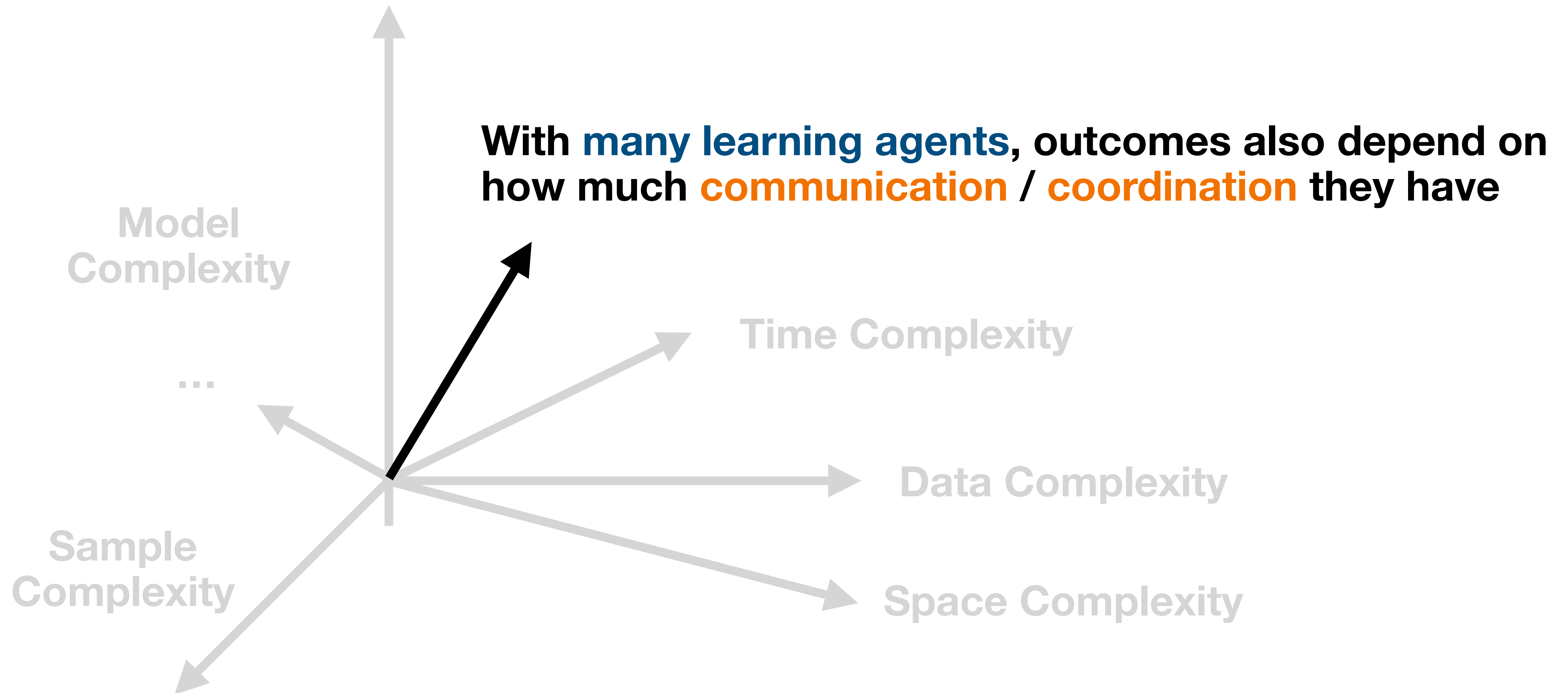
Complex problems often require a division of labor or decentralization.



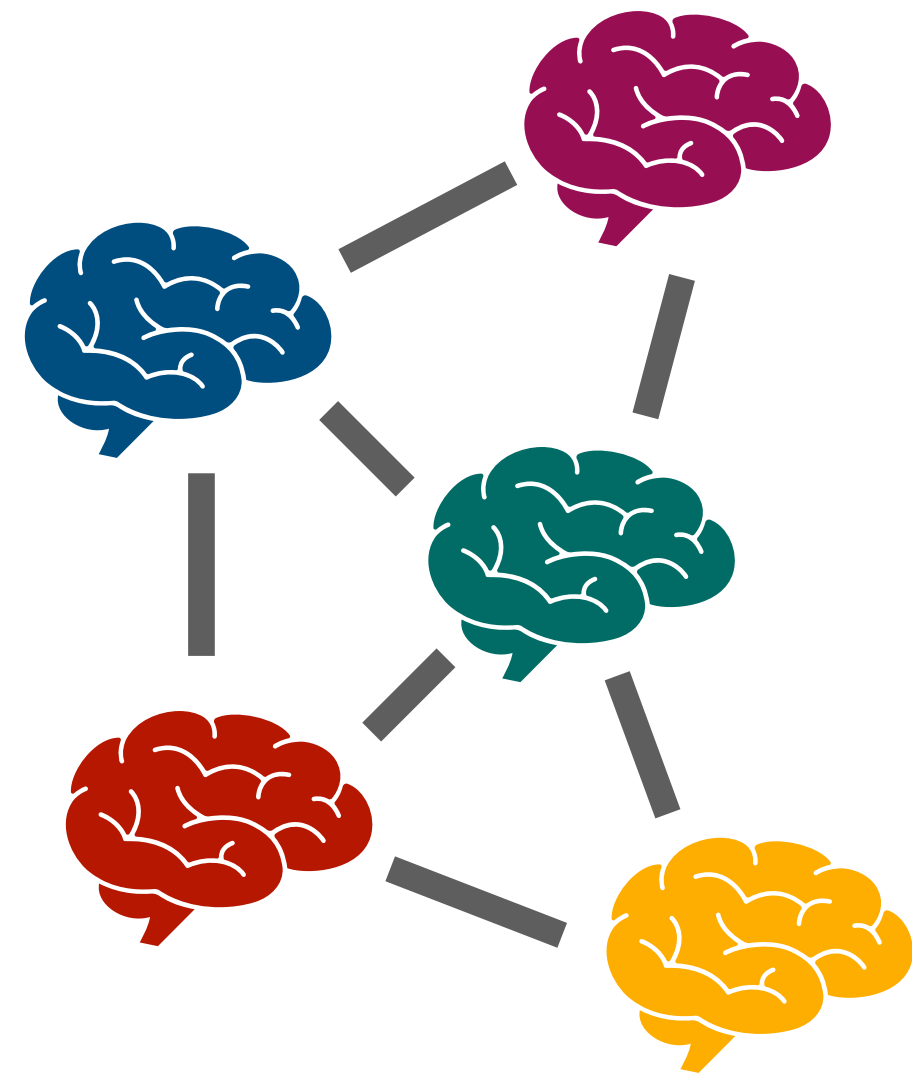
[GANs]
image from MATLAB

The solutions some learning problems can be formulated in terms of a game.

