



DINFK

Detecting when the available data does not allow reliable inference

 Fanny Yang, Statistical Machine Learning Group
Department of Computer Science @ETH Zurich

joint work with students Alex Tifrea, Eric Stavarache,
Piersilvio de Bartolomeis, Javier A. Martinez, Konstantin Donhauser

ETH zürich



Problem of validity of inference

Problem of validity of inference



Problem of validity of inference



Problem **1**
too much hidden
confounding

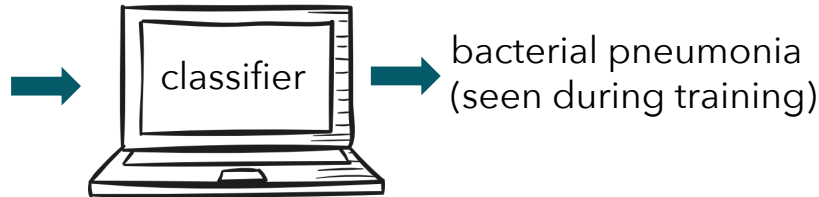
Problem of validity of inference



Problem **1**
too much hidden
confounding



individual test sample
with new disease/class



Problem of validity of inference

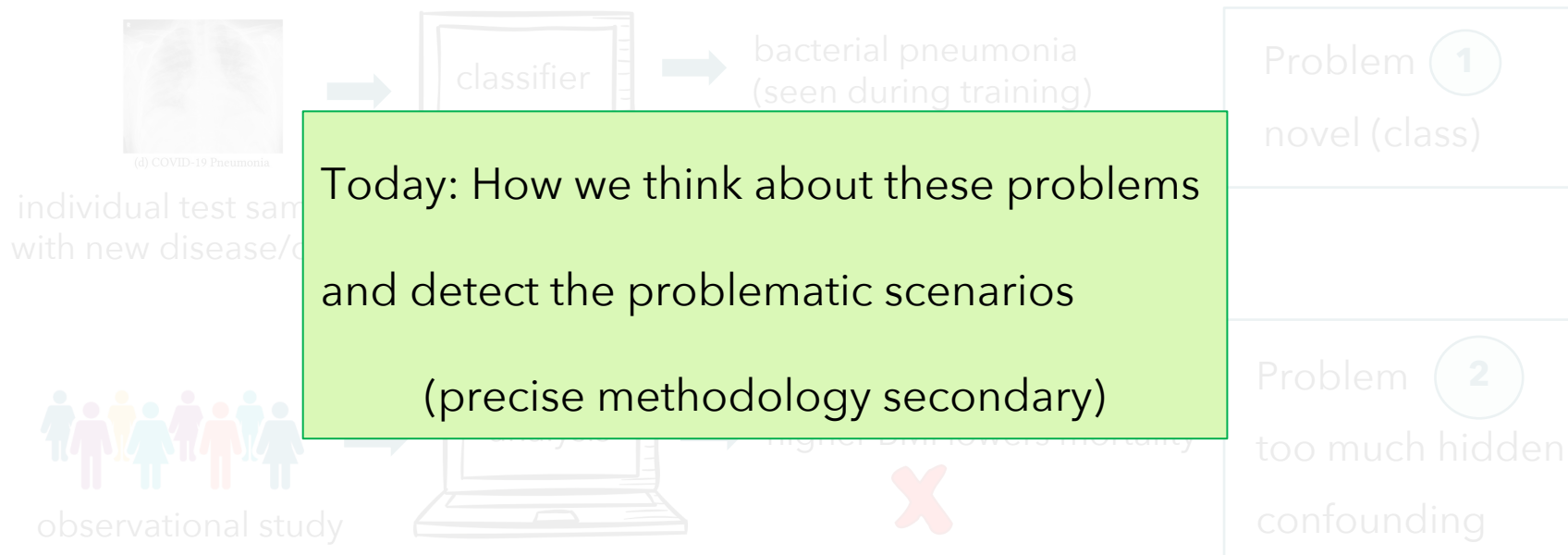


Problem 1
too much hidden
confounding



Problem 2
novel (class)

Problem of validity of inference





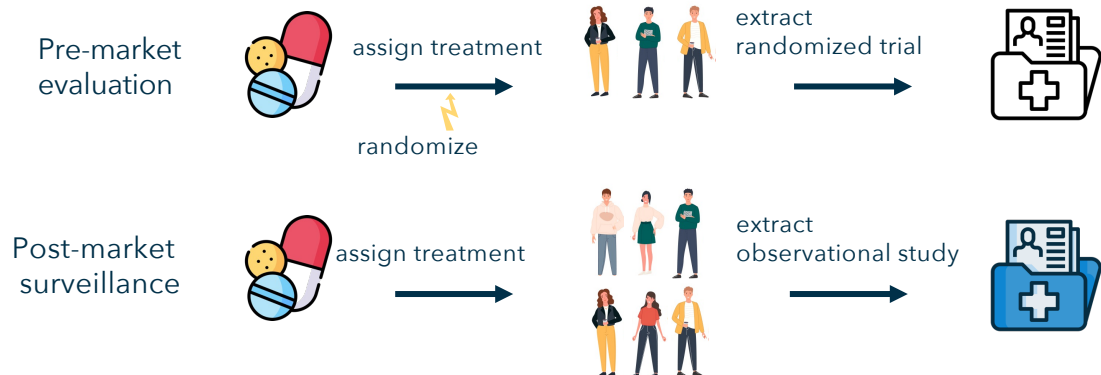
I. A lower bound for hidden confounding using randomized control trials

joint work with Piersilvio de Bartolomeis, Javier Abad Martinez, Konstantin Donhauser

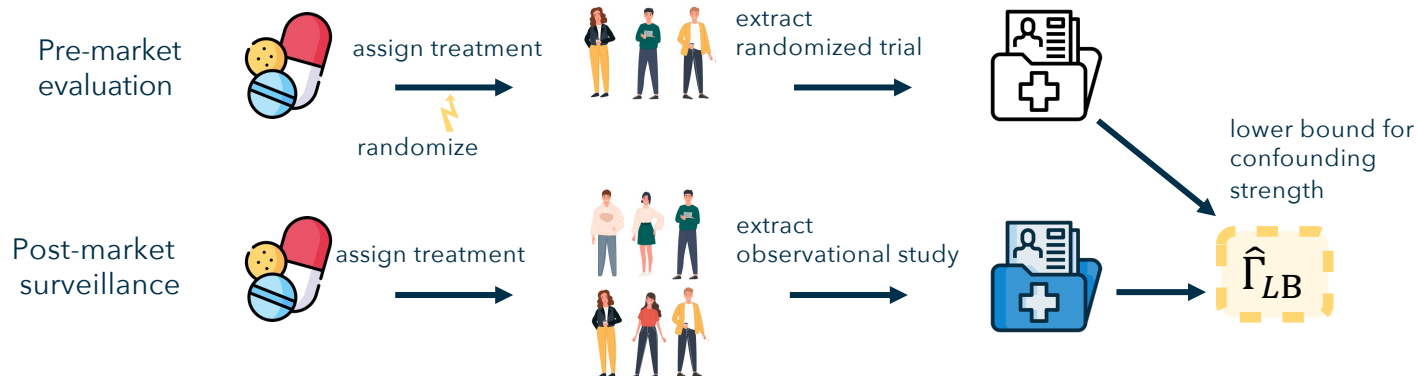
work in progress



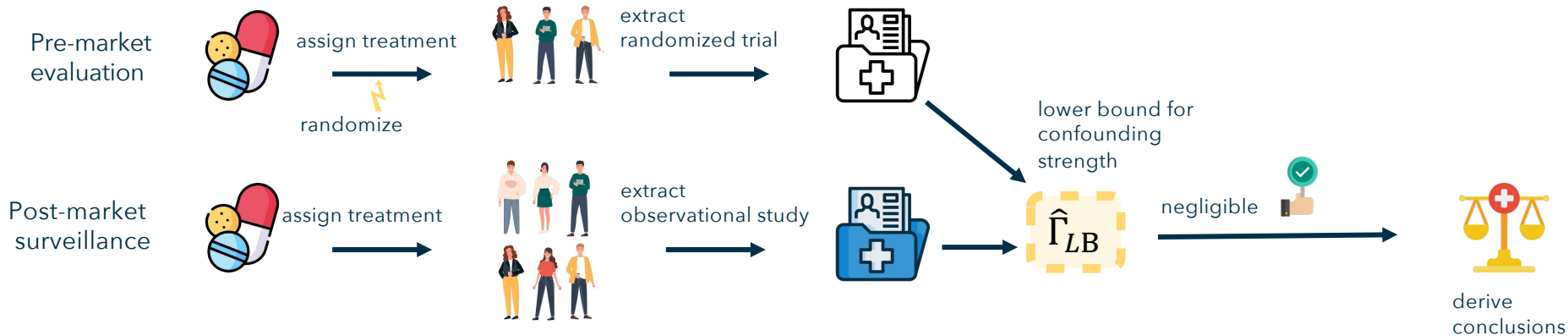
Our goal: lower bounding confounding strength



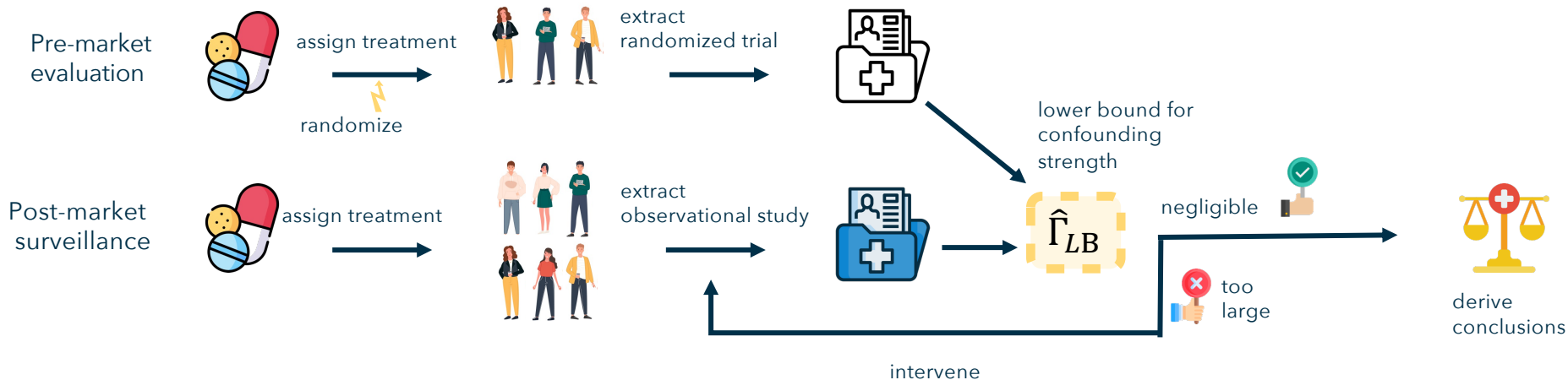
Our goal: lower bounding confounding strength