Some dimensions of what makes learning hard

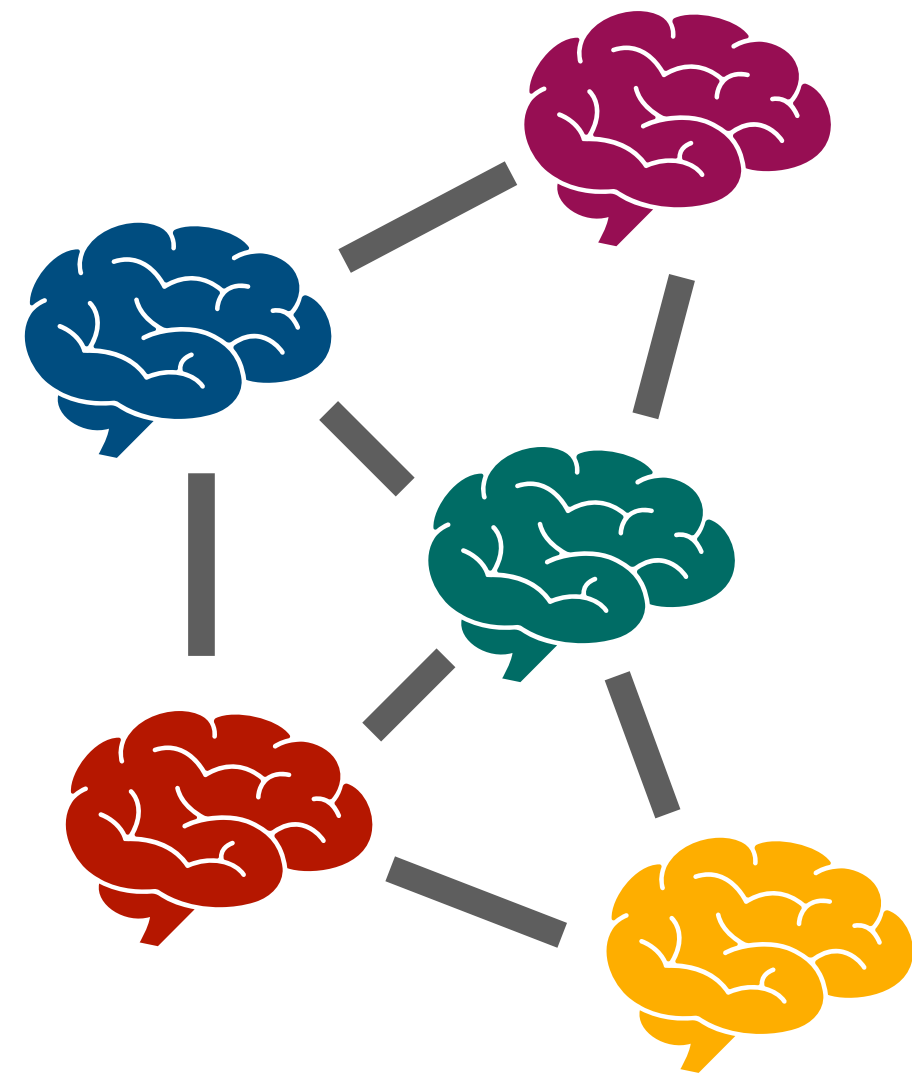


The price of anarchy



Decentralization can **introduce constraints** on communication and coordination

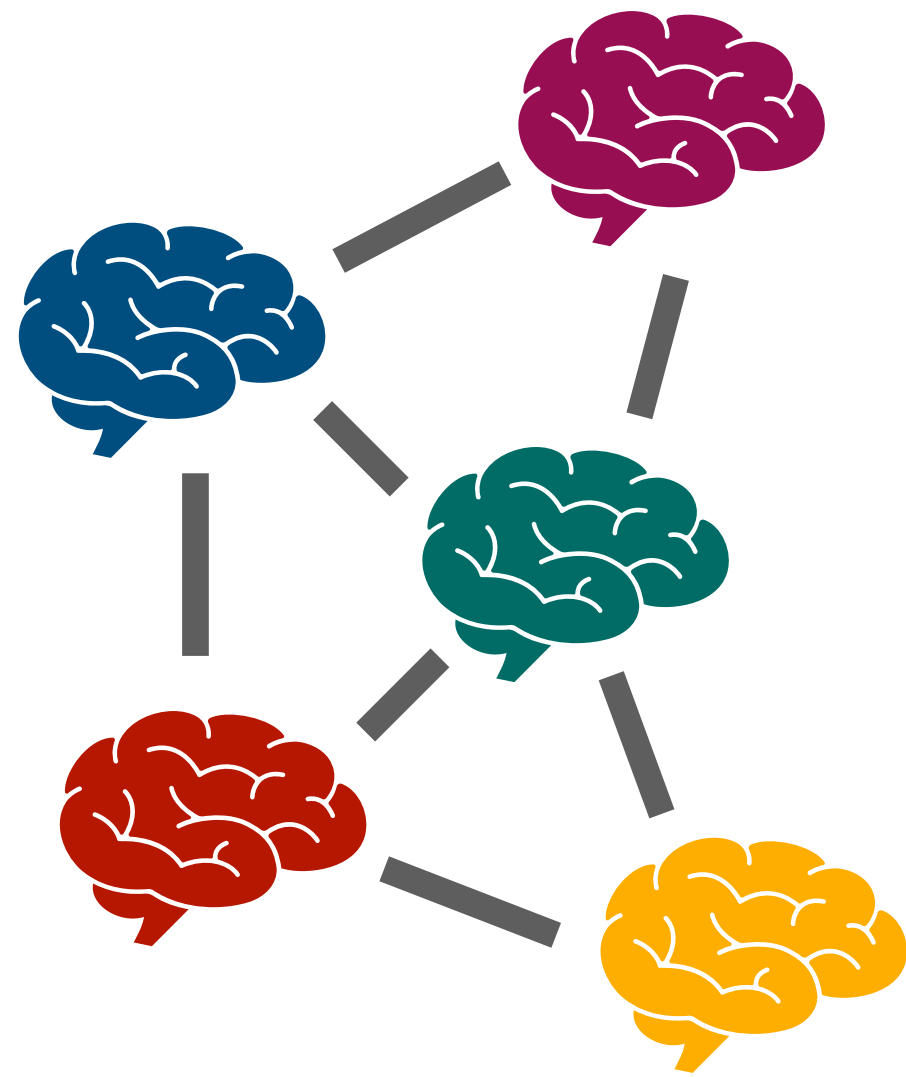
The price of anarchy



Decentralization can introduce constraints on communication and coordination

➡ leading to situations where **everyone loses.**

The price of anarchy

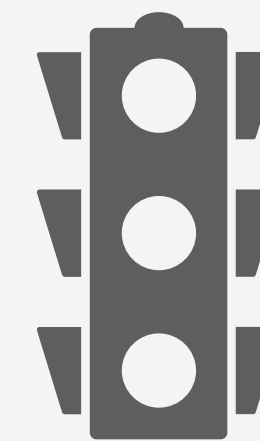


Decentralization can introduce constraints on communication and coordination

➔ **leading to situations where everyone loses.**

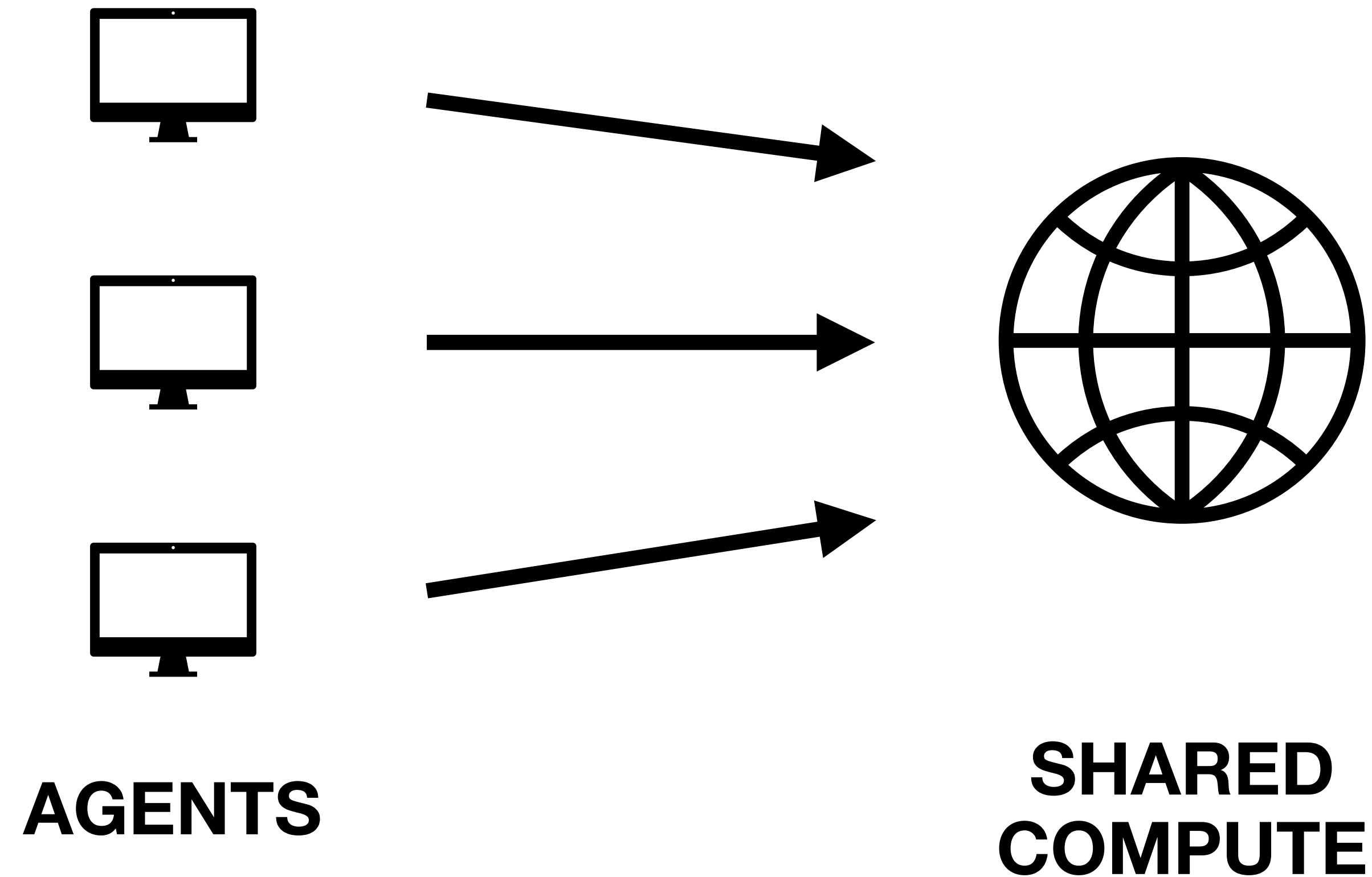


Prisoner's Dilemma

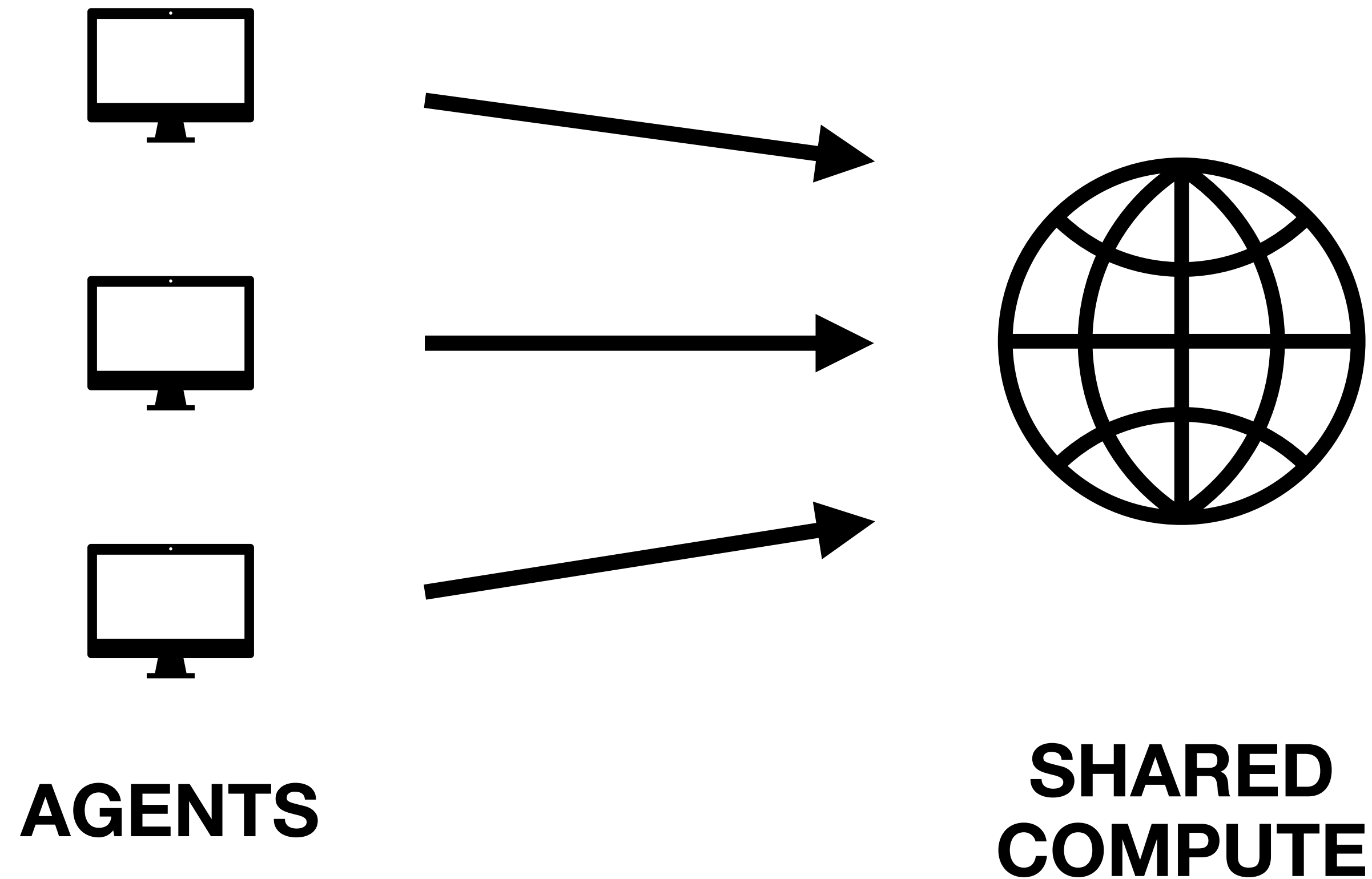


Tragedy of the Commons

Example: Tragedy of the Commons



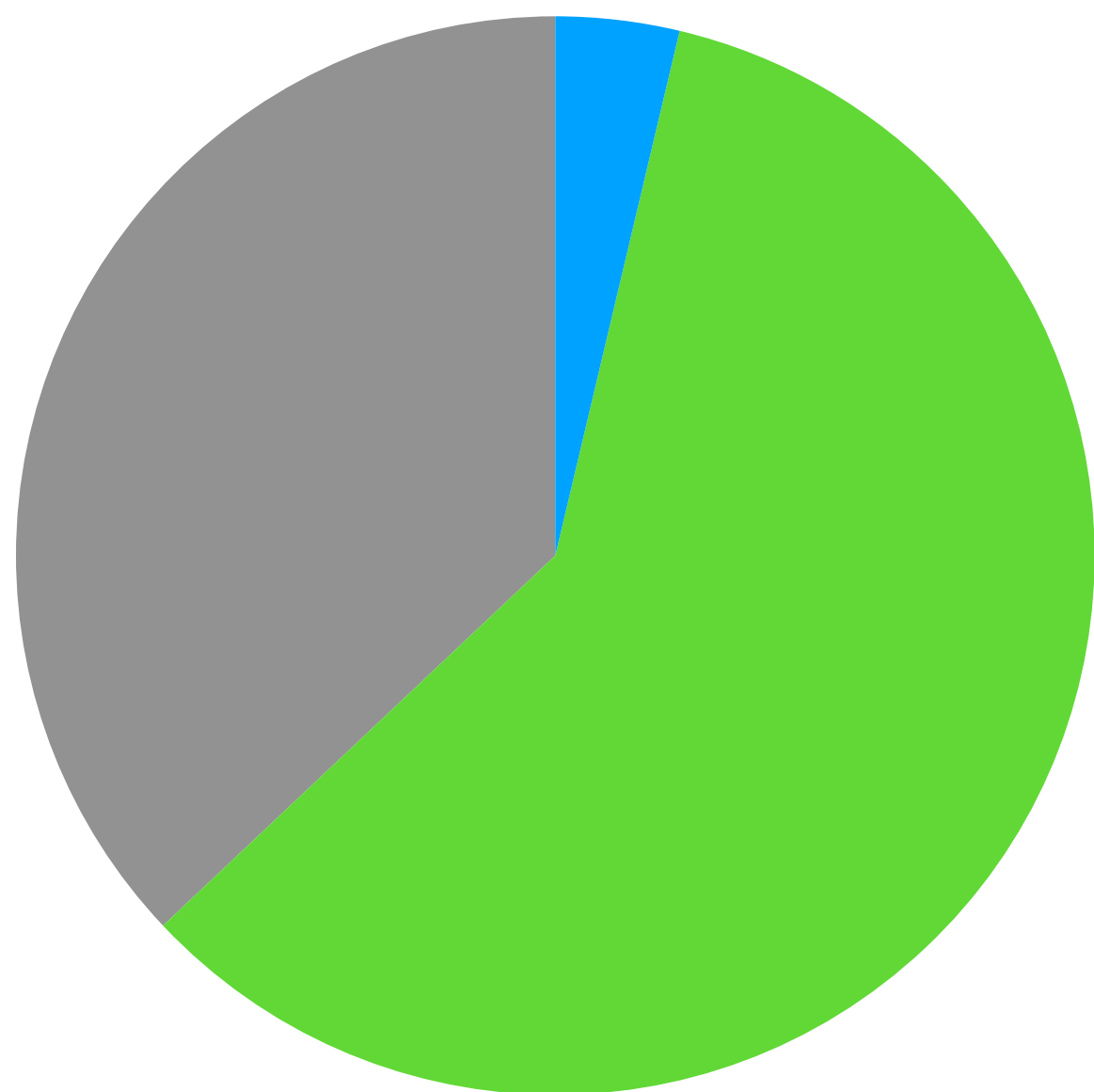
Example: Tragedy of the Commons



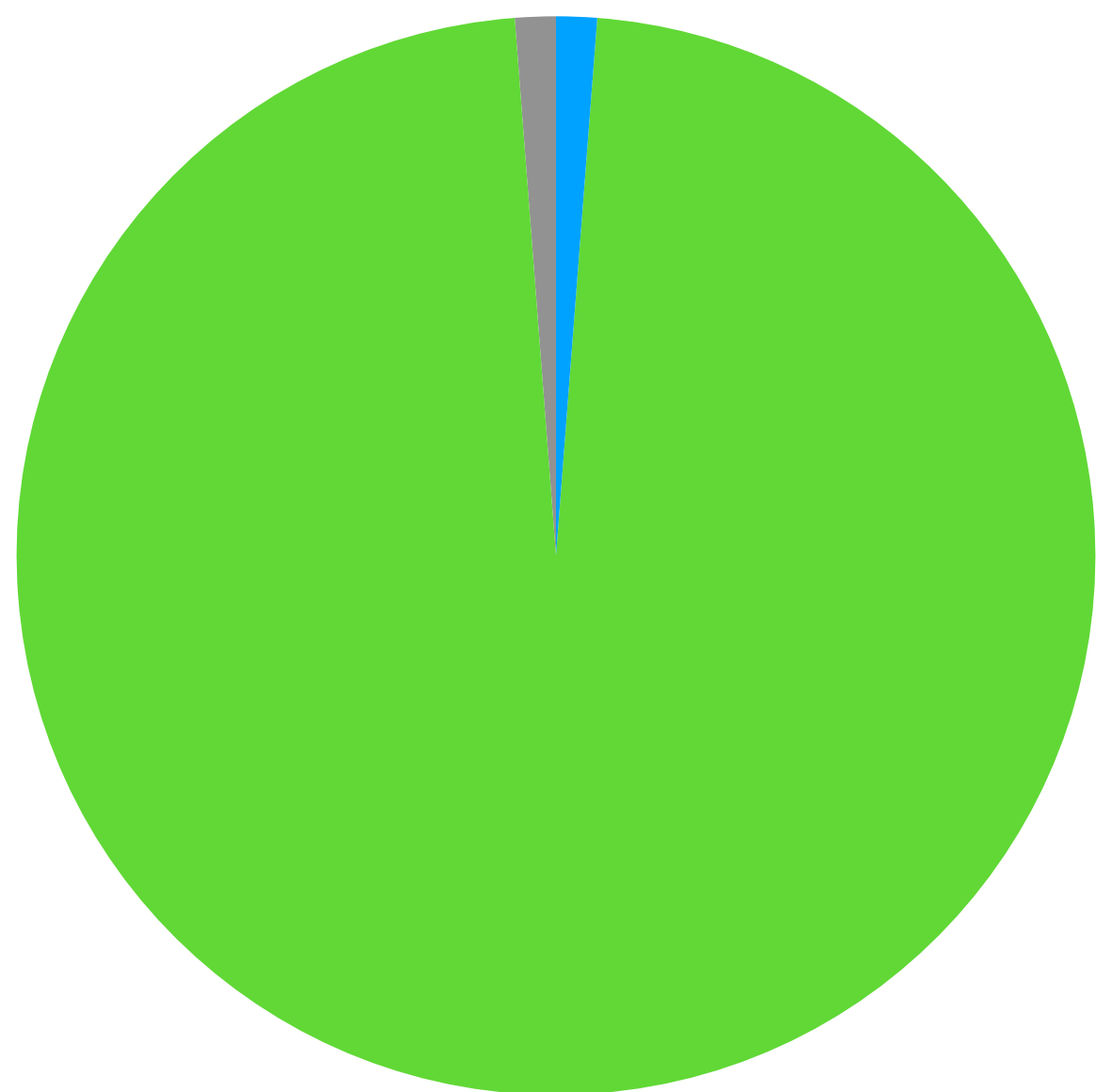
Question: How often should agents send jobs?

Three Possible Worlds

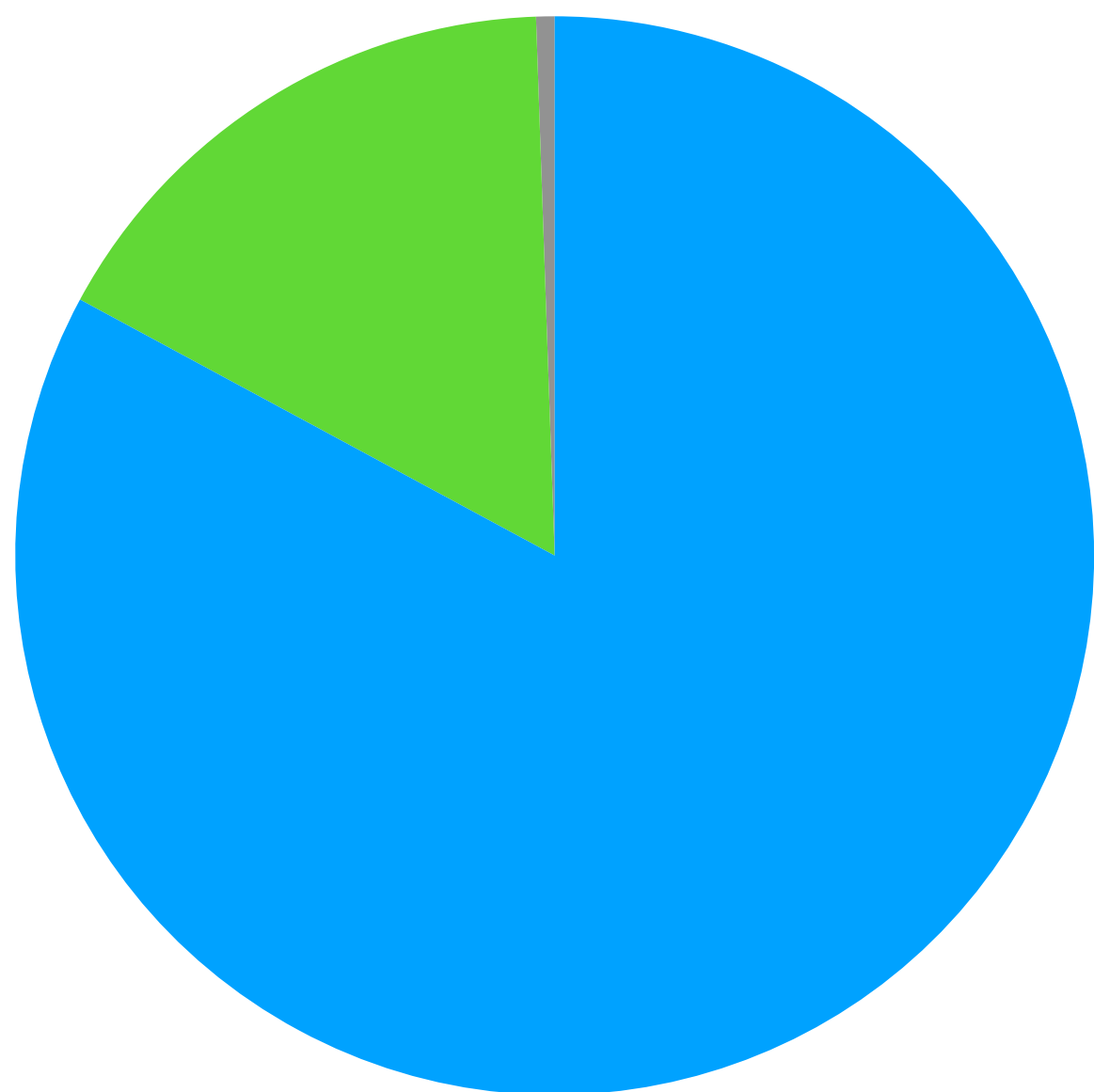
Under-utilization



Ideal load balancing

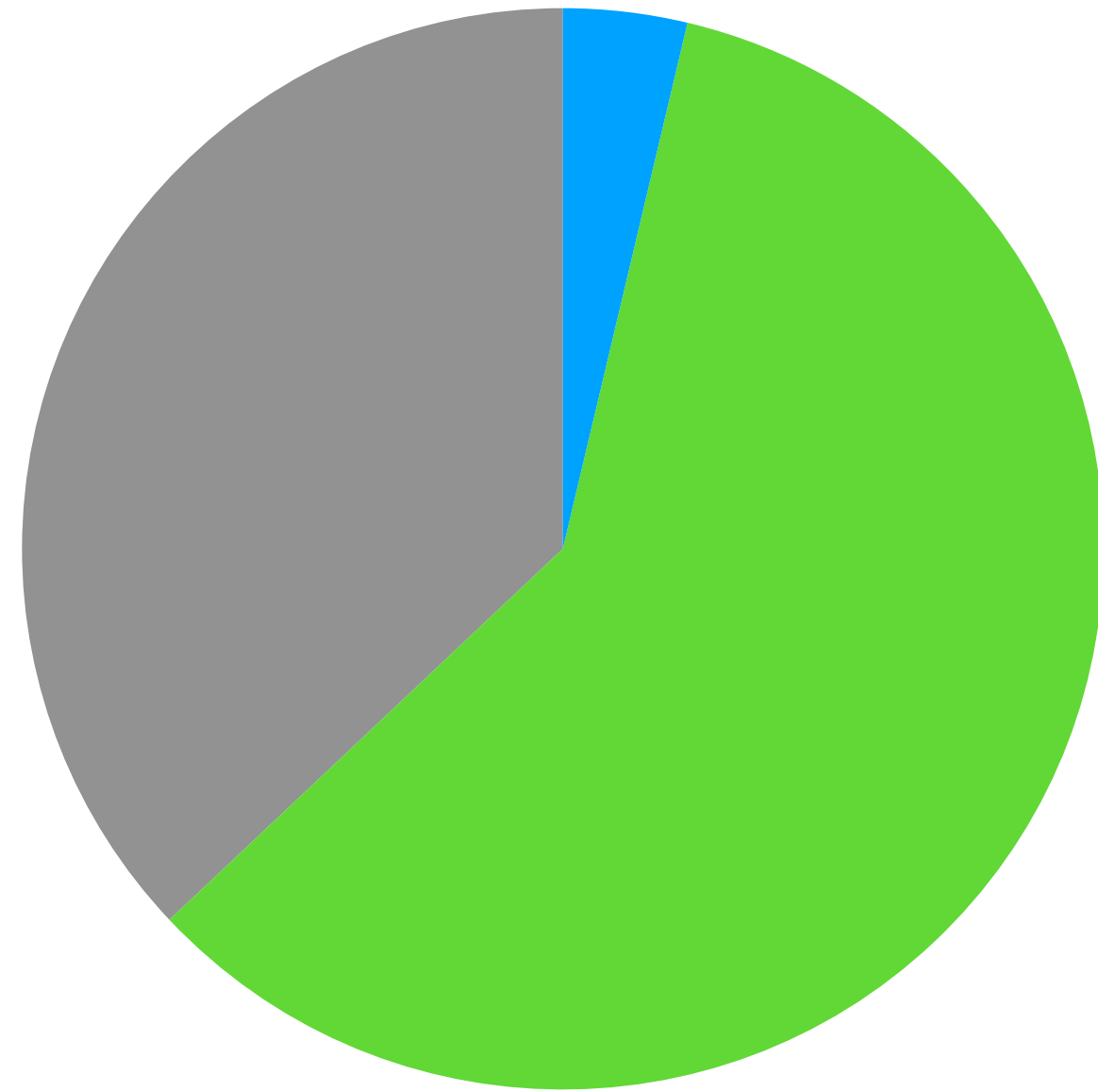


Wasted overhead

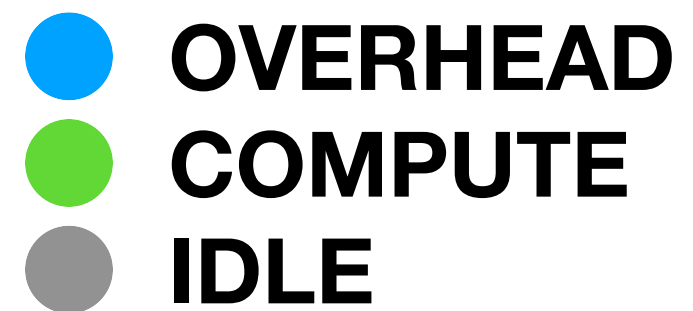


● OVERHEAD ● COMPUTE ● IDLE

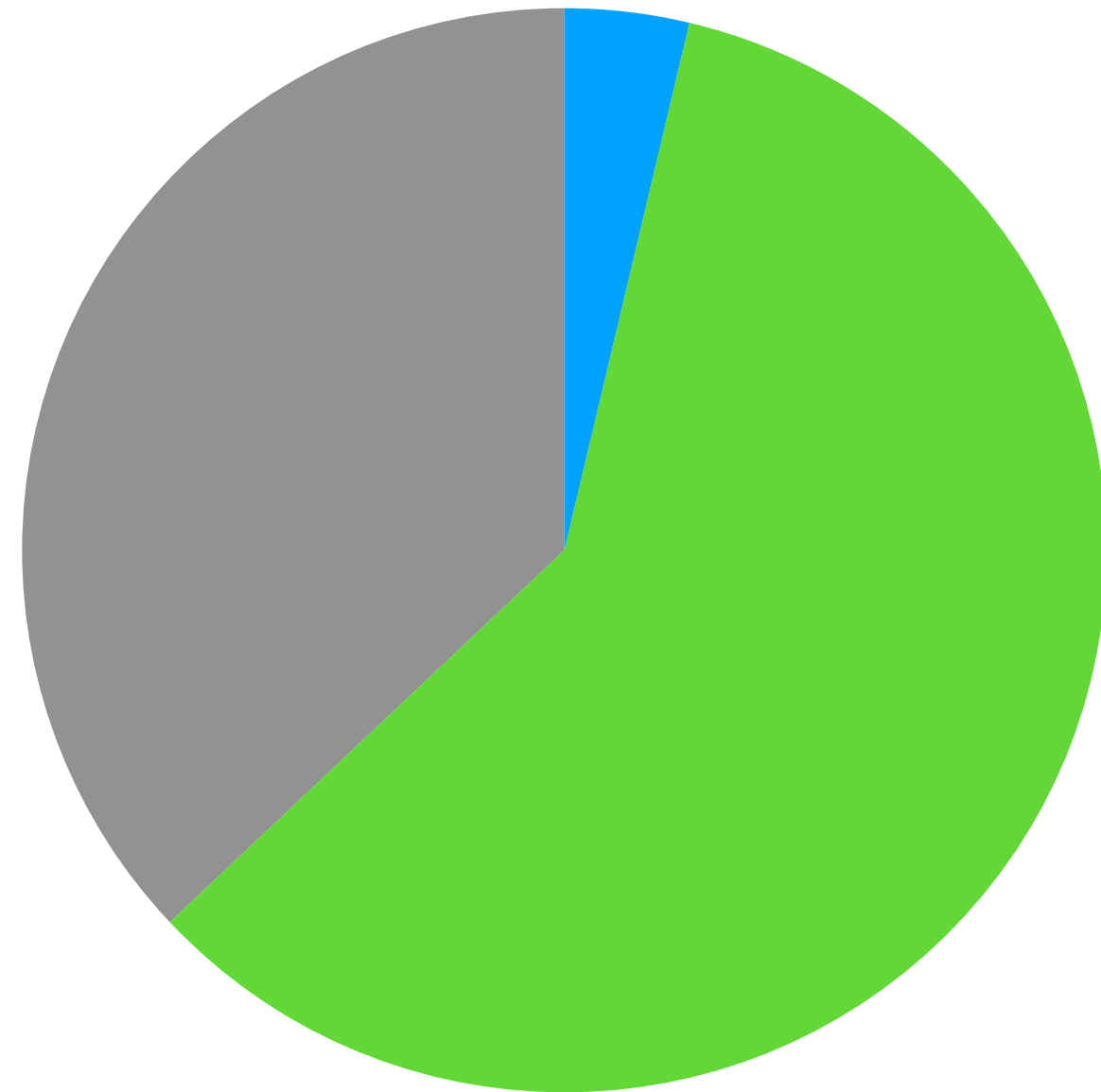
World One: Under-Utilization



Agents send jobs sporadically.

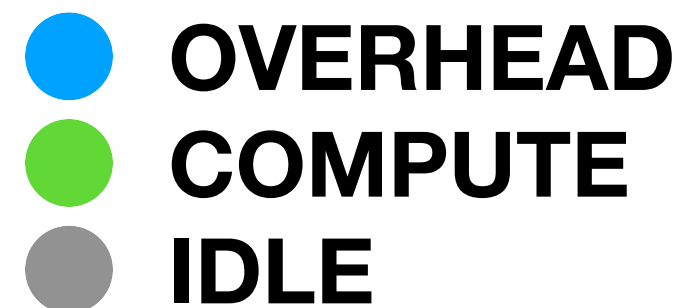


World One: Under-Utilization

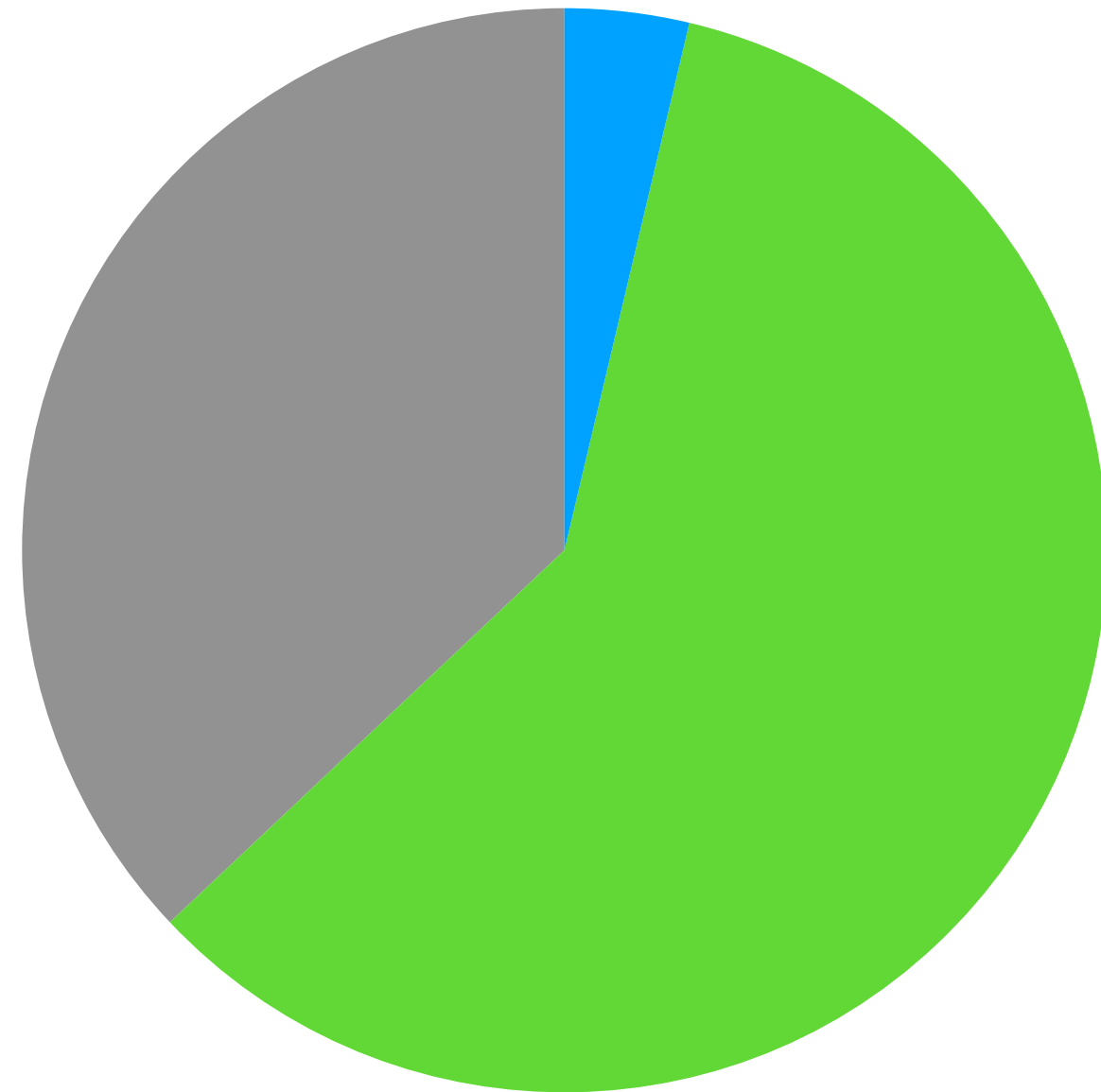


Agents send jobs sporadically.