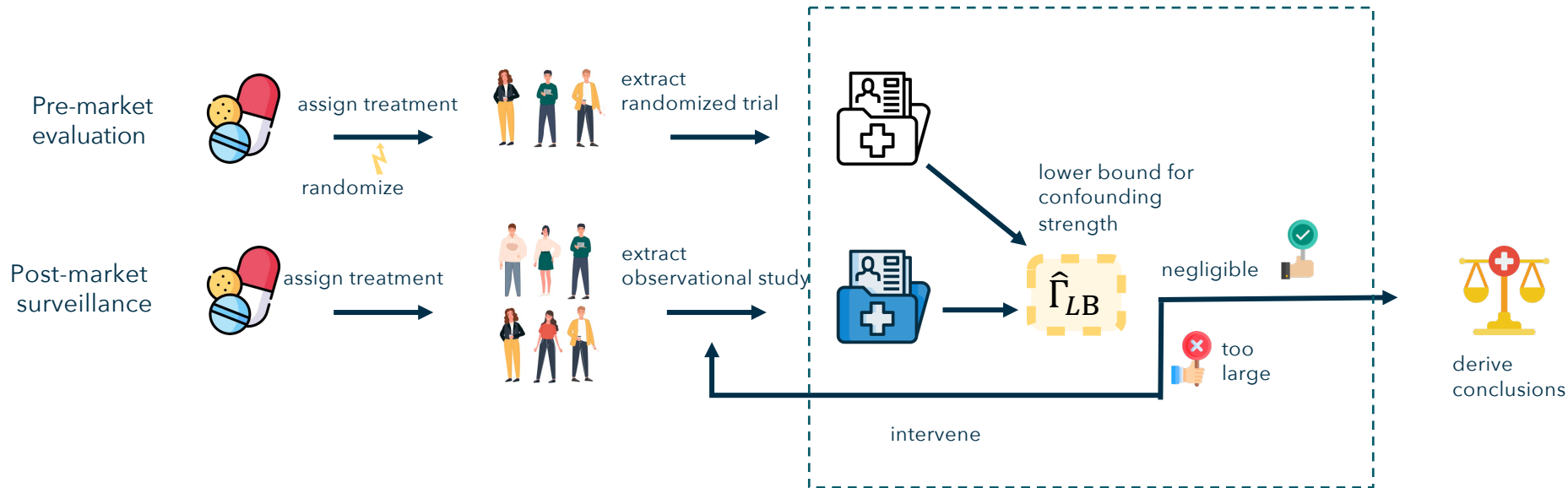
Our goal: lower bounding confounding strength



Our goal: lower bounding confounding strength



Our goal: lower bounding confounding strength



1. Definition of “strong” hidden confounding
2. Approach: How to detect it using RCT?

grey: observed variables,
white: latent variables

Potential outcome framework

Observed samples (X_i, Y_i, T_i) i.i.d. from the following distribution with $Y = Y(1)T + Y(0)(1 - T)$ (SUTVA)