A lot of idle time: **socially inefficient.**



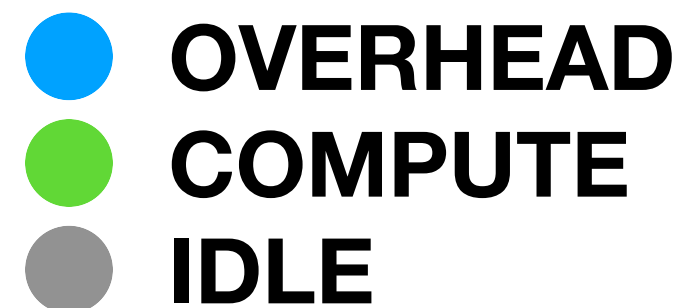
World One: Under-Utilization



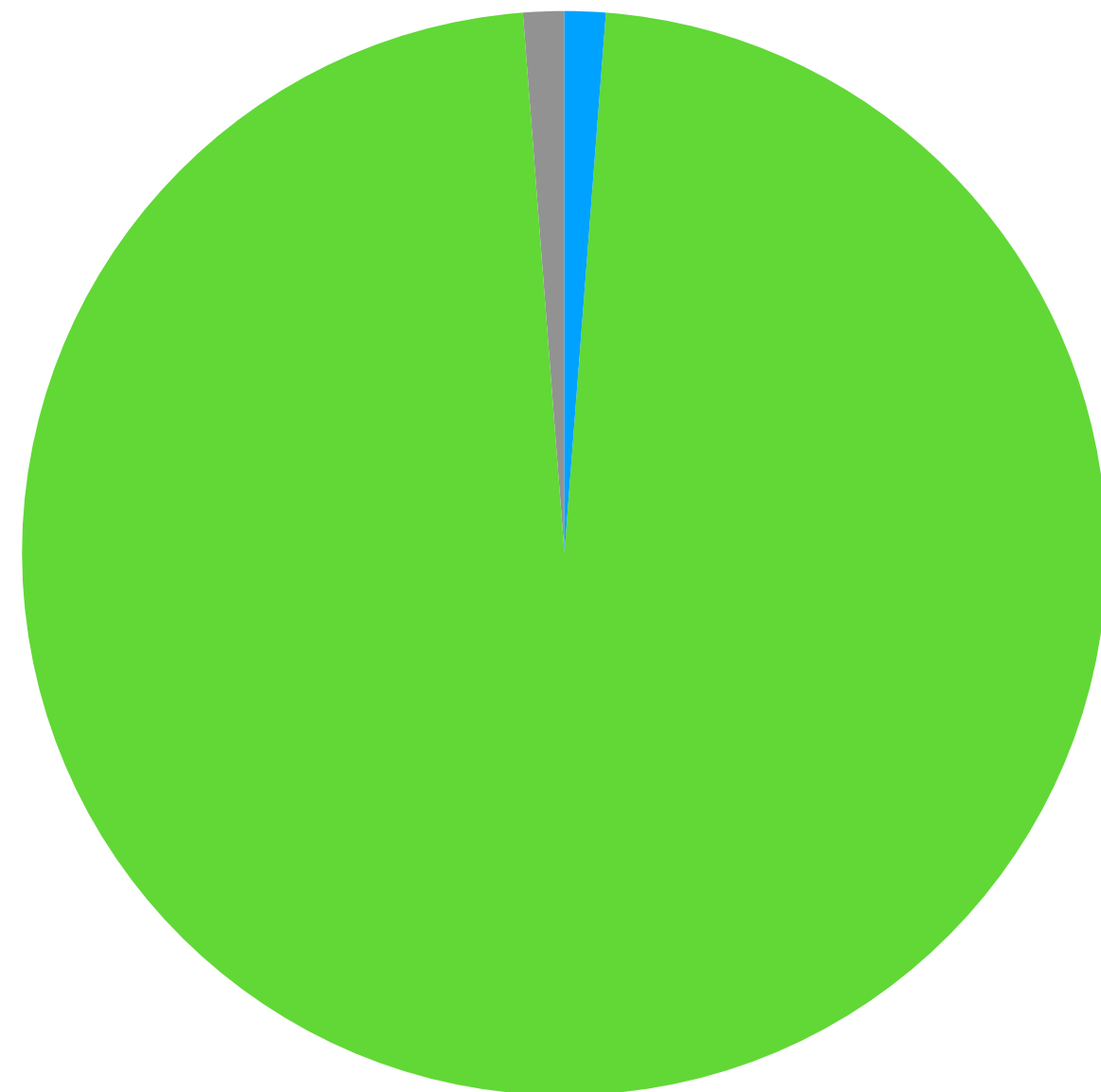
Agents send jobs sporadically.

A lot of idle time: **socially inefficient**.

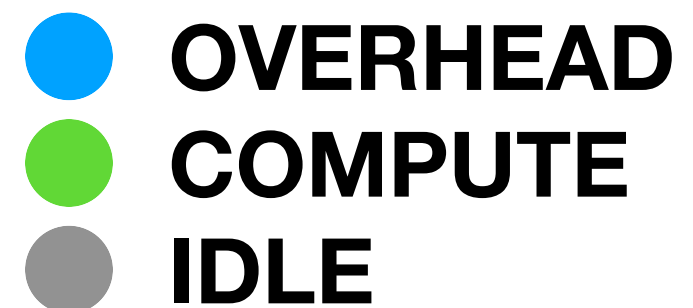
Everyone benefits from slightly increasing their utilization rate: **unstable**.



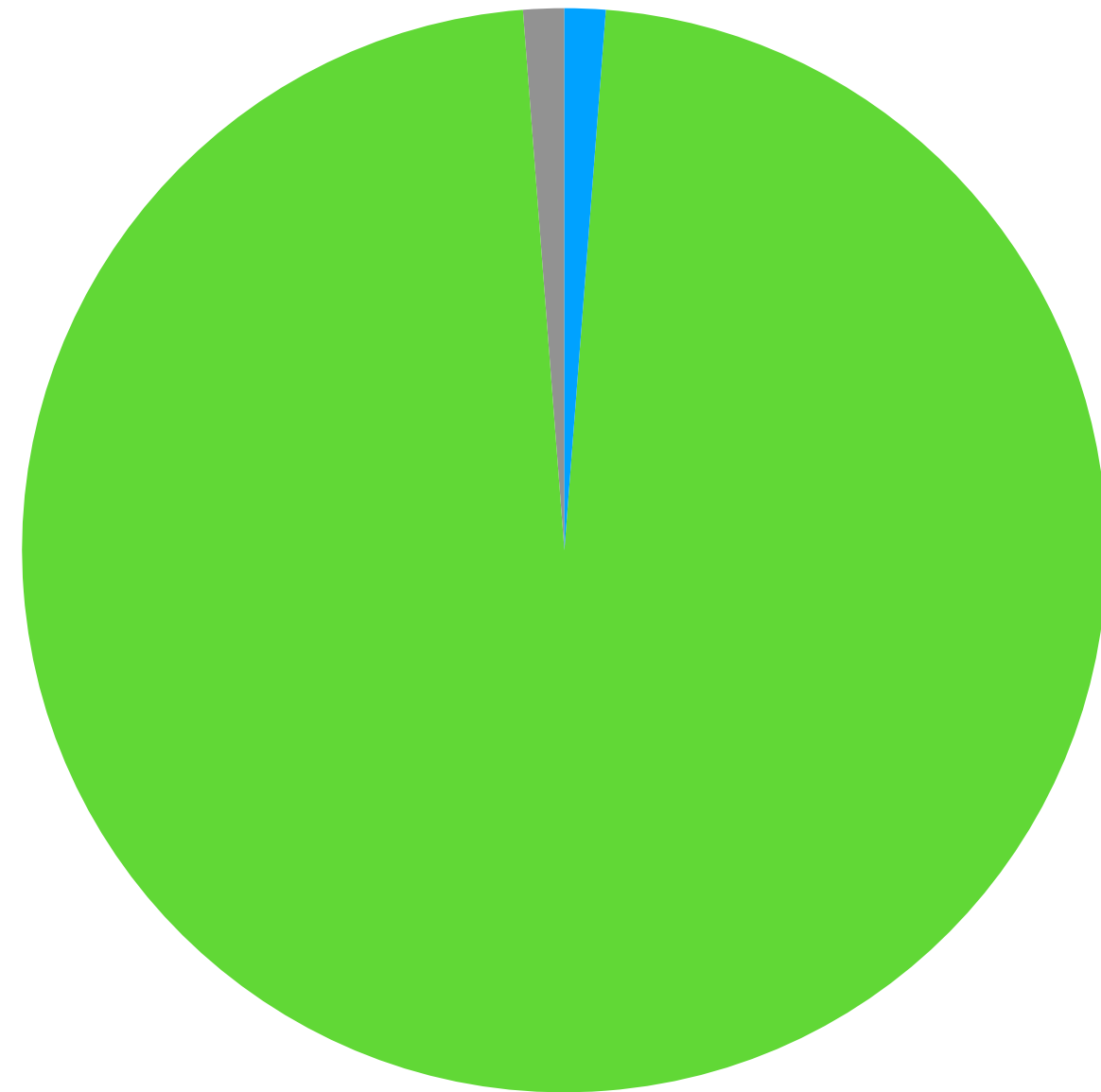
World Two: Perfect Utilization



Demand matches supply.

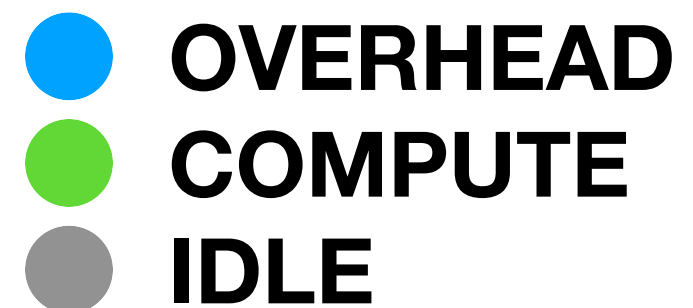


World Two: Perfect Utilization

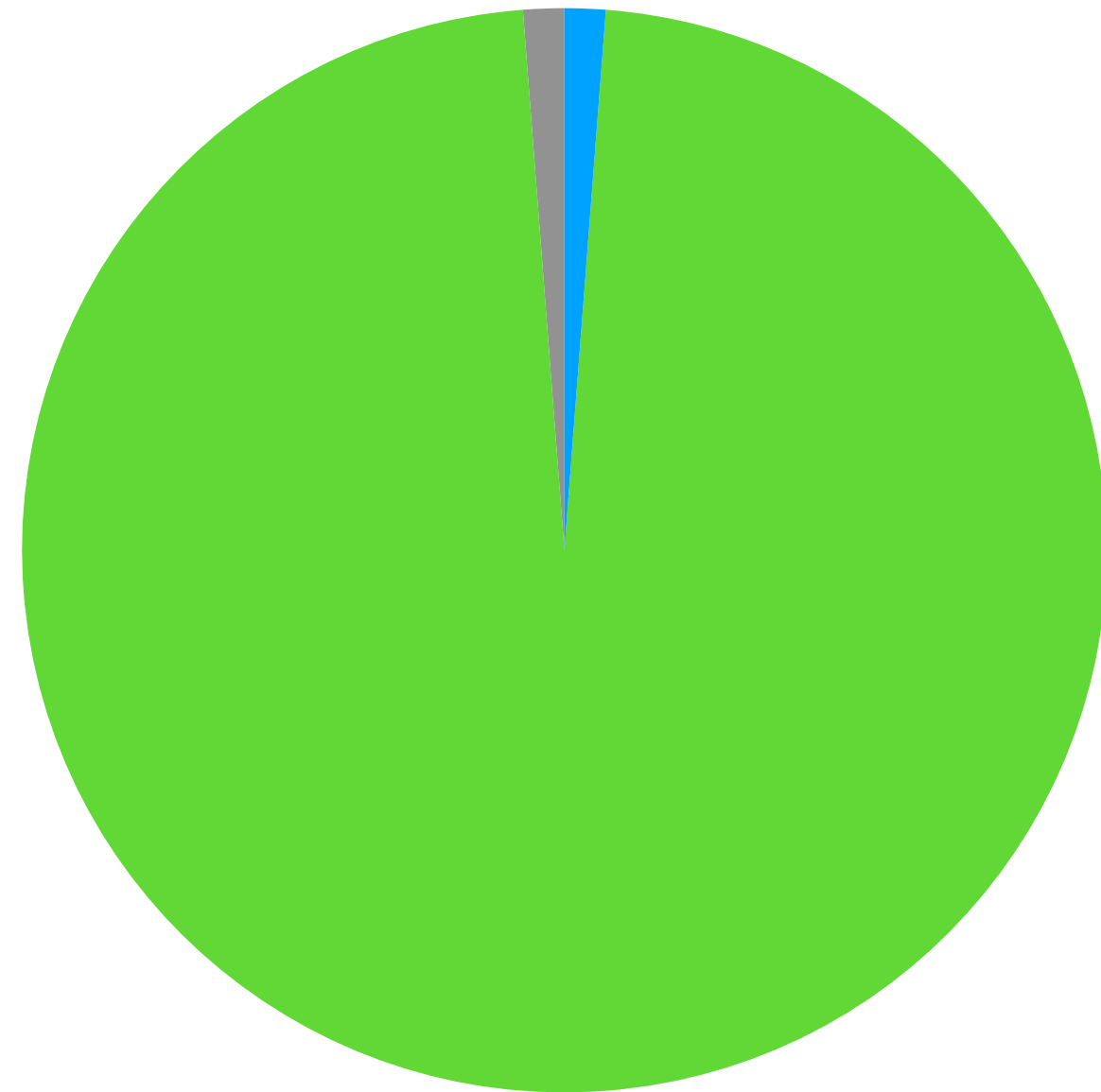


Demand matches supply.

No wasted compute: **socially optimal.**



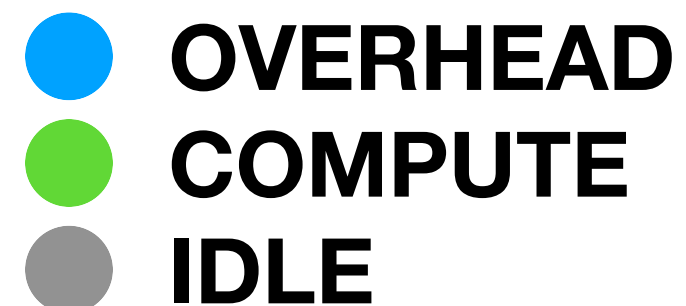
World Two: Perfect Utilization



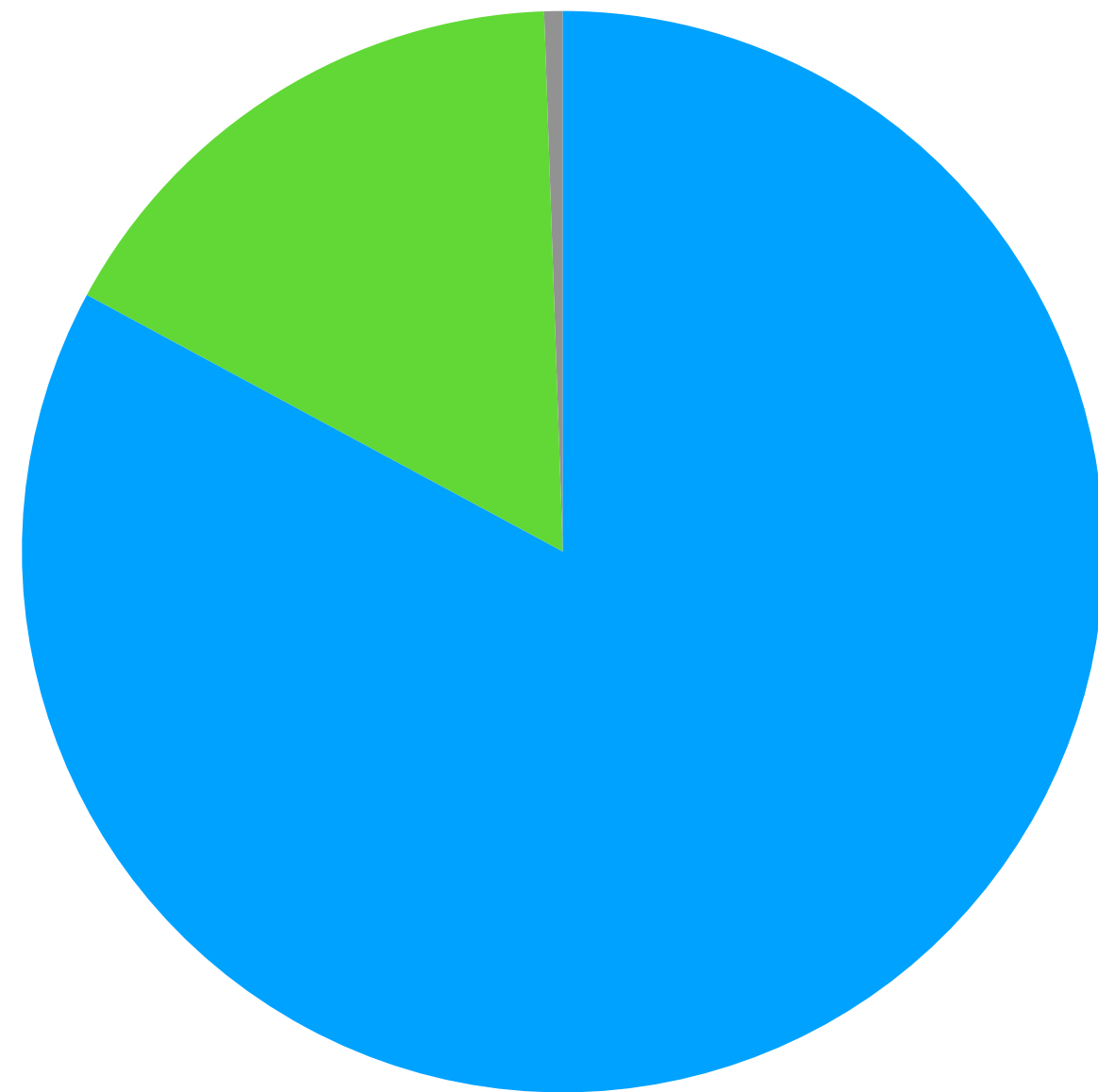
Demand matches supply.

No wasted compute: **socially optimal**.

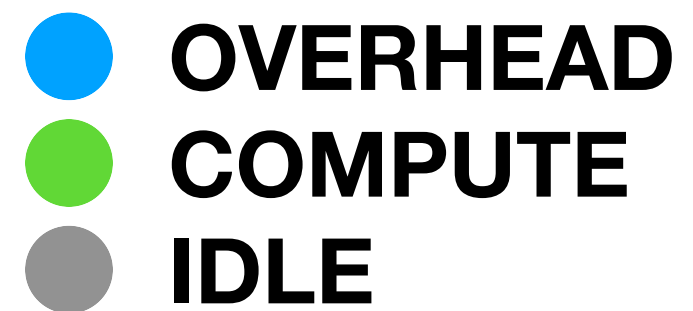
Individuals marginally benefit from increasing usage, as long as others do not: **unstable**.



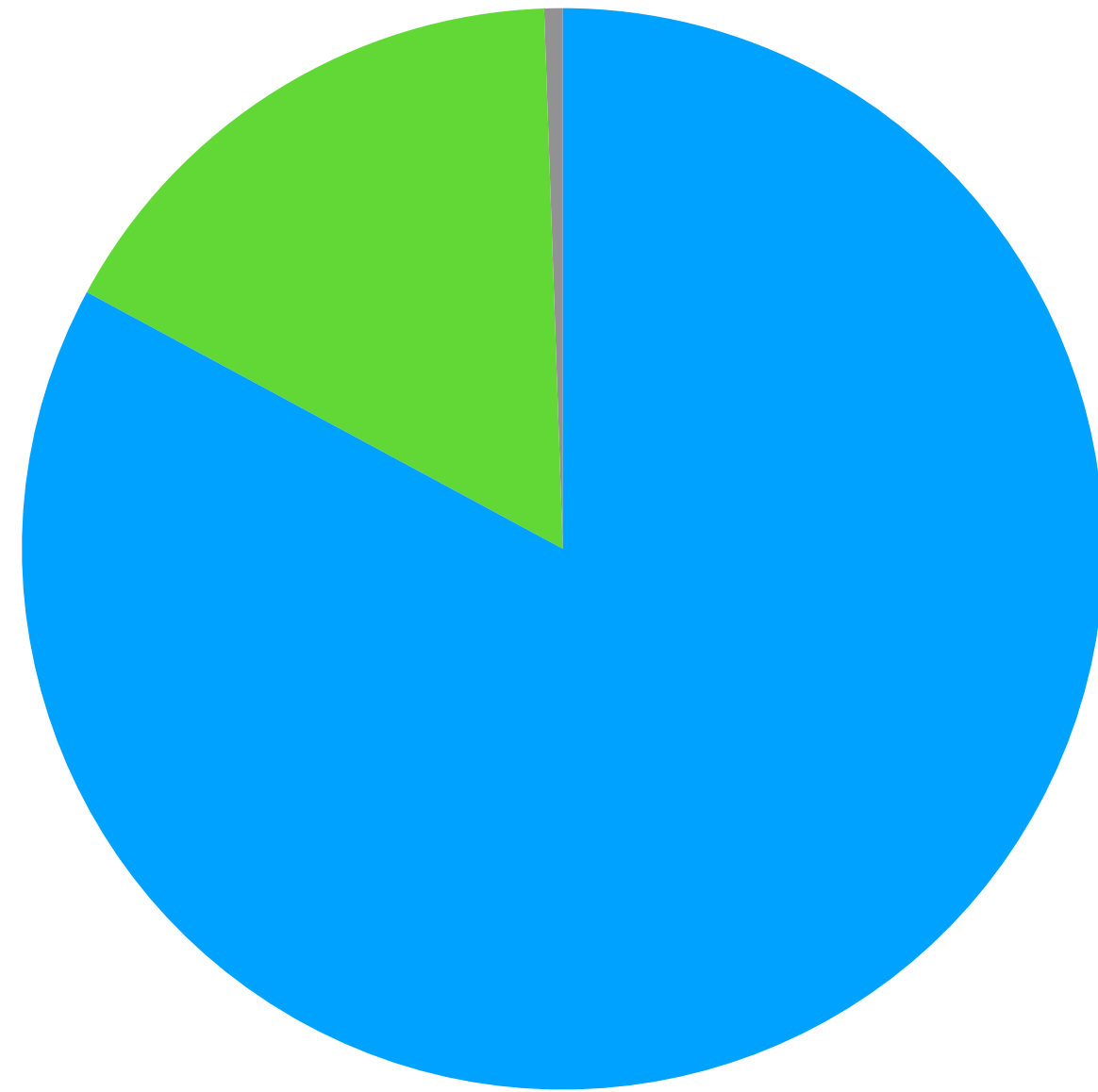
World Three: Over-Utilization



Agents overwhelm the server.

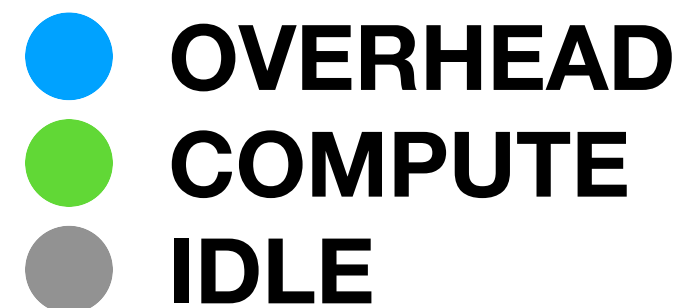


World Three: Over-Utilization

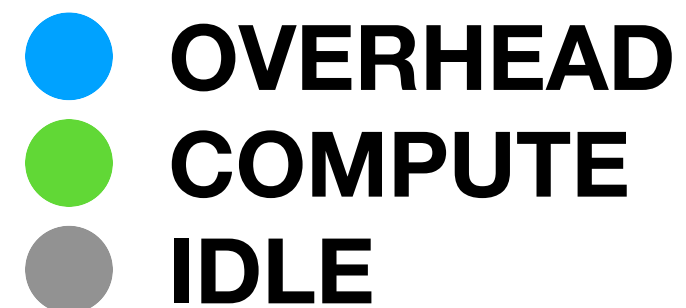
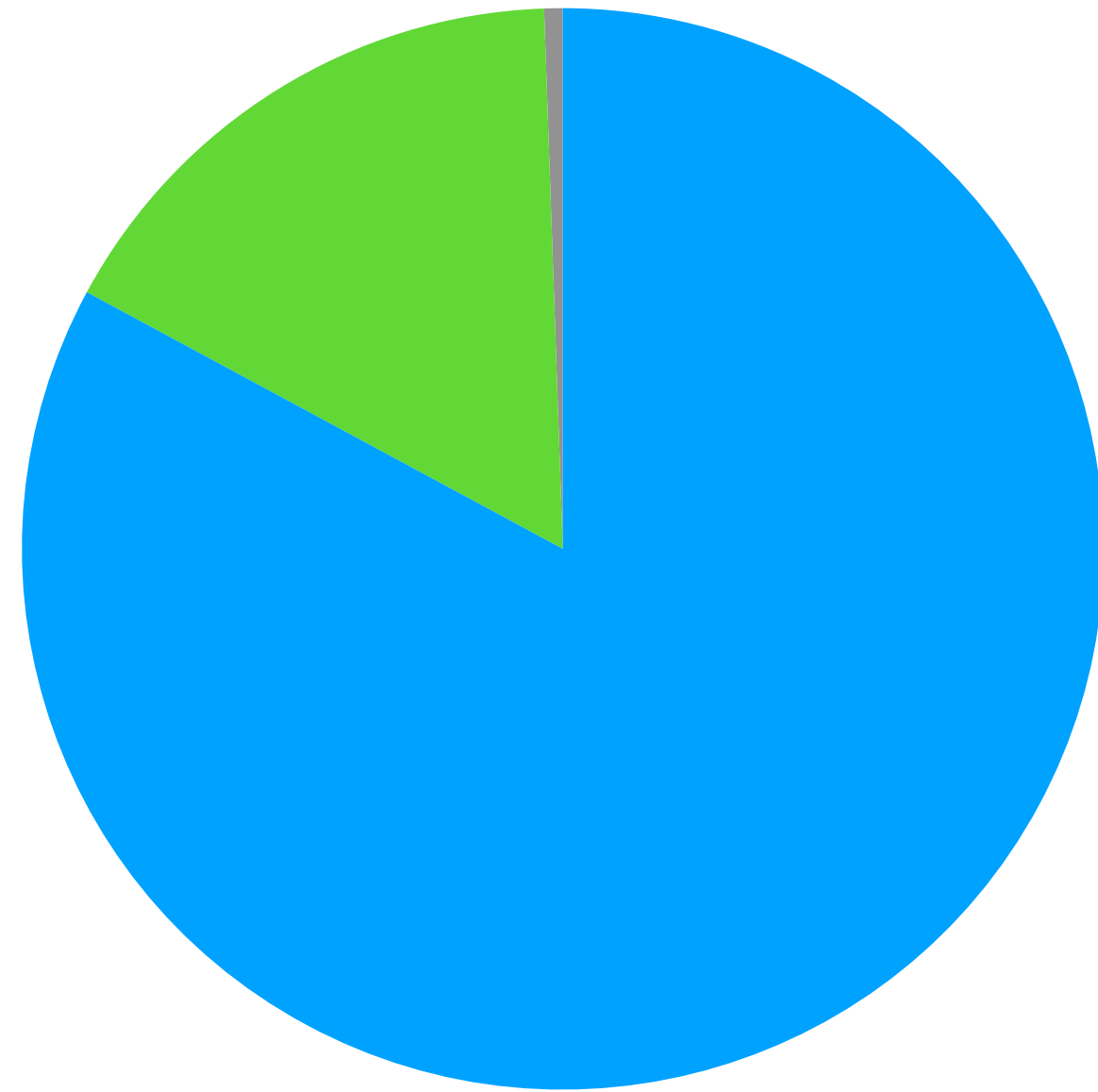


Agents overwhelm the server.

Leads to thrashing—most of time is spent in context switching: very **socially inefficient**.



World Three: Over-Utilization



Agents overwhelm the server.

Leads to thrashing—most of time is spent in context switching: very **socially inefficient**.

The individual is marginally worse off by sending fewer jobs: **stable equilibrium**.

Three Possible Worlds

Under-utilization

Ideal load balancing

Wasted overhead

The Tragedy

Everyone is better off in **World Two**,
but only **World Three** is stable.



OVERHEAD

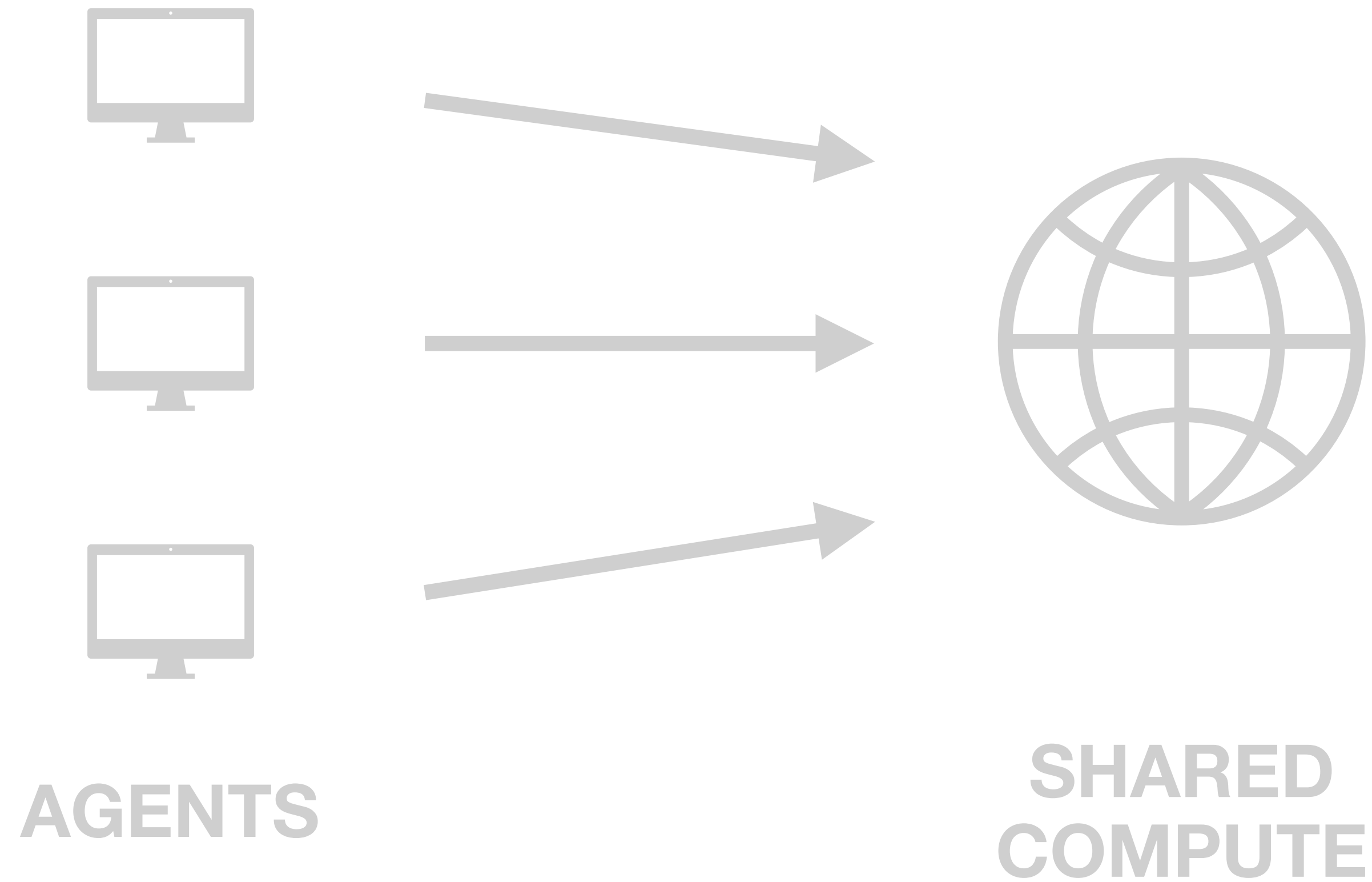


COMPUTE



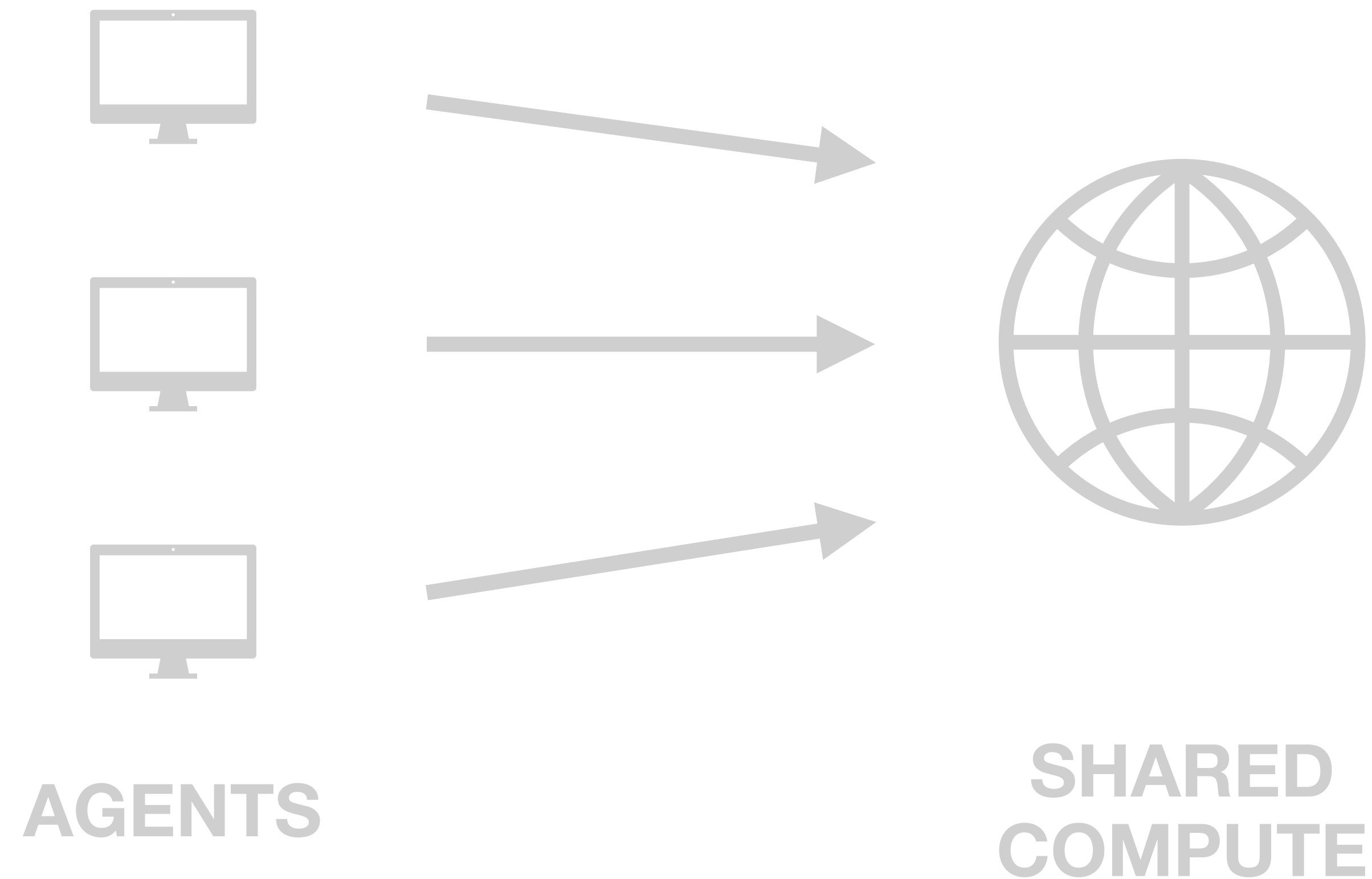
IDLE

Tragedy of the Commons



Guardrails:

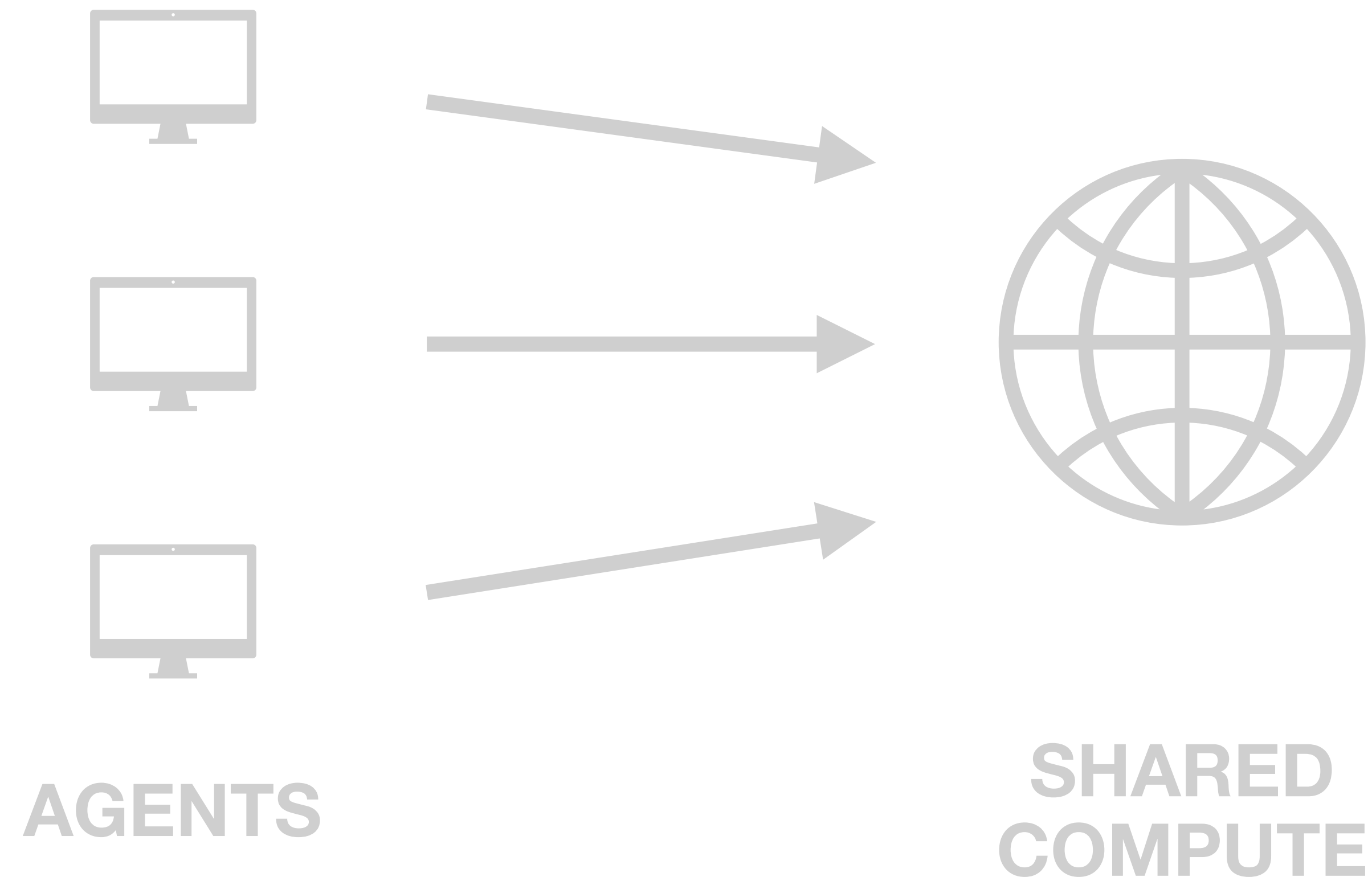
Tragedy of the Commons



Guardrails:

- We can change the **agent behavior** to limit utilization.

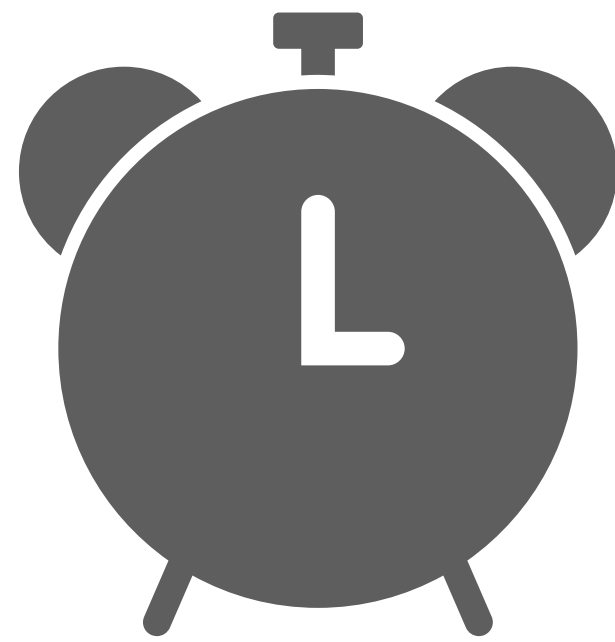
Tragedy of the Commons



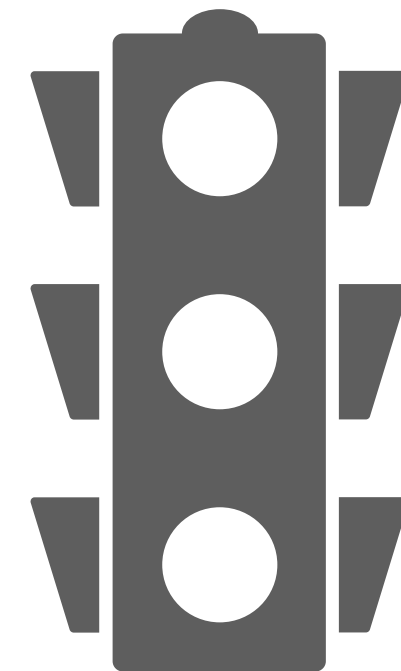
Guardrails:

- We can change the **agent behavior** to limit utilization.
- We can introduce a **central coordinator** to handle usage.

Real-World Solutions



Exponential Backoff
(Client behavior in TCP)



Traffic Light
(Central coordinator)

- **Learnable mixed Nash equilibria are collectively rational**
with Yian Ma (preprint)

This work: the relationship between

- **[learnability]** the collective behavior that emerges out of adaptivity, and
- **[rationality]** the social efficiency of the solutions found.

Game theory background

Hallmark of a Game

Decisions are made **individually**, but the **outcome** depends on **collective** choices.

Mathematical Formalism

An N -player game:

- Players indexed by $n \in \{1, \dots, N\}$

Mathematical Formalism

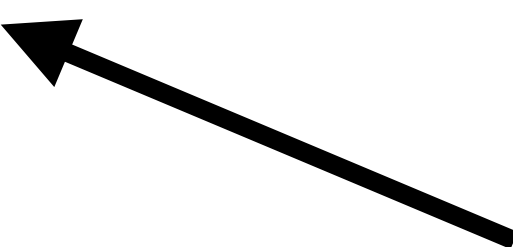
An N -player game:

- Players indexed by $n \in \{1, \dots, N\}$
- Each player has a decision variable $x_n \in \Omega_n$

Mathematical Formalism

An N -player game:

- Players indexed by $n \in \{1, \dots, N\}$
- Each player has a decision variable $x_n \in \Omega_n$
- Joint decision: $\mathbf{x} = (x_1, \dots, x_N) \in \Omega$

$$\Omega := \Omega_1 \times \dots \times \Omega_N$$


Mathematical Formalism

An N -player game:

- Players indexed by $n \in \{1, \dots, N\}$
- Each player has a decision variable $x_n \in \Omega_n$
 - Joint decision: $\mathbf{x} = (x_1, \dots, x_N) \in \Omega$
 - Notation: $\mathbf{x}_{-n} := (x_1, \dots, x_{n-1}, x_{n+1}, \dots, x_N)$

Mathematical Formalism

An N -player game:

- Players indexed by $n \in \{1, \dots, N\}$
- Each player has a decision variable $x_n \in \Omega_n$
 - Joint decision: $\mathbf{x} = (x_1, \dots, x_N) \in \Omega$
 - Notation: $\mathbf{x}_{-n} := (x_1, \dots, x_{n-1}, x_{n+1}, \dots, x_N)$
- Each player wants to maximize their utility $u_n : \Omega \rightarrow \mathbb{R}$

$$u_n(\mathbf{x}) \equiv u_n(x_n; \mathbf{x}_{-n}).$$

Mathematical Formalism

An N -player game:

- Players indexed by $n \in \{1, \dots, N\}$
- Each player has a decision variable $x_n \in \Omega_n$
 - Joint decision: $\mathbf{x} = (x_1, \dots, x_N) \in \Omega$
 - Notation: $\mathbf{x}_{-n} := (x_1, \dots, x_{n-1}, x_{n+1}, \dots, x_N)$
- Each player wants to maximize their utility $u_n : \Omega \rightarrow \mathbb{R}$

$$u_n(\mathbf{x}) \equiv u_n(x_n; \mathbf{x}_{-n}).$$

Unlike economics, in computer science, we often choose the utility for the agents we build.

Repeated Games

Players enter a game with limited knowledge:

- They know what set of actions they/others can take.
- They know their own utilities.

They do not know:

- What other players want.
- How other players learn.



The players are **uncoupled**.

Repeated Games

Players enter a game with limited knowledge:

- They know what set of actions they/others can take.
- They know their own utilities.

They do not know:

- What other players want.
- How other players learn.