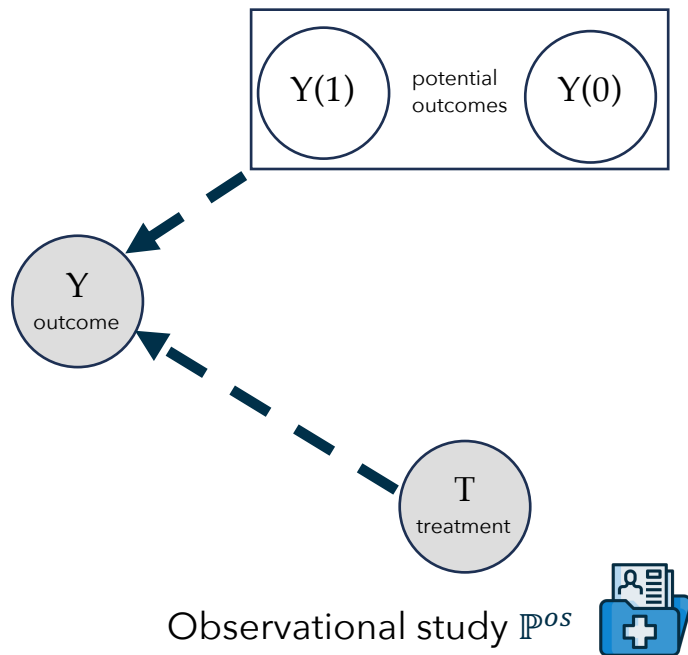
Observational study $\mathbb{P}^{os}$



grey: observed variables,
white: latent variables

Potential outcome framework

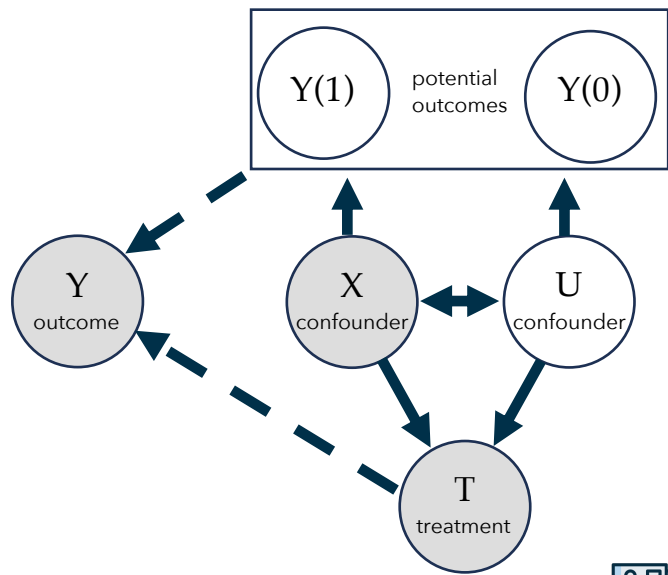
Observed samples (X_i, Y_i, T_i) i.i.d. from the following distribution with $Y = Y(1)T + Y(0)(1 - T)$ (SUTVA)



grey: observed variables,
white: latent variables

Potential outcome framework

Observed samples (X_i, Y_i, T_i) i.i.d. from the following distribution with $Y = Y(1)T + Y(0)(1 - T)$ (SUTVA)



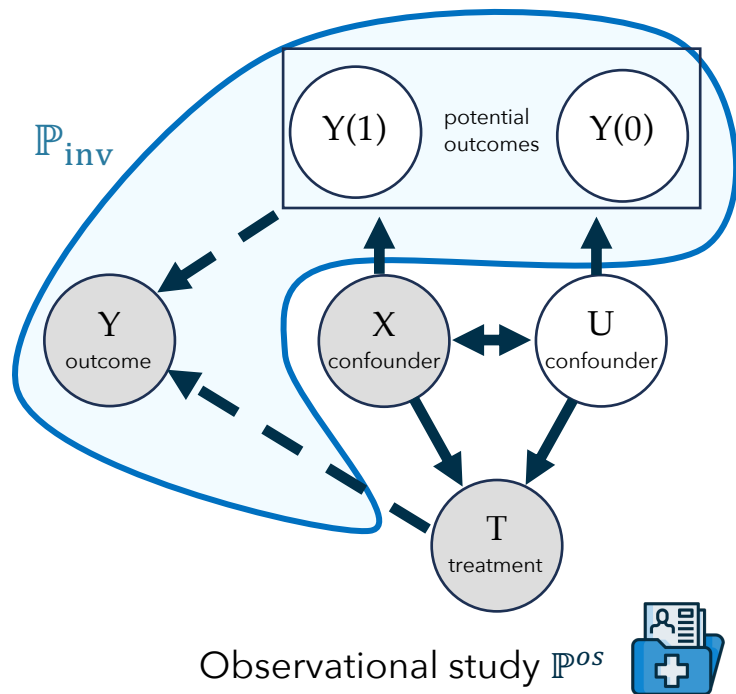
Observational study $\mathbb{P}^{os}$



grey: observed variables,
white: latent variables

Potential outcome framework

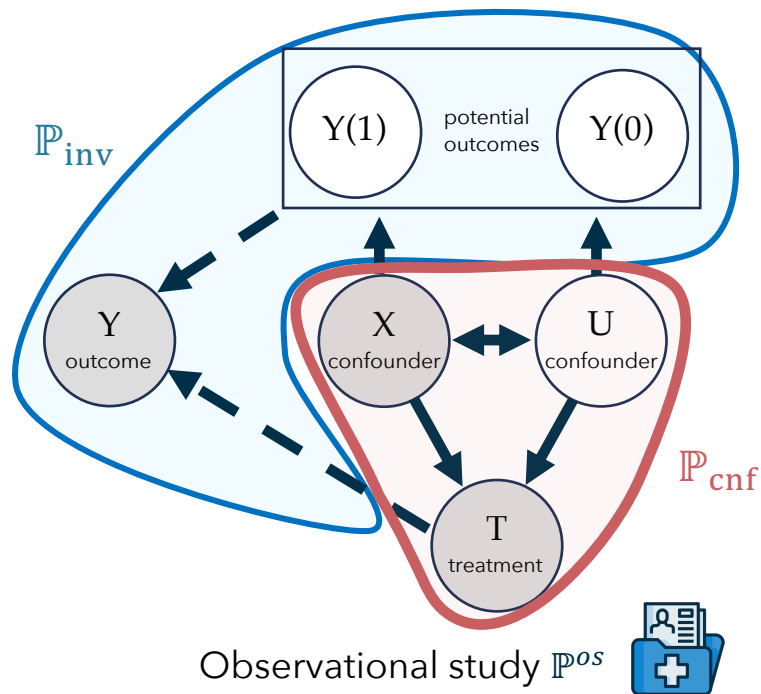
Observed samples (X_i, Y_i, T_i) i.i.d. from the following distribution with $Y = Y(1)T + Y(0)(1 - T)$ (SUTVA)