Players will repeatedly play the game, observe what other do, and **adjust their own behaviors for future rounds.**

Learning Setting

For $t = 1, 2, \dots$:

- Each agent plays an action $x_n(t)$.
- Each agent observes joint strategy played $\mathbf{x}(t)$.
- Each agent receives utility $u_n(\mathbf{x}(t))$

Learning Rule

A learning rule describes how an agent uses past observations to make future decisions:

$$x_n(t + 1) \leftarrow \text{Learning Rule} \left(\mathbf{x}(1), \dots, \mathbf{x}(t); u_n \right)$$

Learning Rule

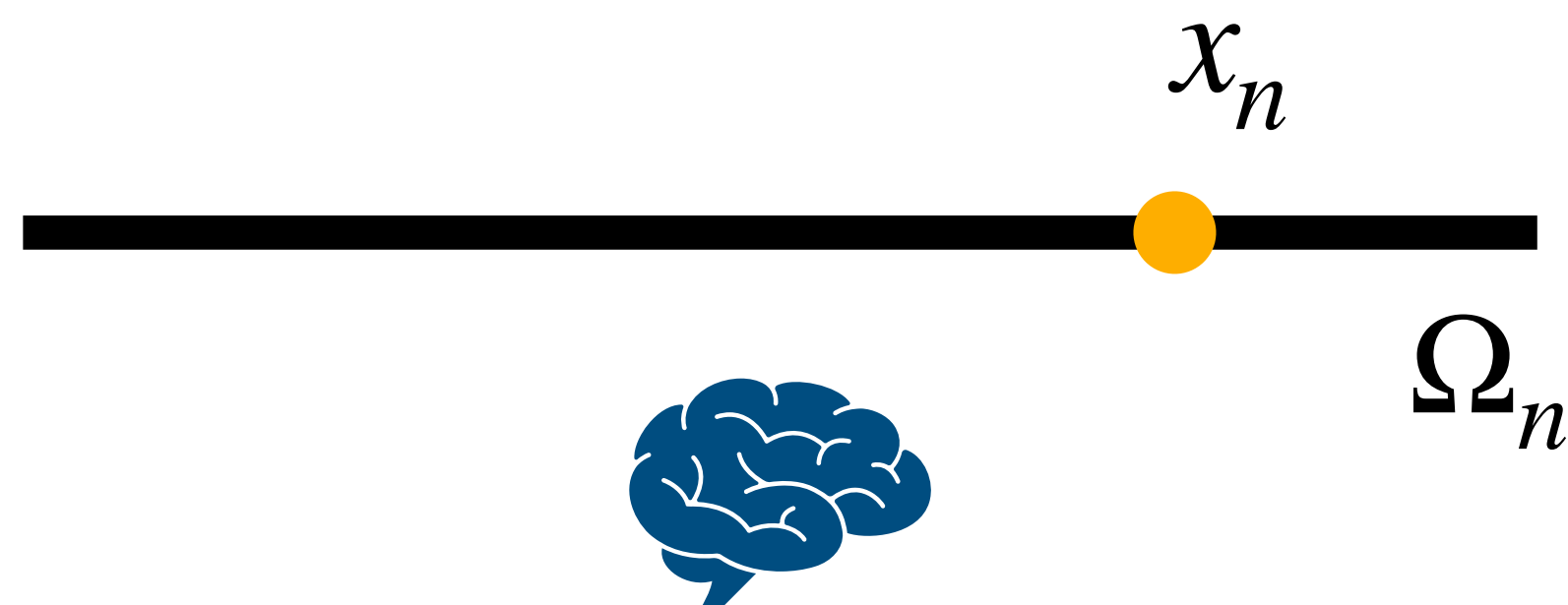
A learning rule describes how an agent uses past observations to make future decisions:

$$x_n(t + 1) \leftarrow \text{Learning Rule} \left(\mathbf{x}(t); u_n \right)$$

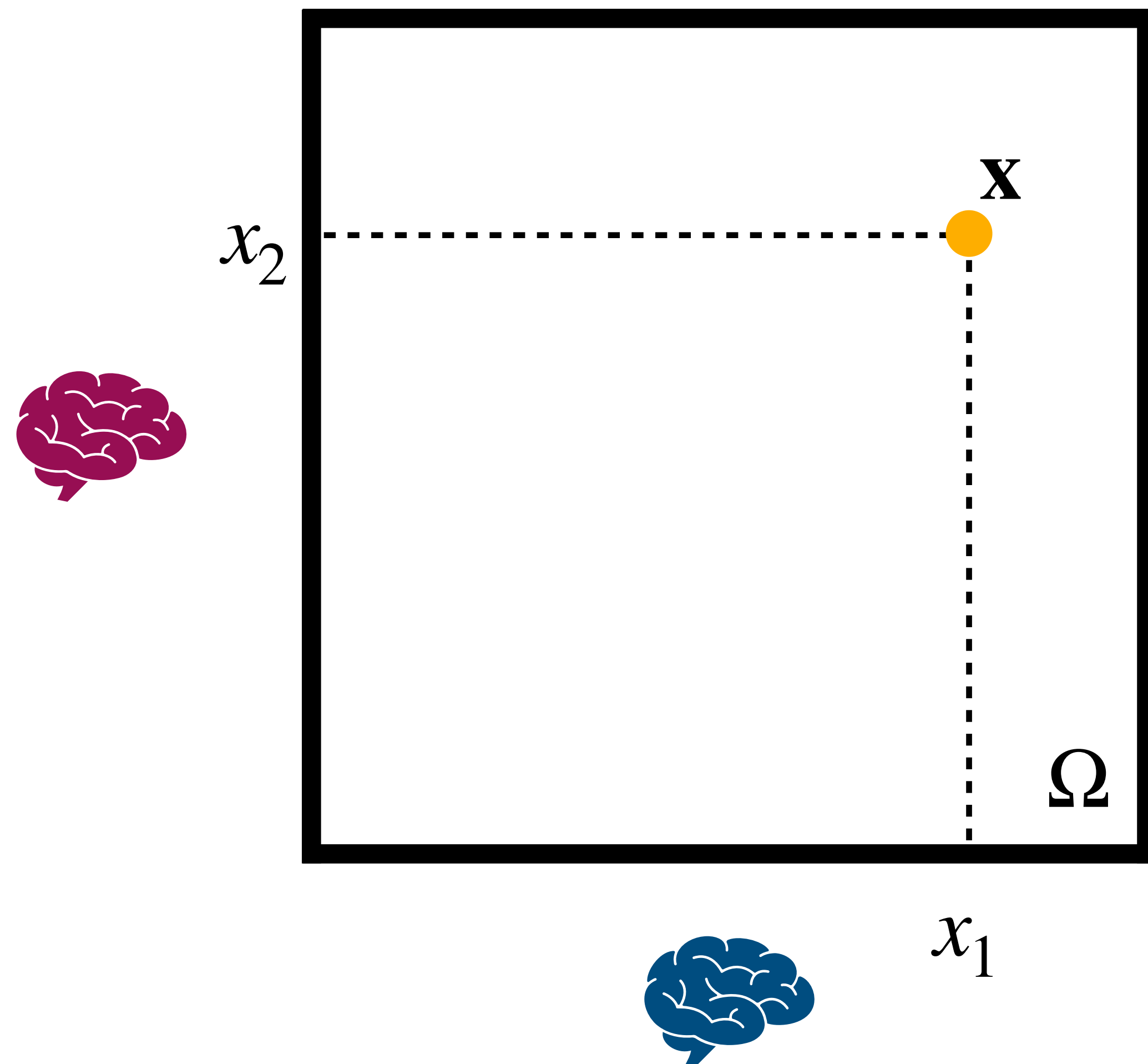
For simplicity, we assume agents are myopic/Markov.

A dynamical-systems perspective

Let Ω_n be player n 's **strategy space**.



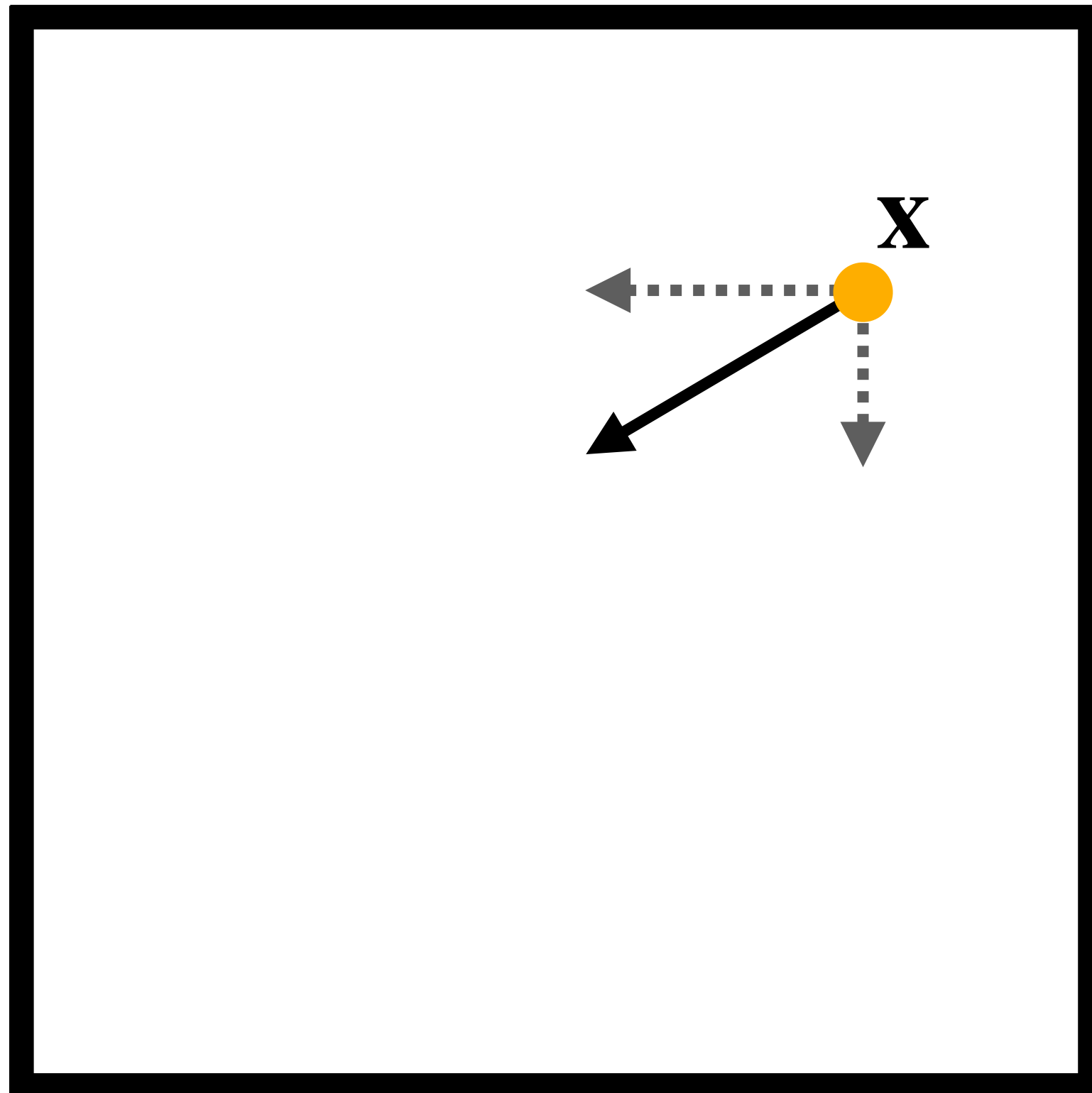
A dynamical-systems perspective



Let Ω be the **joint strategy space**.

- Each point $\mathbf{x} \in \Omega$ corresponds to a configuration of the agents.

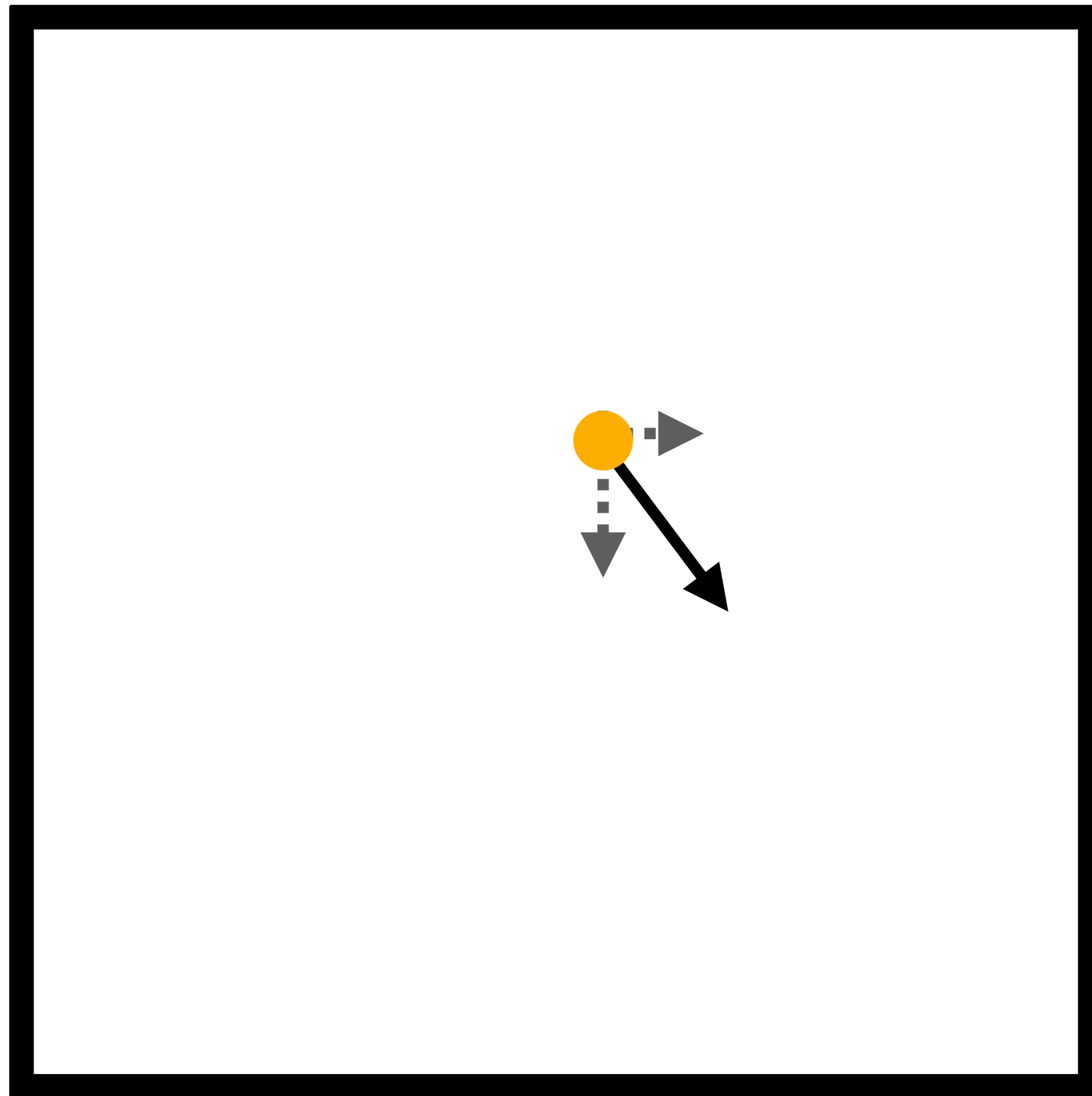
A dynamical-systems perspective



Each agent has a learning rule:
How it adapts under the situation $x \in \Omega$.

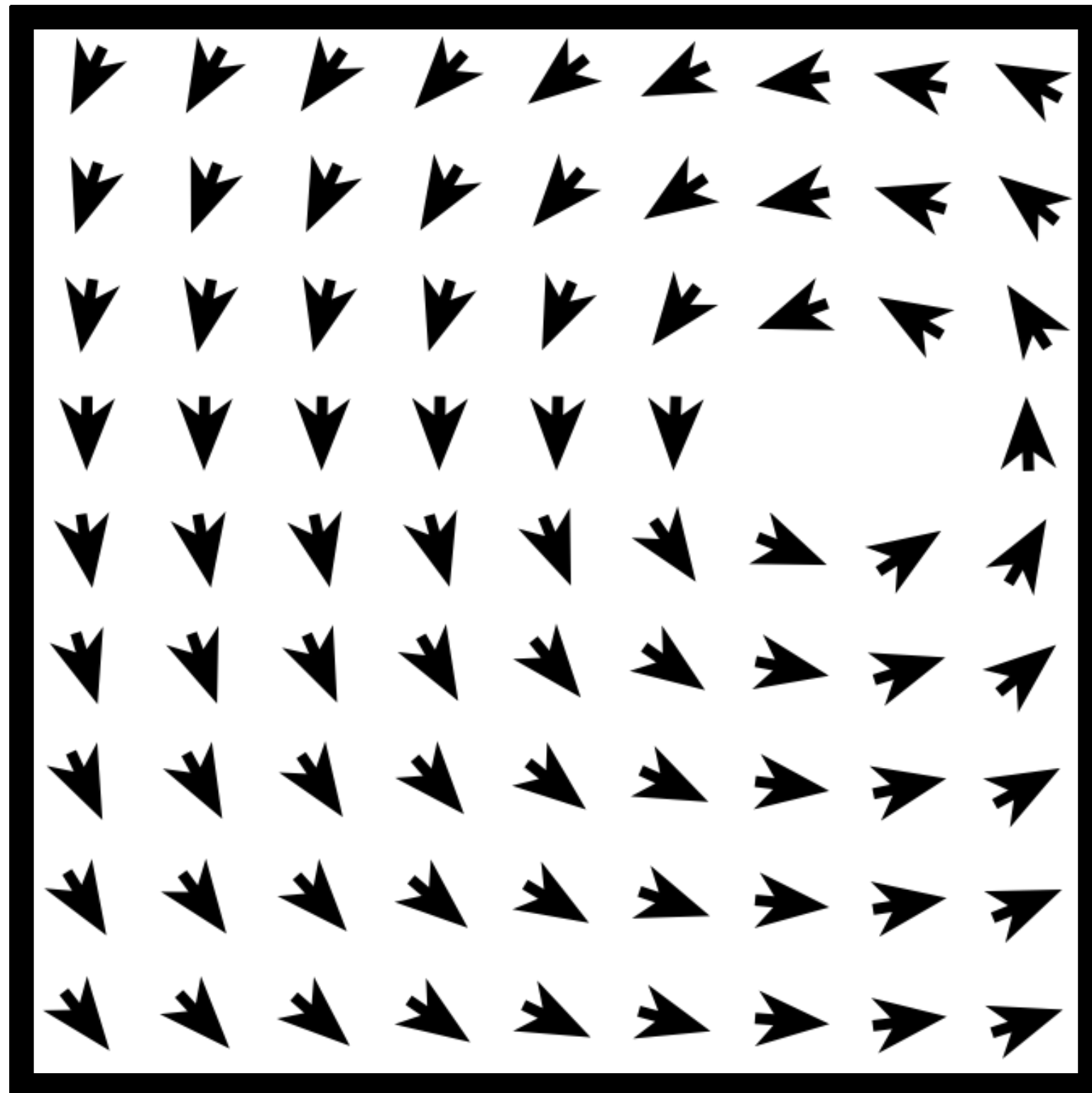
- Each aims to maximize its own utility.

A dynamical-systems perspective



- Each agent has a learning rule:**
How it adapts under the situation $\mathbf{x} \in \Omega$.
- Each aims to maximize its own utility.

A dynamical-systems perspective



- Each agent has a learning rule:**
How it adapts under the situation $\mathbf{x} \in \Omega$.
- Each aims to maximize its own utility.

Simple Learning Rules

- **Best Response:** optimize, assuming that players in the next round will continue to play the current strategy.
- **Follow the Leader:** optimize, assuming that past experience forms a representative sample of strategies that will be played.
- **No-Regret Learning:** apply standard online learning algorithms.

Simple Learning Rules

- **Best-response dynamics**

$$x_n(t + 1) = \arg \max_{z_n \in \Omega_n} u_n(z_n; \mathbf{x}_{-n}(t))$$

- **Smoothed best-response dynamics**
- **Gradient ascent**
- **Positive-definite dynamics**

Simple Learning Rules

- **Best-response dynamics**
- **Smoothed best-response dynamics**

$$x_n(t + 1) = \arg \max_{z_n \in \Omega_n} u_n(z_n; \mathbf{x}_{-n}(t)) + \text{regularizer}(z_n)$$

- **Gradient ascent**
- **Positive-definite dynamics**

Simple Learning Rules

- **Best-response dynamics**
- **Smoothed best-response dynamics**
- **Gradient ascent-ascent**

$$\dot{x}_n(t) = \eta \nabla_n u_n(x_n(t); \mathbf{x}_{-n}(t))$$

- **Positive-definite dynamics**

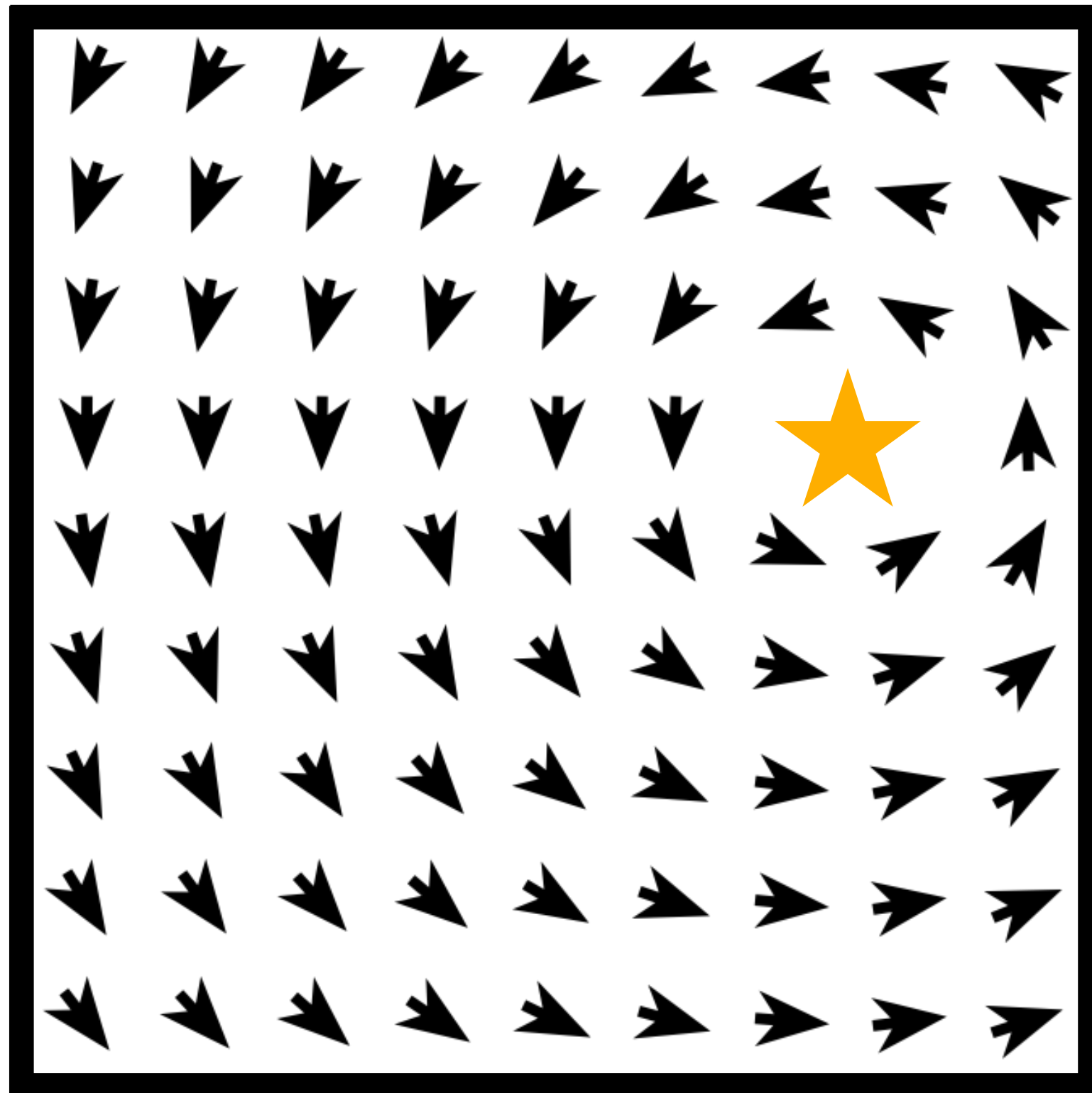
Simple Learning Rules

- **Best-response dynamics**
- **Smoothed best-response dynamics**
- **Gradient ascent**
- **Positive-definite dynamics**

$$\dot{\mathbf{x}}(t) = H_n^{-1} \nabla_n u_n(\mathbf{x}(t))$$

where H_n is positive-definite (also, preconditioned gradient ascent, mirror ascent,...).

Learning in games



Goal

Understand how the joint system evolves.

Fixed points / equilibria

What are potential behaviors that can persist and have long-term relevance?

Nash equilibria

Definition. A joint strategy $\mathbf{x}^\star$ is a Nash equilibrium if no player can unilaterally improve their own utility

Nash equilibria

Definition. A joint strategy $\mathbf{x}^\star$ is a Nash equilibrium if no player can unilaterally improve their own utility

$$u_n(\mathbf{x}^\star) = \max_{x_n \in \Omega_n} u_n(x_n; \mathbf{x}_{-n}^\star).$$

Nash equilibria

Definition. A joint strategy $\mathbf{x}^\star$ is a Nash equilibrium if no player can unilaterally improve their own utility

$$u_n(\mathbf{x}^\star) = \max_{x_n \in \Omega_n} u_n(x_n; \mathbf{x}_{-n}^\star).$$

This is a notion of **individual rationality**: players would need to *collaborate* to further improve their utilities.

Nash equilibria

Nash equilibria (NE) describe behaviors of multi-agent systems that may have **long-term relevance**.

- ▶ No agent has any incentive to deviate from the status quo.

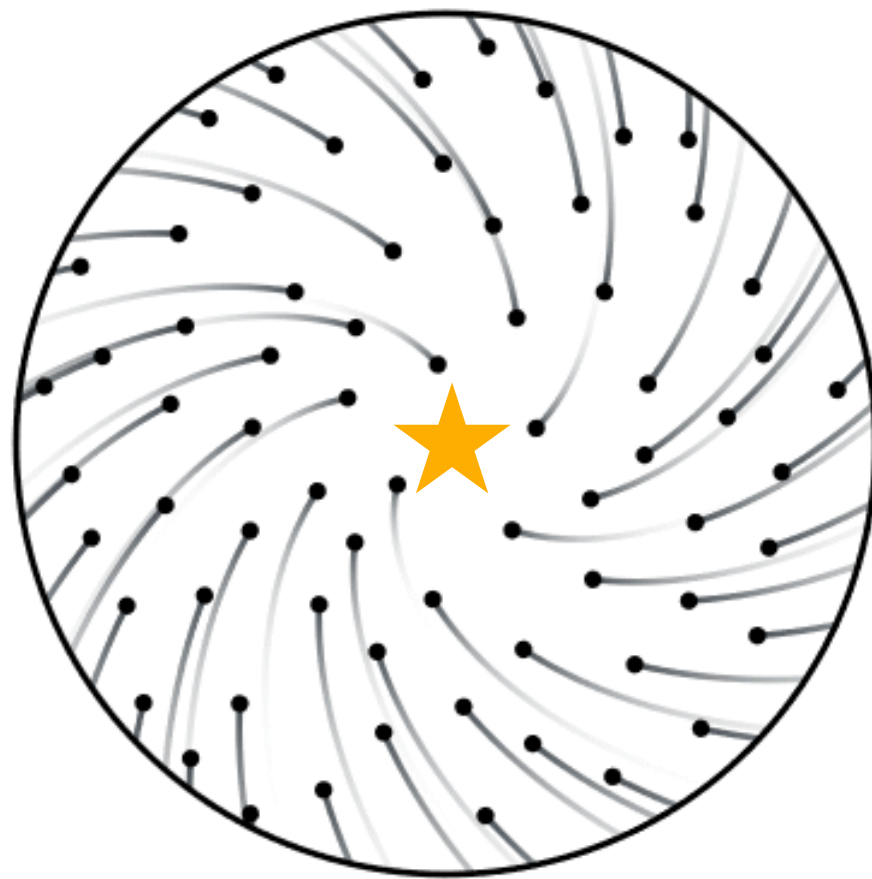
Nash equilibria

Nash equilibria (NE) describe behaviors of multi-agent systems that may have **long-term relevance**.

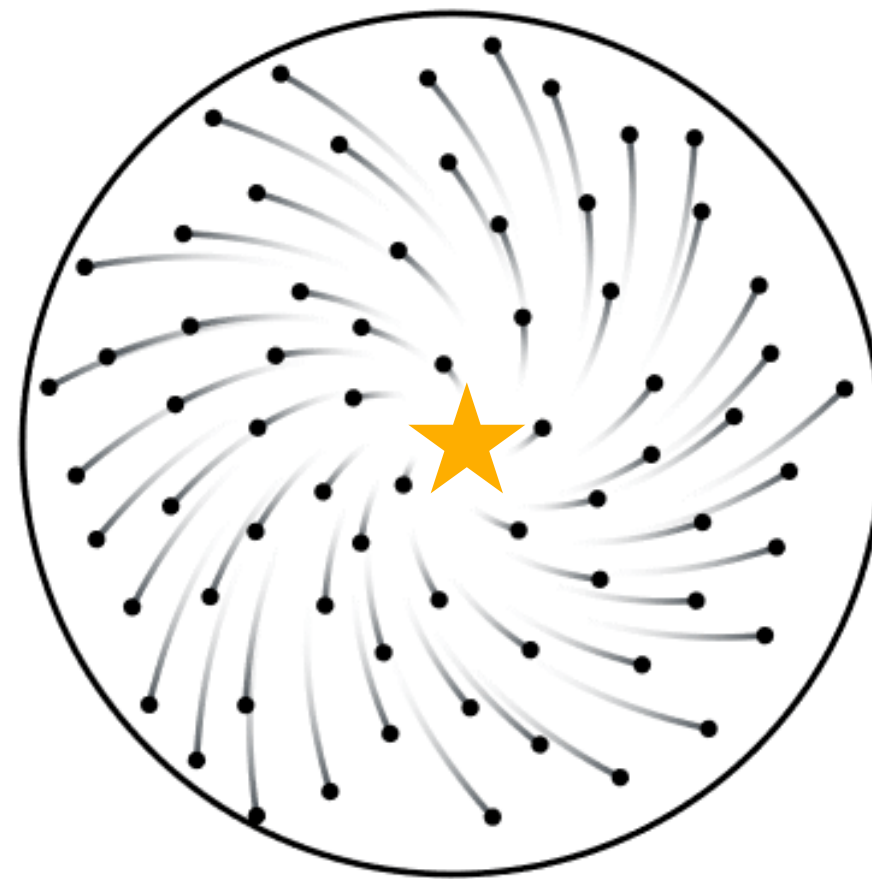
- ▶ No agent has any incentive to deviate from the status quo.
- ▶ NE are *stationary* points, but not necessarily *stable* (attractors).

Classic notions of stability (learnability)

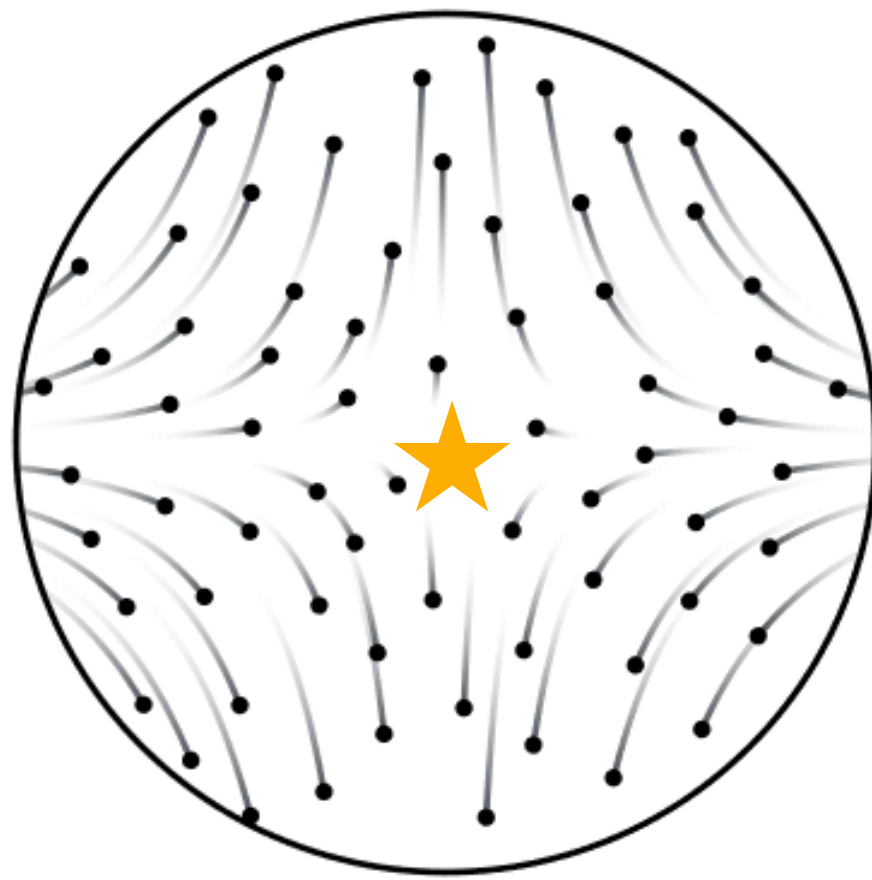
Stable fixed point



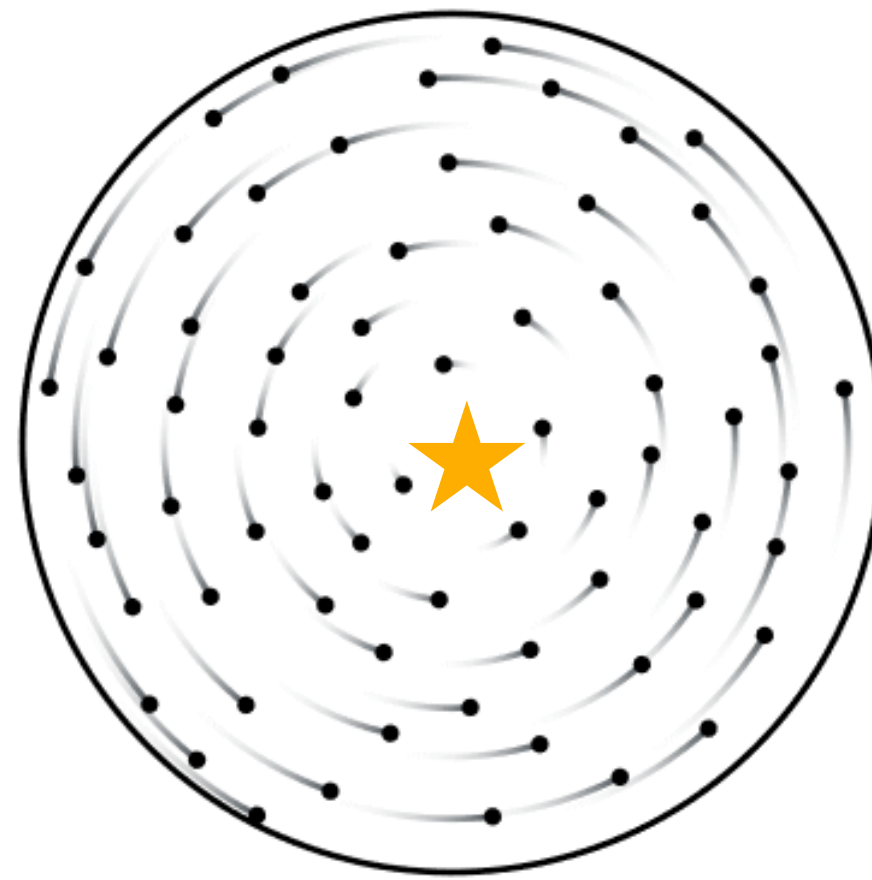
Unstable fixed point



Hyperbolic fixed point



Elliptic fixed point



Asymptotic stability

The dynamics will eventually converge to ★.

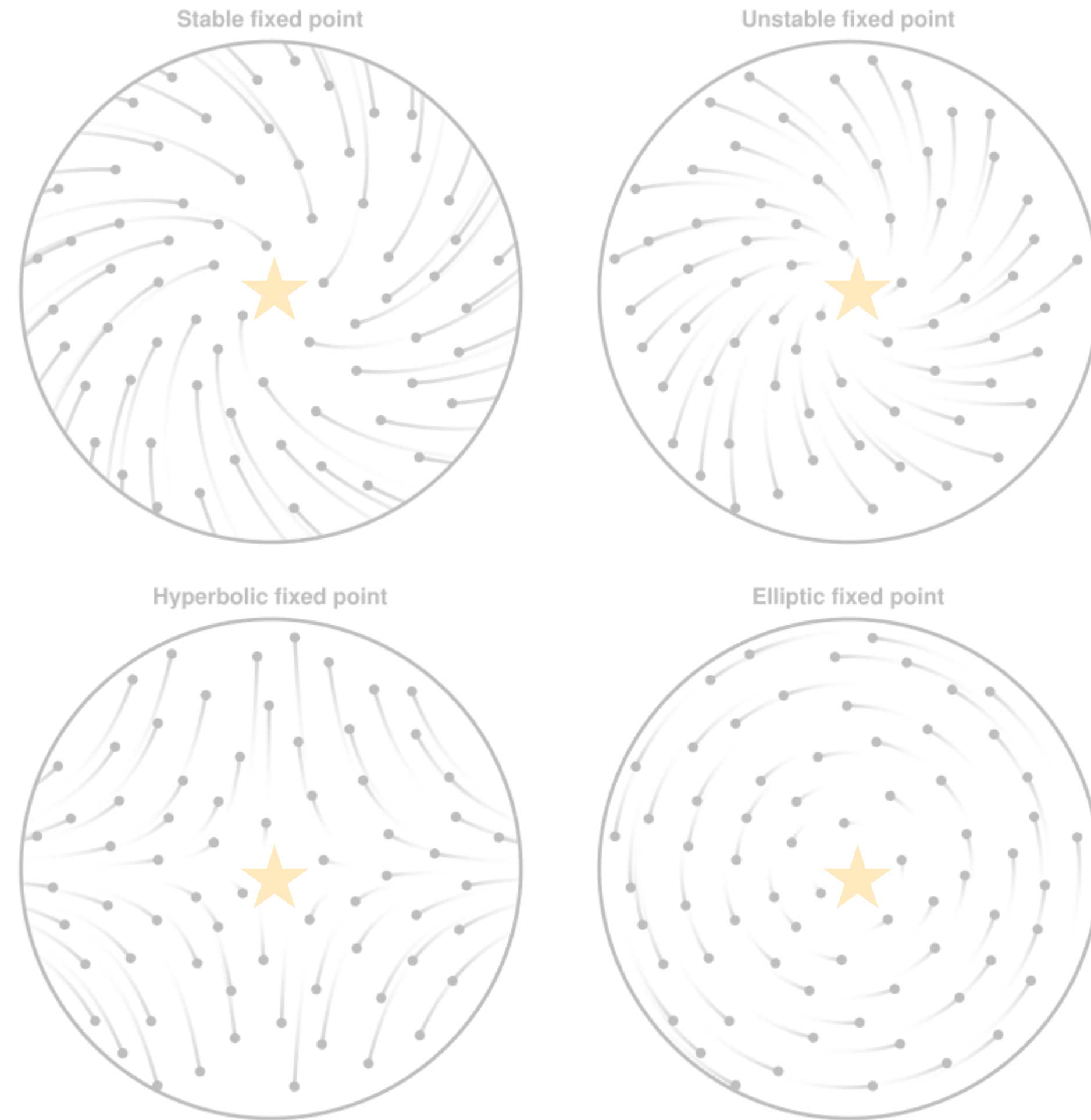
Non-asymptotic stability

The dynamics starting nearby will stay close.

Instability

The dynamics will eventually escape ★.

Classic notions of stability



Asymptotic stability

The dynamics will eventually converge to ★.

Existing work only focus on this “strong” notion of learnability.

Asymptotic stability

A notion of “strong learnability”

- Dynamics are robust to perturbations
- It will eventually converge to the Nash equilibrium
- Guarantees last-iterate convergence

An impossibility result on finding Nash equilibria

Uncoupled dynamics do not lead to Nash equilibrium, Hart and Mas-Colell 2003

Main result: there is a family of games $\mathcal{U}$ such that:

An impossibility result on finding Nash equilibria

Uncoupled dynamics do not lead to Nash equilibrium, Hart and Mas-Colell 2003

Main result: there is a family of games $\mathcal{U}$ such that:

1. Each game has a unique Nash equilibrium.

An impossibility result on finding Nash equilibria

Uncoupled dynamics do not lead to Nash equilibrium, Hart and Mas-Colell 2003

Main result: there is a family of games $\mathcal{U}$ such that:

1. Each game has a unique Nash equilibrium.
2. For every uncoupled dynamics, the dynamics will fail to converge to the Nash equilibrium with asymptotic stability for at least one game in this family.

An impossibility result on finding Nash equilibria

Uncoupled dynamics do not lead to Nash equilibrium, Hart and Mas-Colell 2003

Main result: there is a family of games $\mathcal{U}$ such that:

1. Each game has a unique Nash equilibrium.
2. For every uncoupled dynamics, the dynamics will fail to converge to the Nash equilibrium with asymptotic stability for at least one game in this family.

Read: due to *informational constraints*, not all Nash equilibria are “strongly learnable”.

Which Nash equilibria can we find?

Convergent dynamics exist for certain important subclasses of equilibria:

Which Nash equilibria can we find?

Convergent dynamics exist for certain important subclasses of equilibria:

- For many standard learning dynamics, there are result of the form:

*An equilibrium admits asymptotic convergence
if and only if it is a **strict** Nash equilibrium.*

Which Nash equilibria can we find?

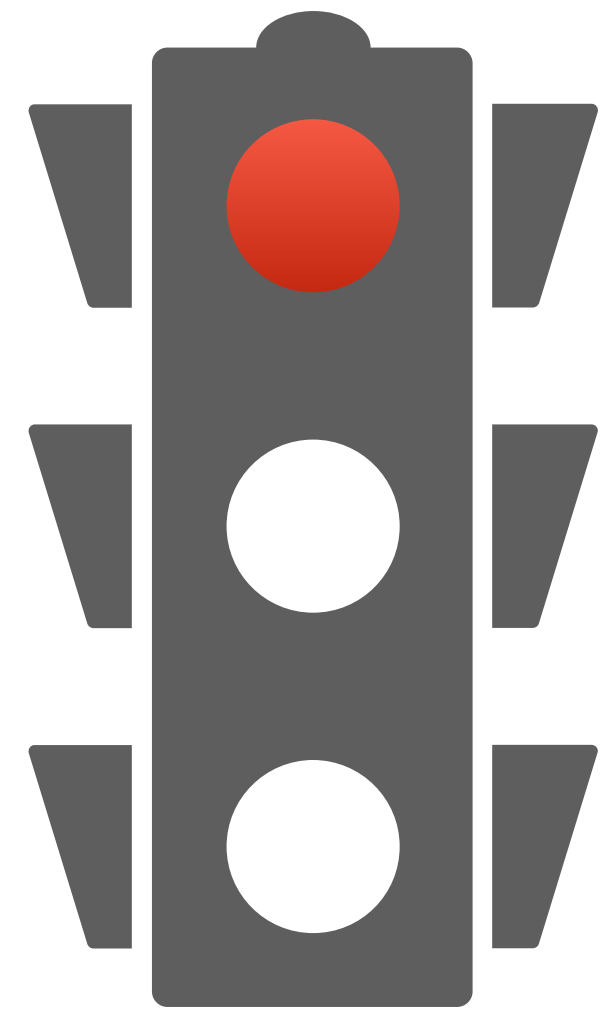
Convergent dynamics exist for certain important subclasses of equilibria:

- For many standard learning dynamics, there are result of the form:

*An equilibrium admits asymptotic convergence
if and only if it is a **strict** Nash equilibrium.*

- Fictitious play, replicator dynamics, regularized learning, no-regret learning...
Samuelson and Zhang (1992); Vlastakis-Gkaragkounis et. al. 2020; Giannou et. al. (2021) ...

Stability of Social Conventions



Strict Nash equilibria $\approx$ social conventions

Each player has a deterministic strategy that is locally optimal.

Example. Once enough people agree that **RED** means **STOP**, everyone else is forced to adopt this convention.



Asymptotic stability does not guarantee **collective rationality**.

Stability in Non-Strict Equilibria

What about dynamics around
non-strict Nash equilibria?



Stability in Non-Strict Equilibria

What about dynamics around non-strict Nash equilibria?

Example. When the job market is at equilibrium, many career options are viable, leading to a mixed population.



Instability for non-strict/mixed NE

Many learning dynamics exhibit **instability** around **mixed** NE.

Convergence in Zero-Sum Games

However, there are uncoupled dynamics that can be stabilized to certain mixed equilibria (e.g. those in zero-sum games):

- Optimistic gradient descent-ascent
- Extragradient method
- Smoothed best-response

Behavior around mixed equilibria unclear

- Asymptotic stability ruled out
- Instability can be observed
- But, some mixed equilibria can be stabilized

Some foreshadowing

Uncoupled dynamics do not lead to Nash equilibrium, Hart and Mas-Colell 2003

Asymptotic stability is too strong of a criterion of learnability

Some foreshadowing

Uncoupled dynamics do not lead to Nash equilibrium, Hart and Mas-Colell 2003

Asymptotic stability is too strong of a criterion of learnability

- We will introduce a non-asymptotic notion called uniform stability.

Some foreshadowing

Uncoupled dynamics do not lead to Nash equilibrium, Hart and Mas-Colell 2003

Asymptotic stability is too strong of a criterion of learnability

- We will introduce a non-asymptotic notion called uniform stability.
- Perhaps it is also not so bad we cannot find certain Nash equilibria.