grey: observed variables,
white: latent variables

Potential outcome framework

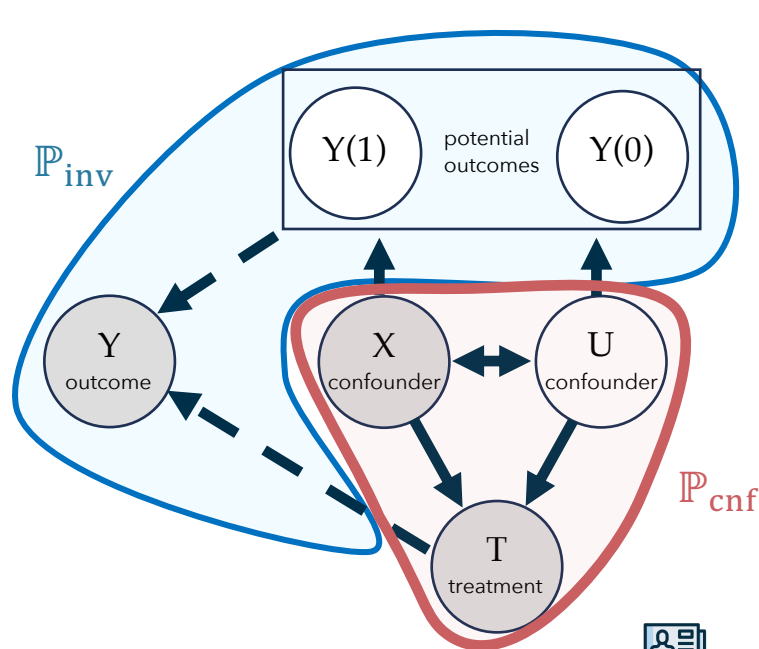
Observed samples (X_i, Y_i, T_i) i.i.d. from the following distribution with $Y = Y(1)T + Y(0)(1 - T)$ (SUTVA)



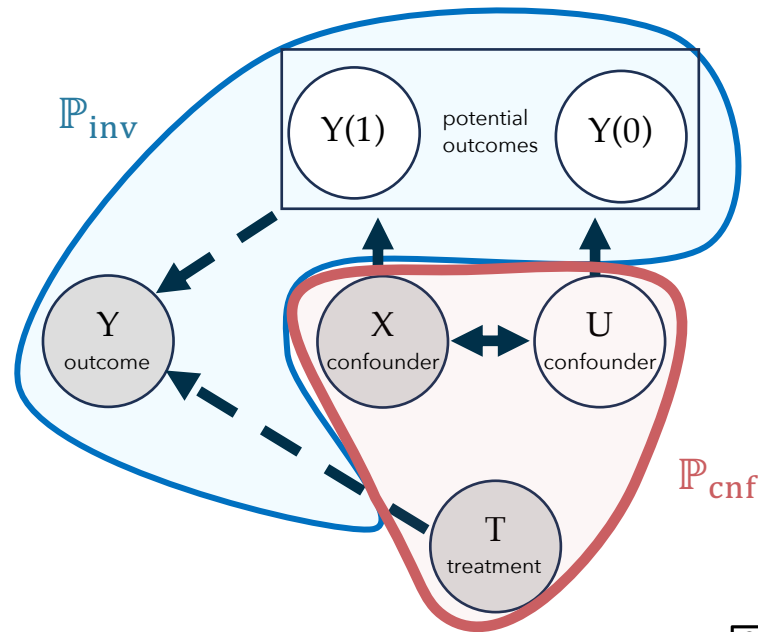
grey: observed variables,
white: latent variables

Potential outcome framework

Observed samples (X_i, Y_i, T_i) i.i.d. from the following distribution with $Y = Y(1)T + Y(0)(1 - T)$ (SUTVA)



Observational study $\mathbb{P}^{os}$



Randomized control trial $\mathbb{P}^{rct}$

Marginal sensitivity model

Additional assumptions:

- Transportability of CATE, i.e. $\mathbb{E}_{\mathbb{P}^{os}}[Y(1) - Y(0) \mid X] = \mathbb{E}_{\mathbb{P}^{rct}}[Y(1) - Y(0) \mid X]$
- Support inclusion $supp(\mathbb{P}^{rct}) \subseteq supp(\mathbb{P}^{os})$

Marginal sensitivity model

Additional assumptions:

- Transportability of CATE, i.e. $\mathbb{E}_{\mathbb{P}^{os}}[Y(1) - Y(0) \mid X] = \mathbb{E}_{\mathbb{P}^{rct}}[Y(1) - Y(0) \mid X]$
- Support inclusion $\text{supp}(\mathbb{P}^{rct}) \subseteq \text{supp}(\mathbb{P}^{os})$

Definitions:

- $\mathbb{P}^{os}$ satisfies MSM(Γ) if $\Gamma^{-1} \leq \frac{\mathbb{P}^{os}(T=1|X,U)}{\mathbb{P}^{os}(T=0|X,U)} / \frac{\mathbb{P}^{os}(T=1|X)}{\mathbb{P}^{os}(T=0|X)} \leq \Gamma$ almost surely (Tan-06)

Marginal sensitivity model

Additional assumptions:

- Transportability of CATE, i.e. $\mathbb{E}_{\mathbb{P}^{os}}[Y(1) - Y(0) | X] = \mathbb{E}_{\mathbb{P}^{rct}}[Y(1) - Y(0) | X]$
- Support inclusion $supp(\mathbb{P}^{rct}) \subseteq supp(\mathbb{P}^{os})$

Definitions:

- $\mathbb{P}^{os}$ satisfies MSM(Γ) if $\Gamma^{-1} \leq \frac{\mathbb{P}^{os}(T=1|X,U)}{\mathbb{P}^{os}(T=0|X,U)} / \frac{\mathbb{P}^{os}(T=1|X)}{\mathbb{P}^{os}(T=0|X)} \leq \Gamma$ almost surely (Tan-06)

$\Gamma = 1 \triangleq$ unconfoundedness



Marginal sensitivity model

Additional assumptions:

- Transportability of CATE, i.e. $\mathbb{E}_{\mathbb{P}^{os}}[Y(1) - Y(0) | X] = \mathbb{E}_{\mathbb{P}^{rct}}[Y(1) - Y(0) | X]$
- Support inclusion $supp(\mathbb{P}^{rct}) \subseteq supp(\mathbb{P}^{os})$

Definitions:

$\Gamma = 1 \triangleq$ unconfoundedness

- $\mathbb{P}^{os}$ satisfies MSM(Γ) if $\Gamma^{-1} \leq \frac{\mathbb{P}^{os}(T=1|X,U)}{\mathbb{P}^{os}(T=0|X,U)} / \frac{\mathbb{P}^{os}(T=1|X)}{\mathbb{P}^{os}(T=0|X)} \leq \Gamma$ almost surely (Tan-06)
- true confounding strength $\Gamma^*(\mathbb{P}^{os})$: The smallest Γ for which $\mathbb{P}^{os}$ satisfies MSM(Γ)

Marginal sensitivity model

Additional assumptions:

- Transportability of CATE, i.e. $\mathbb{E}_{\mathbb{P}^{os}}[Y(1) - Y(0) | X] = \mathbb{E}_{\mathbb{P}^{rct}}[Y(1) - Y(0) | X]$
- Support inclusion $supp(\mathbb{P}^{rct}) \subseteq supp(\mathbb{P}^{os})$

Definitions:

$\Gamma = 1 \triangleq$ unconfoundedness

- $\mathbb{P}^{os}$ satisfies MSM(Γ) if $\Gamma^{-1} \leq \frac{\mathbb{P}^{os}(T=1|X,U)}{\mathbb{P}^{os}(T=0|X,U)} / \frac{\mathbb{P}^{os}(T=1|X)}{\mathbb{P}^{os}(T=0|X)} \leq \Gamma$ almost surely (Tan-06)
- true confounding strength $\Gamma^*(\mathbb{P}^{os})$: The smallest Γ for which $\mathbb{P}^{os}$ satisfies MSM(Γ)

Scenarios we want to detect: when true confounding Γ^* of $\mathbb{P}^{os}$ is too large

1. Definition of “strong” hidden confounding
2. Approach: How to detect it using RCT?

Our paradigm: finding a lower bound

Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma): \mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$

Our paradigm: finding a lower bound

Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma)$: $\mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$
2. Report $\hat{\Gamma}_{LB} = \inf \{\Gamma: \phi_\alpha(\Gamma) = 0\}$ and flag if $\hat{\Gamma}_{LB} > \Gamma_{\text{thresh}}$

Our paradigm: finding a lower bound

Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma)$: $\mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$
2. Report $\hat{\Gamma}_{LB} = \inf \{\Gamma: \phi_\alpha(\Gamma) = 0\}$ and flag if $\hat{\Gamma}_{LB} > \Gamma_{\text{thresh}}$

- Test $H_0(\Gamma) \Rightarrow$ test whether $\mu \in [\mu_\Gamma^-, \mu_\Gamma^+]$ with ATE $\mu = \mathbb{E}_{\mathbb{P}}[Y(1) - Y(0)]$ and

ATE sensitivity bounds $\mu_\Gamma^- = \inf_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$, $\mu_\Gamma^+ = \sup_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$

Our paradigm: finding a lower bound

Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma)$: $\mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$
2. Report $\hat{\Gamma}_{LB} = \inf \{\Gamma: \phi_\alpha(\Gamma) = 0\}$ and flag if $\hat{\Gamma}_{LB} > \Gamma_{\text{thresh}}$

- Test $H_0(\Gamma) \Rightarrow$ test whether $\mu \in [\mu_\Gamma^-, \mu_\Gamma^+]$ with ATE $\mu = \mathbb{E}_\mathbb{P}[Y(1) - Y(0)]$ and

ATE sensitivity bounds $\mu_\Gamma^- = \inf_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$, $\mu_\Gamma^+ = \sup_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$

all full distributions that yield observed $\mathbb{P}_{X,Y,T}^{os}$ and satisfy $MSM(\Gamma)$

Our paradigm: finding a lower bound

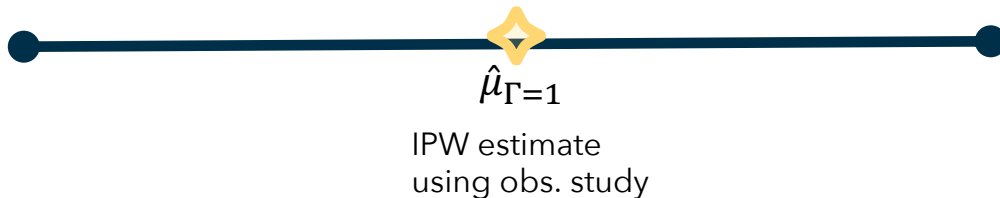
Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma)$: $\mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$
2. Report $\hat{\Gamma}_{LB} = \inf \{\Gamma: \phi_\alpha(\Gamma) = 0\}$ and flag if $\hat{\Gamma}_{LB} > \Gamma_{\text{thresh}}$