Stability in Non-Strict Equilibria



What about dynamics around
non-strict Nash equilibria?

Example. When the job market is at equilibrium, many career options are viable, leading to a mixed population.



Uniform stability guarantees **collective rationality**.

This work

Pareto optimality

A solution is (weakly) **Pareto optimal** if, even with cooperation, there is **no way to strictly improve the outcomes for all agents** at once.

Pareto optimality

Definition. A joint strategy $\mathbf{x}^\star$ is (weakly) Pareto optimal if there is no way to improve all players at the same time

Pareto optimality

Definition. A joint strategy $\mathbf{x}^\star$ is (weakly) Pareto optimal if there is no way to improve all players at the same time

$$\nexists \mathbf{x} \in \Omega \quad \text{s.t.} \quad u_n(\mathbf{x}) > u_n(\mathbf{x}^\star) \text{ for all } n \in [N].$$

Pareto optimality

Definition. A joint strategy $\mathbf{x}^\star$ is (weakly) Pareto optimal if there is no way to improve all players at the same time

$$\nexists \mathbf{x} \in \Omega \quad \text{s.t.} \quad u_n(\mathbf{x}) > u_n(\mathbf{x}^\star) \text{ for all } n \in [N].$$

It would be **collectively irrational** to jointly play $\mathbf{x}^\star$ if everyone is happier with $\mathbf{x}$.

Invisible hand of stability

Theorem, informal (So and Ma 2025)

The mixed Nash equilibria that are **uniformly learnable** are the ones that are **strategically Pareto optimal**.

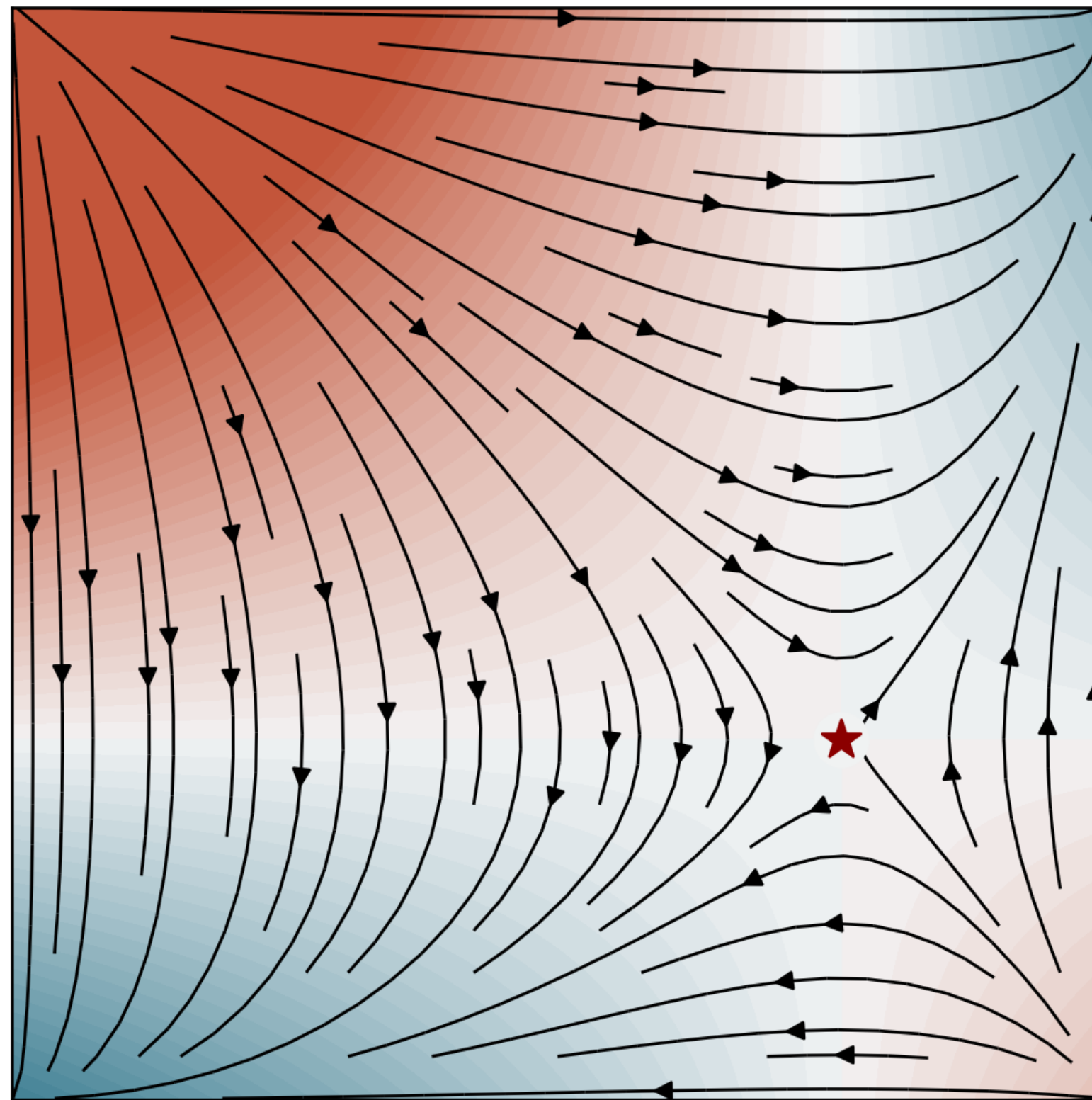
Invisible hand of stability

Theorem, informal (So and Ma 2025)

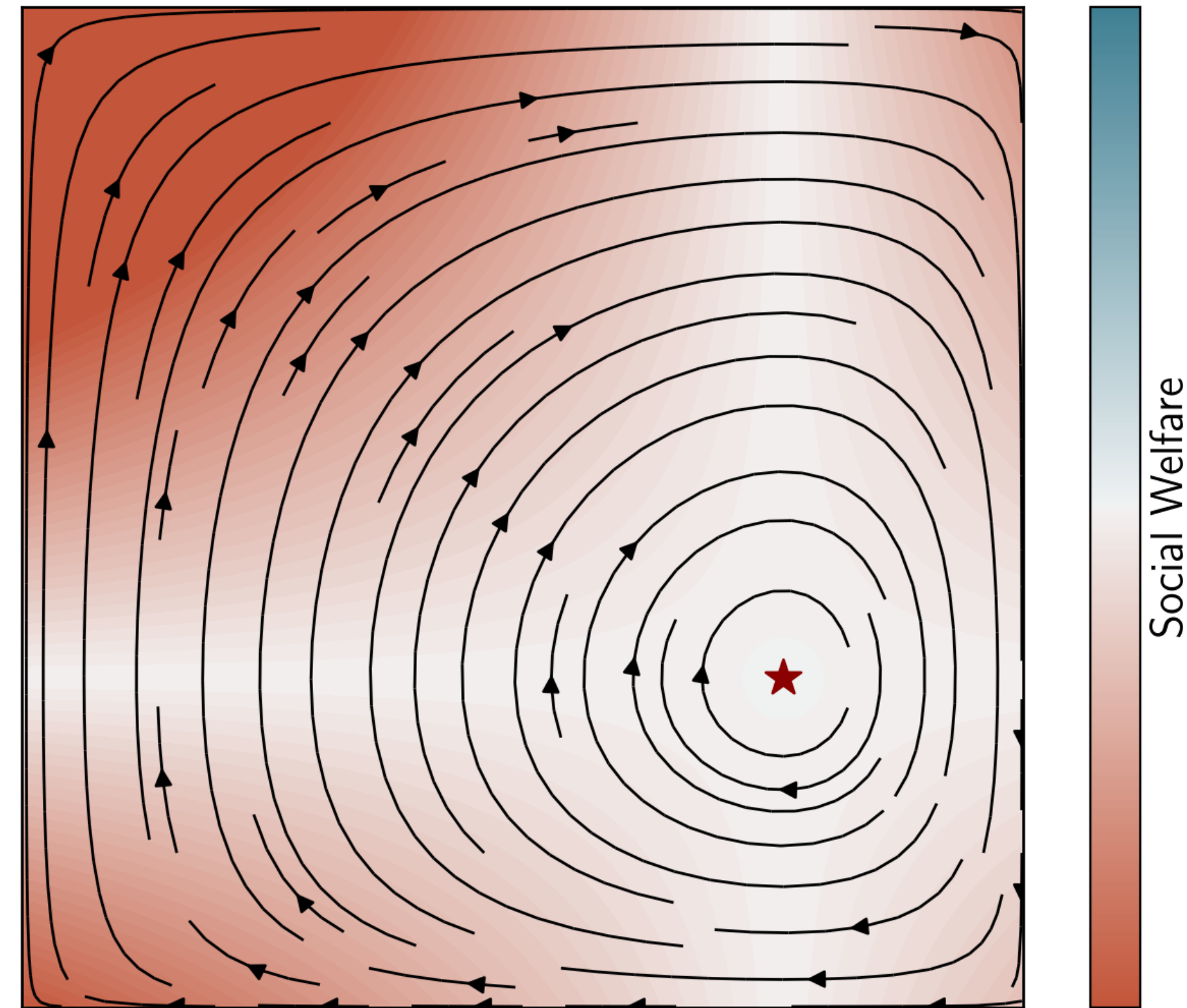
The mixed Nash equilibria that are **uniformly learnable** are the ones that are **strategically Pareto optimal**.

- The notion of learnability is a weaker, non-asymptotic criterion.
- Non-convergence is not necessarily a bad thing if the equilibrium is Pareto dominated.
- When can learning agents act independently, and when is communication or coordination necessary to achieve better collective outcomes?

Comparison of Dynamics Around Non-Strict NE



Unstable Nash equilibria are not strategically Pareto optimal.



Uniformly stable Nash equilibria are strategically Pareto optimal.

Our results

1. Connection between stability and collective rationality.

local uniform stability $\implies$ Pareto optimality* $\implies$ pointwise uniform stability

* Technically, strategic Pareto optimality, which holds up to strategic equivalence.

Our results

1. Connection between stability and collective rationality.

local uniform stability $\implies$ Pareto optimality* $\implies$ pointwise uniform stability

2. Dynamics of a specific class of learning rules.

- **Incremental smooth best-response dynamics**
- **Local uniform stability $\implies$ all dynamics can be stabilized**
- **Not uniformly stable $\implies$ certain dynamics can never be stabilized**

* Technically, strategic Pareto optimality, which holds up to strategic equivalence.

Takeaway

- Modern ML often implements multi-agent solutions.
- Decentralization introduces structural constraints.
- What are the ramifications and when are guardrails needed?

Linear Stability Analysis

Consider the following continuous-time, smooth dynamical system:

$$\dot{\mathbf{x}}(t) = T(\mathbf{x}(t)).$$

- We can study the **stability** of the dynamics around a fixed point $\mathbf{x}^\dagger$ by looking at the **linearized dynamics** around $\mathbf{x}^\dagger$.

Linearized Dynamics

For simplicity, let $\mathbf{x} \in \mathbb{R}^N$. The **Jacobian** $\nabla T(\mathbf{x}) \in \mathbb{R}^{N \times N}$ is given by:

$$\nabla T = \begin{bmatrix} \nabla_1 T_1 & \nabla_2 T_1 & \cdots & \nabla_N T_1 \\ \nabla_1 T_2 & \nabla_2 T_2 & & \\ \vdots & & \ddots & \\ \nabla_1 T_N & & & \nabla_N T_N \end{bmatrix}$$

Linearized Dynamics

For simplicity, let $\mathbf{x} \in \mathbb{R}^N$. The **Jacobian** $\nabla T(\mathbf{x}) \in \mathbb{R}^{N \times N}$ is given by:

$$\nabla T = \begin{bmatrix} \nabla_1 T_1 & \nabla_2 T_1 & \cdots & \nabla_N T_1 \\ \nabla_1 T_2 & \nabla_2 T_2 & & \\ \vdots & & \ddots & \\ \nabla_1 T_N & & & \nabla_N T_N \end{bmatrix}$$

Around the fixed point, the linearized dynamics are defined by:

$$\dot{\mathbf{z}}(t) = \nabla T(\mathbf{x}^\dagger) \mathbf{z}(t)$$

where $\mathbf{z} \equiv \mathbf{x} - \mathbf{x}^\dagger$.

Stability from Spectral Analysis

Stability of **linear** dynamical systems well-understood:

Stability from Spectral Analysis

Stability of **linear** dynamical systems well-understood:

- If all eigenvalues of ∇T have **negative real parts** $\implies$ asymptotic stability

Stability from Spectral Analysis

Stability of **linear** dynamical systems well-understood:

- If all eigenvalues of ∇T have **negative real parts** $\implies$ asymptotic stability
- If there is an eigenvalue with **positive real part** $\implies$ instability

Stability from Spectral Analysis

Stability of **linear** dynamical systems well-understood:

- If all eigenvalues of ∇T have **negative real parts** $\implies$ asymptotic stability
- If there is an eigenvalue with **positive real part** $\implies$ instability
- If all eigenvalues are **purely imaginary** $\implies$ neutral stability

Stability from Spectral Analysis

Stability of **linear** dynamical systems well-understood:

- If all eigenvalues of ∇T have **negative real parts** $\implies$ asymptotic stability
- If there is an eigenvalue with **positive real part** $\implies$ instability
- If all eigenvalues are **purely imaginary** $\implies$ neutral stability
 - This case is much harder to characterize for **non-linear** systems

Linearizing Gradient Ascent Dynamics

The gradient ascent dynamics:

$$\dot{\mathbf{x}}_n(t) = \nabla_n u_n(\mathbf{x}(t))$$

Linearizing Gradient Ascent Dynamics

The gradient ascent dynamics:

$$\dot{x}_n(t) = \nabla_n u_n(\mathbf{x}(t))$$

- The Jacobian of this dynamics is given by:

$$\mathbf{J} = \begin{bmatrix} \nabla_{11}^2 u_1 & \nabla_{12}^2 u_1 & \cdots & \nabla_{1N}^2 u_1 \\ \nabla_{21}^2 u_2 & \nabla_{22}^2 u_2 & & \\ \vdots & & \ddots & \\ \nabla_{N1}^2 u_N & & & \nabla_{NN}^2 u_N \end{bmatrix}$$

Linearizing Gradient Ascent Dynamics

The gradient ascent dynamics:

$$\dot{x}_n(t) = \nabla_n u_n(\mathbf{x}(t))$$

- The Jacobian of this dynamics is given by:

$$\mathbf{J} = \begin{bmatrix} 0 & \nabla_{12}^2 u_1 & \cdots & \nabla_{1N}^2 u_1 \\ \nabla_{21}^2 u_2 & 0 & & \\ \vdots & & \ddots & \\ \nabla_{N1}^2 u_N & & & 0 \end{bmatrix}$$

- The diagonal is zero by multilinearity in normal-form games.

Linearizing Gradient Ascent Dynamics

The gradient ascent dynamics:

$$\dot{x}_n(t) = \nabla_n u_n(\mathbf{x}(t))$$

- The Jacobian of this dynamics is given by:

$$\mathbf{J} = \begin{bmatrix} 0 & \nabla_{12}^2 u_1 & \cdots & \nabla_{1N}^2 u_1 \\ \nabla_{21}^2 u_2 & 0 & & \\ \vdots & & \ddots & \\ \nabla_{N1}^2 u_N & & & 0 \end{bmatrix}$$

- This object is called the **game Jacobian**.

Jacobian of learning dynamics

The Jacobian of common classes of learning dynamics have the form:

$$\nabla T = \mathbf{H}^{-1} \mathbf{J}$$

where $\mathbf{J}$ is the game Jacobian and $\mathbf{H}$ is block-diagonal and positive definite:

$$\mathbf{H} = \begin{bmatrix} H_1 & & & \\ & H_2 & & \\ & & \ddots & \\ & & & H_N \end{bmatrix}$$

Jacobian of learning dynamics

- **Preconditioned gradient ascent**

$$\dot{\mathbf{x}}_n(t) = H_n^{-1} \nabla_n u_n(\mathbf{x}(t))$$

- The preconditioned H_n also captures some local geometry.
 - Player-dependent notion of “steepest ascent”.

The Jacobian of the dynamics are traceless

For learning dynamics $T : \Omega \rightarrow \Omega$ whose Jacobians are of the form:

$$\nabla T = \mathbf{H}^{-1} \mathbf{J} = \begin{bmatrix} 0 & H_1^{-1} J_{12} & \cdots & H_1^{-1} J_{1N} \\ H_2^{-1} J_{21} & 0 & & \\ \vdots & & \ddots & \\ H_N^{-1} J_{N1} & & & 0 \end{bmatrix}$$

the trace is zero, $\text{tr}(\mathbf{H}^{-1} \mathbf{J}) = 0$.

Barrier to Asymptotic Stability

- Suppose the dynamics around a fixed point is **traceless**: $\text{tr}(\nabla T(\mathbf{x}^\dagger)) = 0$.

Barrier to Asymptotic Stability

- Suppose the dynamics around a fixed point is **traceless**: $\text{tr}(\nabla T(\mathbf{x}^\dagger)) = 0$.
 - Recall that the trace is equal to the sum of the eigenvalues:

$$\text{tr}(\nabla T(\mathbf{x}^\dagger)) = \lambda_1 + \cdots + \lambda_N .$$

Barrier to Asymptotic Stability

- Suppose the dynamics around a fixed point is **traceless**: $\text{tr}(\nabla T(\mathbf{x}^\dagger)) = 0$.
 - Recall that the trace is equal to the sum of the eigenvalues:

$$\text{tr}(\nabla T(\mathbf{x}^\dagger)) = \lambda_1 + \cdots + \lambda_N.$$

- Asymptotic stability requires that the real part $\Re(\lambda_n) < 0$ for all $n \in [N]$.

Barrier to Asymptotic Stability

- Suppose the dynamics around a fixed point is **traceless**: $\text{tr}(\nabla T(\mathbf{x}^\dagger)) = 0$.

- Recall that the trace is equal to the sum of the eigenvalues:

$$\text{tr}(\nabla T(\mathbf{x}^\dagger)) = \lambda_1 + \cdots + \lambda_N.$$

- Asymptotic stability requires that the real part $\Re(\lambda_n) < 0$ for all $n \in [N]$.
 - If some λ_n has a negative real part, then another λ_m must have positive real part. Then, the dynamics are **unstable**.

Barrier to Asymptotic Stability

- Suppose the dynamics around a fixed point is **traceless**: $\text{tr}(\nabla T(\mathbf{x}^\dagger)) = 0$.
 - Recall that the trace is equal to the sum of the eigenvalues:

$$\text{tr}(\nabla T(\mathbf{x}^\dagger)) = \lambda_1 + \cdots + \lambda_N.$$

- Asymptotic stability requires that the real part $\Re(\lambda_n) < 0$ for all $n \in [N]$.
 - If some λ_n has a negative real part, then another λ_m must have positive real part. Then, the dynamics are **unstable**.

This is the very high-level argument used across many works showing that certain dynamics fail to converge stably to interior Nash equilibria.

Stability from Spectral Analysis

- If all eigenvalues of ∇T have **negative real parts** $\implies$ asymptotic stability
- If there is an eigenvalue with **positive real part** $\implies$ instability
- If all eigenvalues are **purely imaginary** $\implies$ neutral stability

The remaining possibility

Neutral stability $\implies$ purely imaginary eigenvalues

Uniform Stability: Non-Asymptotic Stability

Definition. The game Jacobian $\mathbf{J}(\mathbf{x})$ is **uniformly stable** at $\mathbf{x}$ if the eigenvalues of $\mathbf{H}^{-1}\mathbf{J}(\mathbf{x})$ are purely imaginary for all positive-definite block-diagonal matrix $\mathbf{H}$,

$$\text{spec}(\mathbf{H}^{-1}\mathbf{J}(\mathbf{x})) \subset i\mathbb{R}.$$

Uniform Stability: Non-Asymptotic Stability

Definition. The game Jacobian $\mathbf{J}(\mathbf{x})$ is **uniformly stable** at $\mathbf{x}$ if the eigenvalues of $\mathbf{H}^{-1}\mathbf{J}(\mathbf{x})$ are purely imaginary for all positive-definite block-diagonal matrix $\mathbf{H}$,

$$\text{spec}(\mathbf{H}^{-1}\mathbf{J}(\mathbf{x})) \subset i\mathbb{R}.$$

- A Nash equilibrium $\mathbf{x}^\star$ is pointwise uniformly stable if $\mathbf{J}(\mathbf{x}^\star)$ is uniformly stable.

Uniform Stability: Non-Asymptotic Stability

Definition. The game Jacobian $\mathbf{J}(\mathbf{x})$ is **uniformly stable** at $\mathbf{x}$ if the eigenvalues of $\mathbf{H}^{-1}\mathbf{J}(\mathbf{x})$ are purely imaginary for all positive-definite block-diagonal matrix $\mathbf{H}$,

$$\text{spec}(\mathbf{H}^{-1}\mathbf{J}(\mathbf{x})) \subset i\mathbb{R}.$$

- A Nash equilibrium $\mathbf{x}^\star$ is pointwise uniformly stable if $\mathbf{J}(\mathbf{x}^\star)$ is uniformly stable.
- It is locally uniformly stable if there is an open set U containing $\mathbf{x}^\star$ such that $\mathbf{J}(\mathbf{x})$ is uniformly stable for all $\mathbf{x} \in U$.

Uniform Stability: Motivation

Definition. The game Jacobian $\mathbf{J}(\mathbf{x})$ is **uniformly stable** at $\mathbf{x}$ if the eigenvalues of $\mathbf{H}^{-1}\mathbf{J}(\mathbf{x})$ are purely imaginary for all positive-definite block-diagonal matrix $\mathbf{H}$,

$$\text{spec}(\mathbf{H}^{-1}\mathbf{J}(\mathbf{x})) \subset i\mathbb{R}.$$

- The imaginary spectrum captures the only chance for the dynamics to exhibit neutral stability (instead of instability).

Uniform Stability: Motivation

Definition. The game Jacobian $\mathbf{J}(\mathbf{x})$ is **uniformly stable** at $\mathbf{x}$ if the eigenvalues of $\mathbf{H}^{-1}\mathbf{J}(\mathbf{x})$ are purely imaginary for all positive-definite block-diagonal matrix $\mathbf{H}$,

$$\text{spec}(\mathbf{H}^{-1}\mathbf{J}(\mathbf{x})) \subset i\mathbb{R}.$$

- The imaginary spectrum captures the only chance for the dynamics to exhibit neutral stability (instead of instability).
- The quantifier “for all $\mathbf{H}$ ” expands the notion of *uncoupledness*: players do not know how others learn from their experiences.

Jacobian of learning dynamics

- **Smoothed best-response dynamics**

$$x_n(t+1) = \arg \max_{z_n \in \Omega_n} u_n(z_n; \mathbf{x}_{-n}(t)) - \phi_n(z_n)$$

- Usually ϕ_n is a smooth, strictly convex regularizer
 - Entropy map, negative-definite quadratic.
 - It defines the “local geometry” on Ω_n .

Jacobian of learning dynamics

- Smoothed best-response dynamics

$$x_n(t+1) = \arg \max_{z_n \in \Omega_n} u_n(z_n; \mathbf{x}_{-n}(t)) - \phi_n(z_n)$$

- Usually ϕ_n is a smooth, strictly convex regularizer

$$\nabla T(\mathbf{x}) = \mathbf{H}^{-1} \mathbf{J} \quad \text{where} \quad \mathbf{H} = \begin{bmatrix} \nabla^2 \phi_1 & & & \\ & \nabla^2 \phi_2 & & \\ & & \ddots & \\ & & & \nabla^2 \phi_N \end{bmatrix}$$

Technical Detail: Strategic Equivalence

All of the dynamics mentioned factor through *strategic equivalence*:

- Take two utilities u_n and $u'_n = u_n + v_n$ where $v_n(\mathbf{x}_{-n})$ does not depend on x_n . The dynamics for player n are the same for both. For example:

$$\nabla_n u_n \equiv \nabla_n (u_n + v_n).$$

- Parts of the utility can be “external” to the strategic component of the game.
- We focus on **purely strategic** utilities (or, define strategic Pareto optimality to depend only on the strategic component of the utility).

Technical Detail: Non-Degenerate Game

Our results hold under somewhat mild degeneracy assumptions.

- Consider the Taylor expansion of $u_n(\mathbf{x}) - u_n(\mathbf{x}^\star)$ about the equilibrium.
 - At the equilibrium, the linear terms are zero; lowest-order terms are *bilinear*.
 - The game Jacobian $\mathbf{J}$ precisely captures this bilinear information.
 - Our assumptions roughly make sure that the lower-order terms dominate.

The Meaning of Uniform Stability

An Algebraic Characterization

Theorem. Under non-degeneracy, the game Jacobian $\mathbf{J}$ is uniformly stable if and only if there are weights $\lambda_1, \dots, \lambda_N > 0$ so that $\mathbf{\Lambda J}$ is skew-symmetric where:

$$\mathbf{\Lambda} = \begin{bmatrix} \lambda_1 I & & & \\ & \lambda_2 I & & \\ & & \ddots & \\ & & & \lambda_N I \end{bmatrix}$$

- Thus, we can rescale the utilities $\tilde{u}_n = \lambda_n u_n$ and the corresponding Jacobian $\tilde{\mathbf{J}}$ is skew-symmetric.

Local Uniform Stability implies Pareto Optimality

Theorem. Assuming non-degeneracy. If $\mathbf{x}^\star$ be a mixed Nash equilibrium, then:

local uniform stability $\implies$ strategic Pareto optimality.

A simple class: polymatrix games

In a polymatrix game, each utility u_n decomposes as:

$$u_n(\mathbf{x}) = \sum_{m \neq n} (x_n - x_n^\star)^\top J_{nm} (x_m - x_m^\star).$$

- Players are simply playing multiple two-player games at once.
- There are only bilinear interactions, and no higher-order interactions.
- The game Jacobian fully characterizes (the strategic component of) the game.

Uniform Stability implies Pareto Optimality

Theorem (polymatrix version). Consider a non-degenerate, purely strategic, polymatrix game. Let $\mathbf{x}^\star$ be a mixed Nash equilibrium. Then:

uniform stability $\implies$ Pareto optimality.

Proof

1. By the algebraic characterization, there are $\lambda_1, \dots, \lambda_N > 0$ such that:

$$\lambda_n J_{nm} + \lambda_m J_{mn}^\top = 0.$$

2. It follows that the sum of rescaled utilities is zero:

$$\sum_{n \in [N]} \lambda_n u_n(\mathbf{x}) = \sum_{n \in [N]} \sum_{m \neq n} \lambda_n (x_n - x_n^\star)^\top J_{nm} (x_m - x_m^\star) = 0.$$

▸ The (n, m) -term in $\lambda_n J_{nm}$ cancels with the (m, n) -term of $\lambda_m J_{mn}$.

3. Thus, there is no way to choose $\mathbf{x}$ so that $u_n(\mathbf{x}) > 0$ for all $n \in [N]$.

▸ It follows that $\mathbf{x}^\star$ is Pareto optimal.

A comment for general games

In general-sum games, there may be higher-order interactions.

- It turns out that local uniform stability imposes a lot of structure on the utilities.
- Proof is quite non-trivial, but the general scheme is to construct functions $\lambda_1(\mathbf{x}), \dots, \lambda_N(\mathbf{x}) > 0$ such that in a neighborhood around $\mathbf{x}^\star$, we also have:

$$\sum_{n \in [N]} \lambda_n(\mathbf{x}) u_n(\mathbf{x}) = 0.$$

Pareto Optimality implies Uniform Stability

Theorem. Assume non-degeneracy. If $\mathbf{x}^\star$ is a mixed Nash equilibrium. Then:

Pareto optimality $\implies$ Uniform Stability.

Proof (Two-Player Case)

1. Recenter the utilities $u_1 - u_1^\star$ and $u_2 - u_2^\star$ so that $u_n(\mathbf{x}^\star) = 0$.
2. Pareto optimality implies that $\text{sign}(u_1) = -\text{sign}(u_2)$.
3. The zero set of u_1 and u_2 coincide.
4. These utilities are polynomials, so by Hilbert's nullstellensatz, they are scalar multiples of each other (up to scaling, the game is zero-sum).
5. Their Jacobians are also scalar multiples of each other.
6. Apply algebraic characterization of uniform stability.

* Technically, the nullstellensatz does not apply. The actual proof uses SVD and is only slightly longer.

A comment for general games

For N -player games, we need to show the following result:

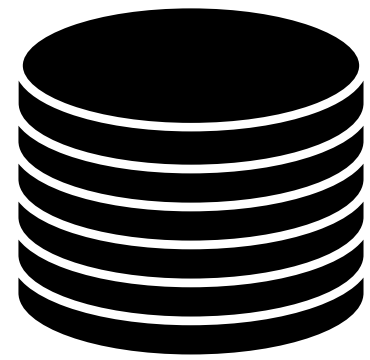
Theorem. Assume non-degeneracy. A mixed Nash equilibrium $\mathbf{x}^\star$ is weakly Pareto optimal if and only if it is a strong Nash equilibrium.

- A strong Nash equilibrium is one where coalitions of players cannot improve their utilities.
- The result proceeds by induction on coalitions.

RESEARCH OVERVIEW: TASK

Learning with Many Objectives

10,000 Foot View of Machine Learning



DATA

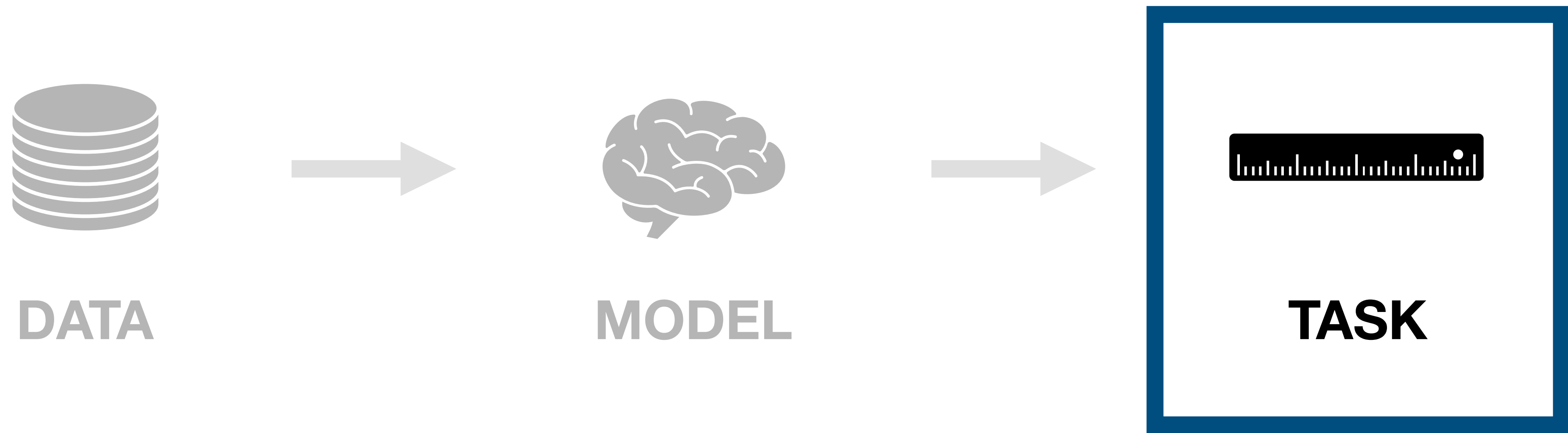


MODEL



TASK

10,000 Foot View of Machine Learning



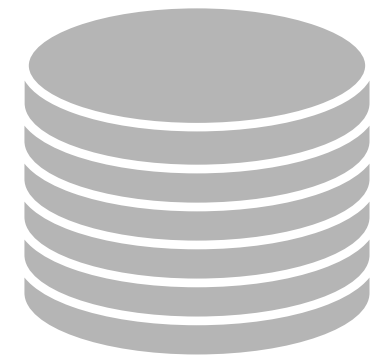
What happens if there is no clear way to define success?

Complex Tasks

- Autonomous Driving
- Chip Design
- Language Model
- Automatic Decision Making

⋮

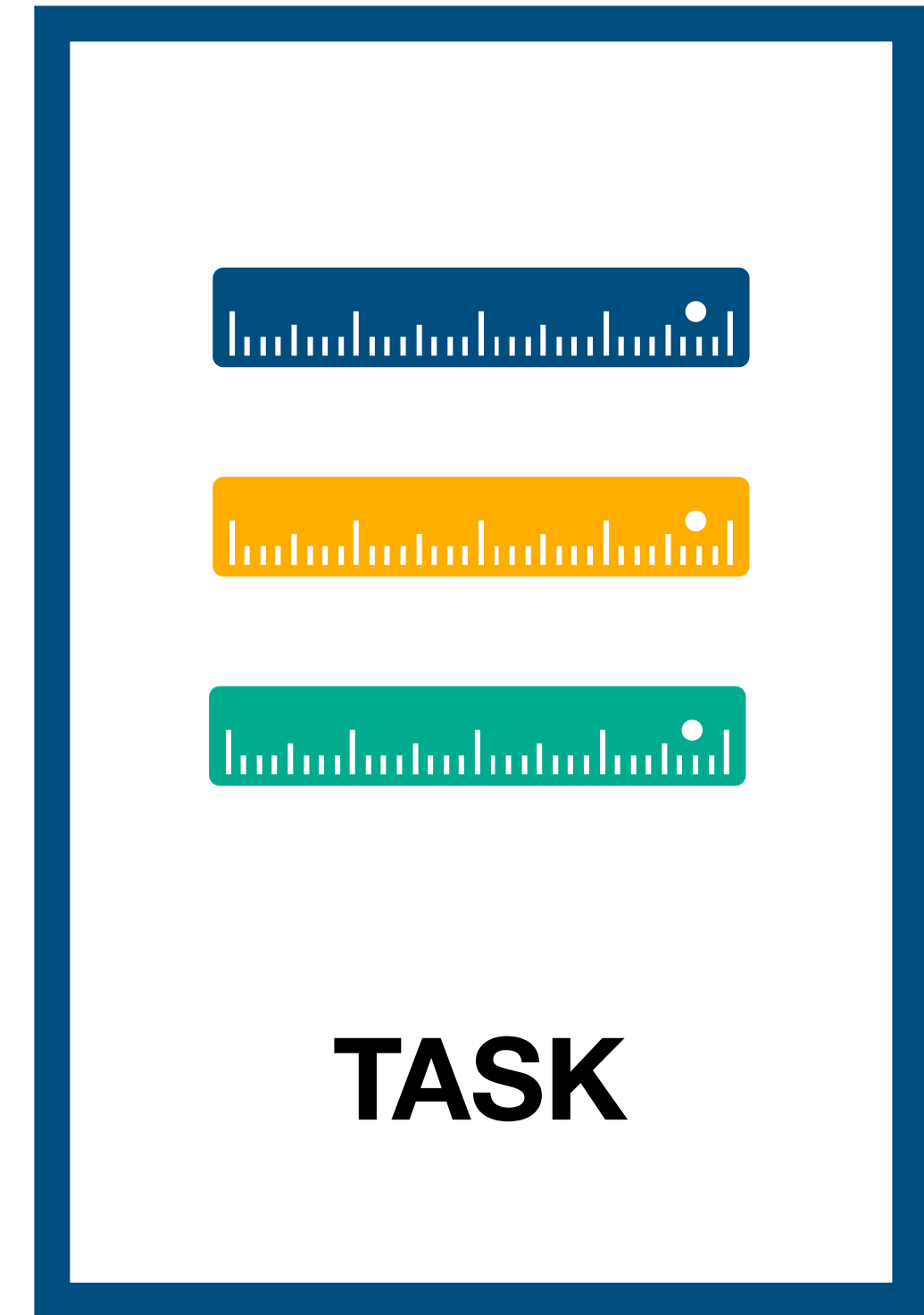
Complex Tasks



DATA



MODEL



There is often multi-objective structure to complex problems.

Complex Tasks

- Autonomous Driving
- Chip Design
- Language Model
- Automatic Decision Making
- ⋮

OBJECTIVES

- Safety
- Fuel Efficiency
- Comfort
- Legality

Complex Tasks

- Autonomous Driving
- Chip Design
- Language Model
- Automatic Decision Making
- ⋮

OBJECTIVES

- Efficiency
- Latency
- Cost
- Robustness

Complex Tasks

- Autonomous Driving
- Chip Design
- Language Model
- Automatic Decision Making
- ⋮

OBJECTIVES

- Informativeness
- Safety
- Succinctness
- Accuracy

Complex Tasks

- Autonomous Driving
- Chip Design
- Language Model
- Automatic Decision Making

⋮

OBJECTIVES

- Accuracy
- Individual Utility
- Social Welfare
- Fairness

Large-Scale Multi-Objective Optimization

- Modern ML settings require scalable optimization methods.

Large-Scale Multi-Objective Optimization

- Modern ML settings require scalable optimization methods.

Gradient descent is the standard in the [single-objective setting](#).

Large-Scale Multi-Objective Optimization

- Modern ML settings require scalable optimization methods.

Gradient descent is the standard in the [single-objective setting](#).

But there is **no multi-objective counterpart**.

Large-Scale Multi-Objective Optimization

- Modern ML settings require scalable optimization methods.

Gradient descent is the standard in the **single-objective setting**.

But there is **no multi-objective counterpart**.



(For **optimization on Pareto sets**).

Large-Scale Multi-Objective Optimization

- Modern ML settings require scalable optimization methods.

Preference Optimization on the Pareto Set: On the Theory of Multi-Objective Optimization

Abhishek Roy*, Geelon So*, Yi-An Ma, NeurIPS 2025

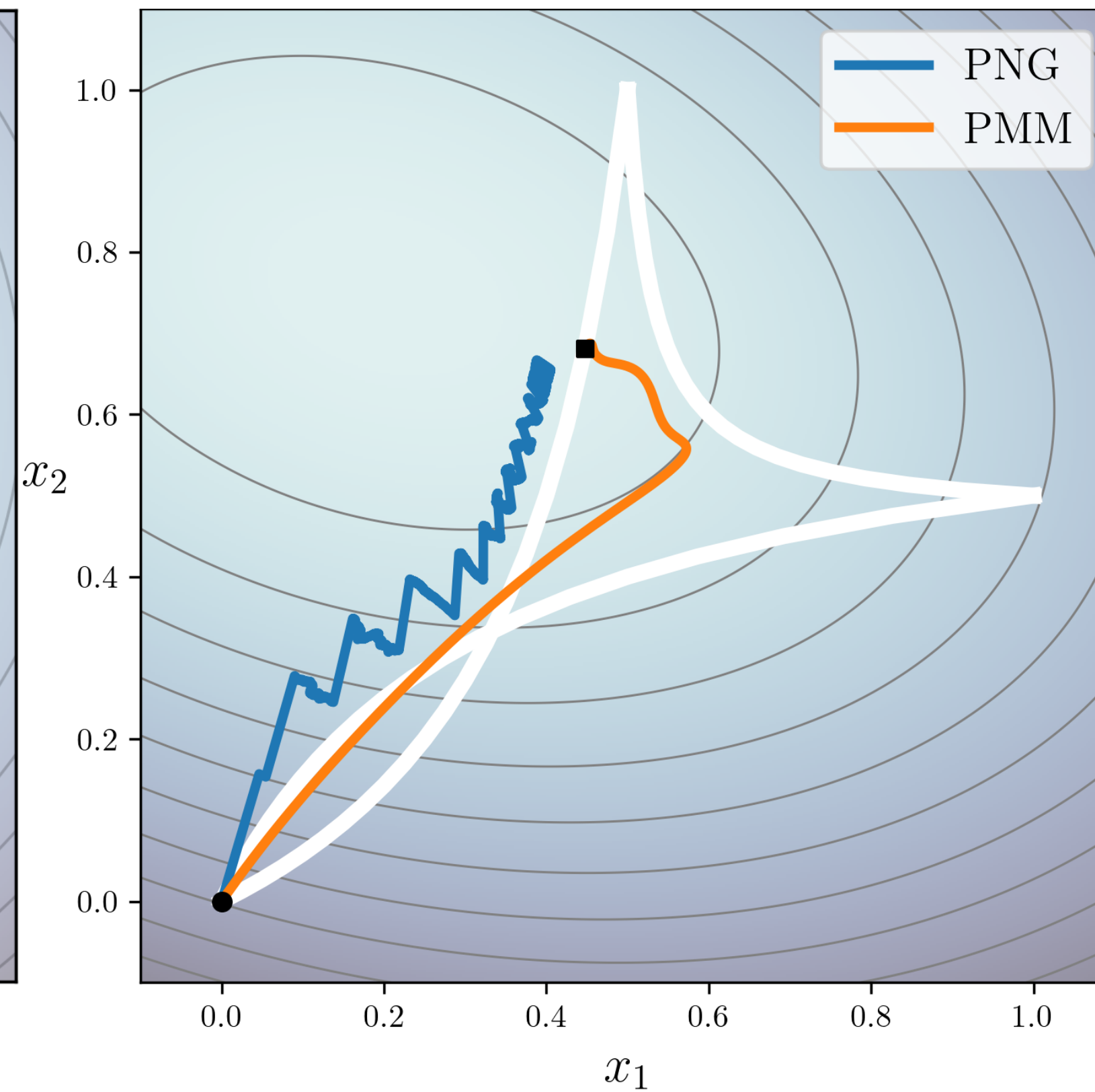
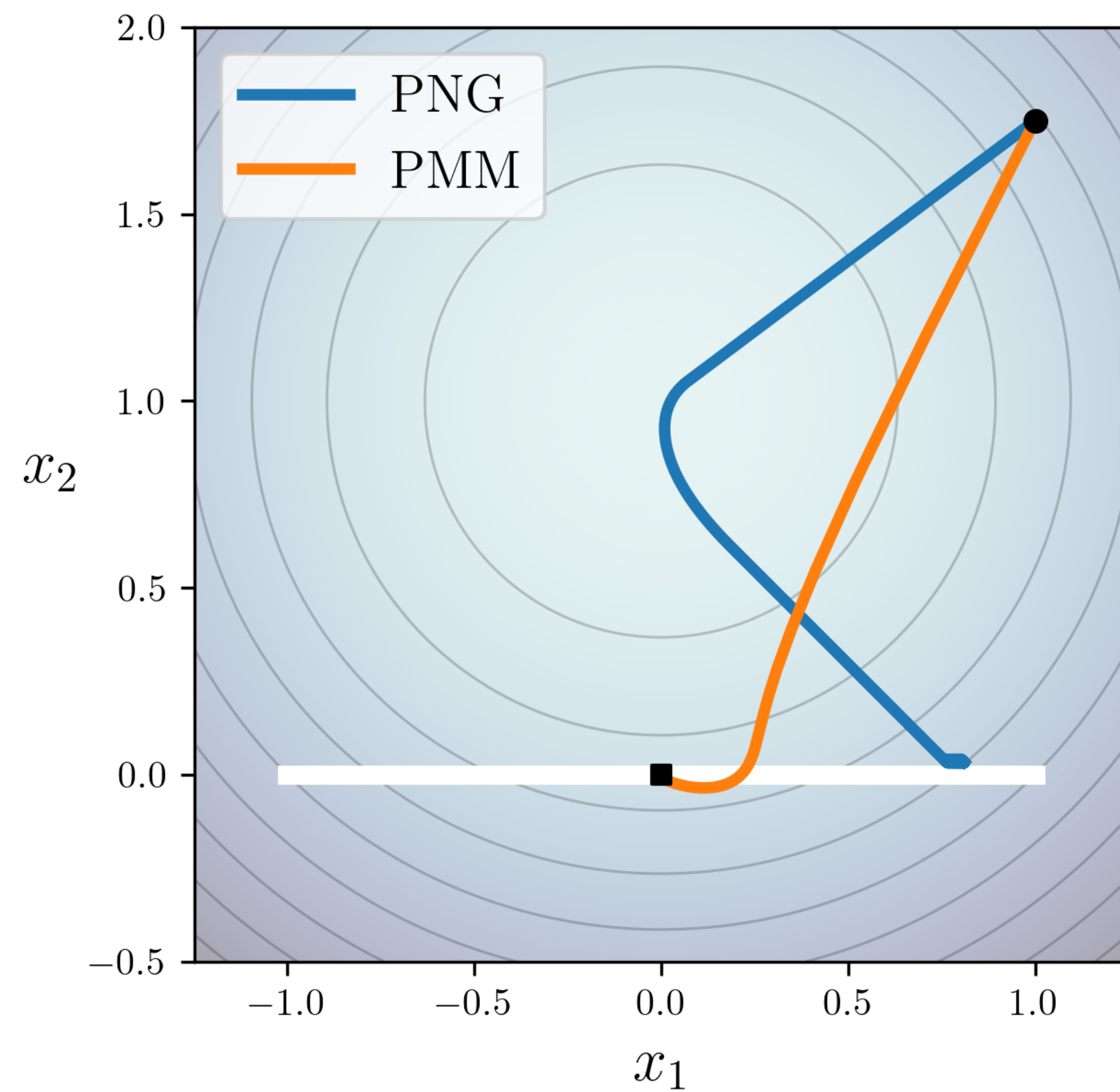
Outcomes

- Introduce a **solution concept** that is computationally-attainable and meaningful.
- Design a simple, **general-purpose optimization algorithm** (GD-like).
- Approach works with **human preference** feedback.

Preference Optimization on the Pareto Set: On the Theory of Multi-Objective Optimization

Abhishek Roy*, Geelon So*, Yi-An Ma, NeurIPS 2025

Comparison with Prior Method



Dynamics of **our algorithm** (PMM)
vs. **prior algorithm** (PNG).

Ground-truth is marked by a square.

Takeaway

- ML problems in practice are usually multi-objective.
- We need computationally-efficient optimization methods.

Other Works

- On the semi-supervised sample complexity of multi-objective learning (NeurIPS 2025)
- Transfer learning through multi-objective monotonicity (ICML 2026)