- Test $H_0(\Gamma) \Rightarrow$ test whether $\mu \in [\mu_\Gamma^-, \mu_\Gamma^+]$ with ATE $\mu = \mathbb{E}_\mathbb{P}[Y(1) - Y(0)]$ and

ATE sensitivity bounds $\mu_\Gamma^- = \inf_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$, $\mu_\Gamma^+ = \sup_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$

Test using consistent
estimates $\hat{\mu}, \hat{\mu}_\Gamma^+, \hat{\mu}_\Gamma^-$
in literature
(experts in audience)



Our paradigm: finding a lower bound

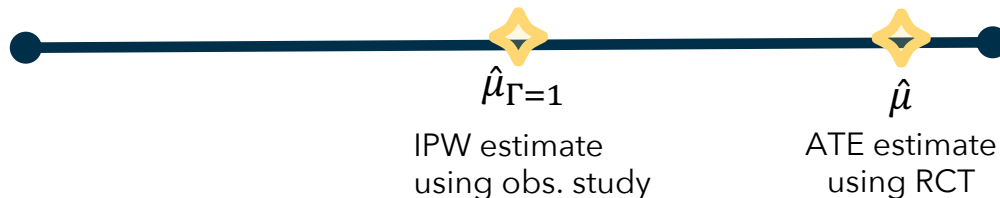
Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma)$: $\mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$
2. Report $\hat{\Gamma}_{LB} = \inf \{\Gamma: \phi_\alpha(\Gamma) = 0\}$ and flag if $\hat{\Gamma}_{LB} > \Gamma_{\text{thresh}}$

- Test $H_0(\Gamma) \Rightarrow$ test whether $\mu \in [\mu_\Gamma^-, \mu_\Gamma^+]$ with ATE $\mu = \mathbb{E}_{\mathbb{P}}[Y(1) - Y(0)]$ and

ATE sensitivity bounds $\mu_\Gamma^- = \inf_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$, $\mu_\Gamma^+ = \sup_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$

Test using consistent estimates $\hat{\mu}, \hat{\mu}_\Gamma^+, \hat{\mu}_\Gamma^-$ in literature (experts in audience)



Our paradigm: finding a lower bound

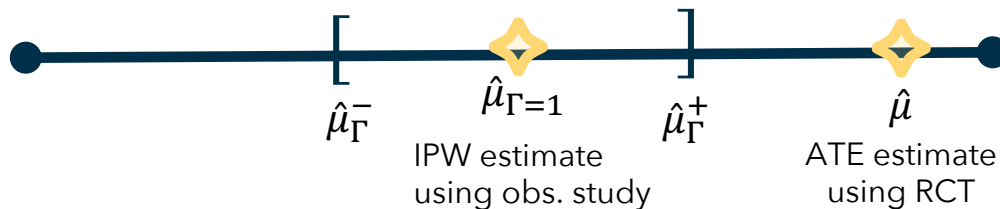
Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma)$: $\mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$
2. Report $\hat{\Gamma}_{LB} = \inf \{\Gamma: \phi_\alpha(\Gamma) = 0\}$ and flag if $\hat{\Gamma}_{LB} > \Gamma_{\text{thresh}}$

- Test $H_0(\Gamma) \Rightarrow$ test whether $\mu \in [\mu_\Gamma^-, \mu_\Gamma^+]$ with ATE $\mu = \mathbb{E}_{\mathbb{P}}[Y(1) - Y(0)]$ and

ATE sensitivity bounds $\mu_\Gamma^- = \inf_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$, $\mu_\Gamma^+ = \sup_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$

Test using consistent estimates $\hat{\mu}, \hat{\mu}_\Gamma^+, \hat{\mu}_\Gamma^-$ in literature (experts in audience)



Our paradigm: finding a lower bound

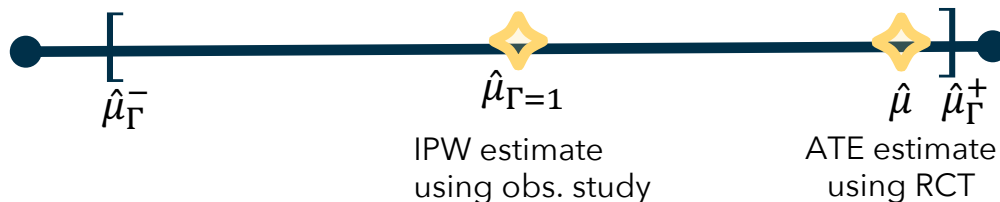
Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma)$: $\mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$
2. Report $\hat{\Gamma}_{LB} = \inf \{\Gamma: \phi_\alpha(\Gamma) = 0\}$ and flag if $\hat{\Gamma}_{LB} > \Gamma_{\text{thresh}}$

- Test $H_0(\Gamma) \Rightarrow$ test whether $\mu \in [\mu_\Gamma^-, \mu_\Gamma^+]$ with ATE $\mu = \mathbb{E}_{\mathbb{P}}[Y(1) - Y(0)]$ and

ATE sensitivity bounds $\mu_\Gamma^- = \inf_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$, $\mu_\Gamma^+ = \sup_{\tilde{\mathbb{P}} \in P_\Gamma(\mathbb{P}_{X,Y,T}^{os})} \mathbb{E}_{\tilde{\mathbb{P}}}[Y(1) - Y(0)]$

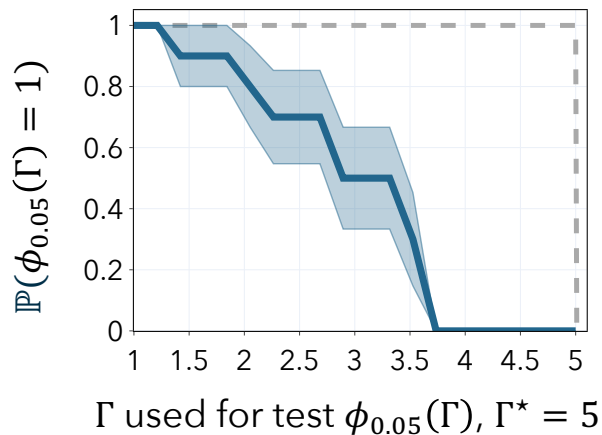
Test using consistent estimates $\hat{\mu}, \hat{\mu}_\Gamma^+, \hat{\mu}_\Gamma^-$ in literature (experts in audience)



Our paradigm: finding a lower bound

Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma)$: $\mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$
2. Report $\hat{\Gamma}_{LB} = \inf \{\Gamma: \phi_\alpha(\Gamma) = 0\}$ and flag if $\hat{\Gamma}_{LB} > \Gamma_{\text{thresh}}$

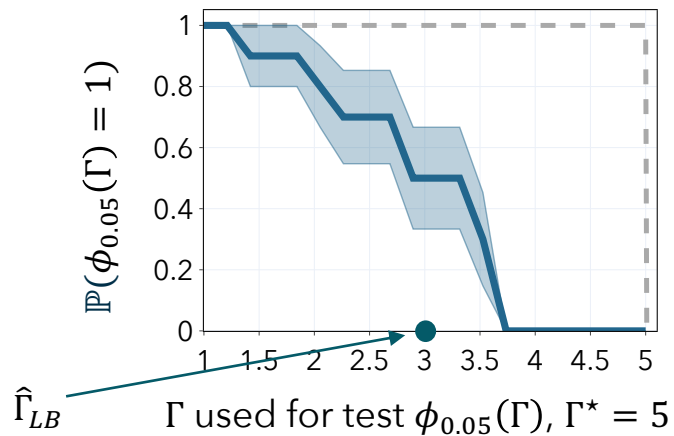


- probability of rejection over 20 runs on semi-synthetic data

Our paradigm: finding a lower bound

Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma)$: $\mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$
2. Report $\hat{\Gamma}_{LB} = \inf \{\Gamma: \phi_\alpha(\Gamma) = 0\}$ and flag if $\hat{\Gamma}_{LB} > \Gamma_{\text{thresh}}$

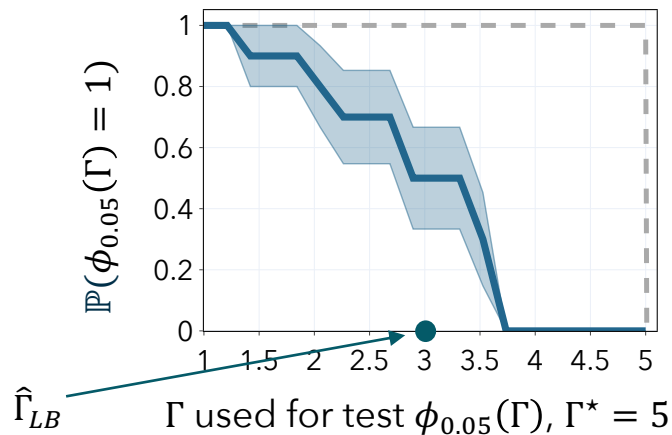


- probability of rejection over 20 runs on semi-synthetic data

Our paradigm: finding a lower bound

Our plug-and-play approach for desired significance α :

1. Test $\phi_\alpha(\Gamma)$ of the null $H_0(\Gamma)$: $\mathbb{P}^{os}$ satisfies $MSM(\Gamma) \Leftrightarrow \Gamma^* \leq \Gamma$
2. Report $\hat{\Gamma}_{LB} = \inf \{\Gamma: \phi_\alpha(\Gamma) = 0\}$ and flag if $\hat{\Gamma}_{LB} > \Gamma_{\text{thresh}}$



- probability of rejection over 20 runs on semi-synthetic data
- asymptotically $\mathbb{P}(\hat{\Gamma}_{LB} > \Gamma^*) \leq \alpha$ implied by $\mathbb{P}(\phi_\alpha(\Gamma^*) = 1) \leq \alpha$

Previous paradigms that could be used to “flag”

- without RCT and using sensitivity bounds
 - quantification via critical value $\hat{\Gamma}_{ct}$ that changes causal conclusions
e.g. vanderWeele-Ding-17, Jin-Ren-Candes-23 etc.

Previous paradigms that could be used to “flag”

- without RCT and using sensitivity bounds
 - quantification via critical value $\hat{\Gamma}_{ct}$ that changes causal conclusions
e.g. vanderWeele-Ding-17, Jin-Ren-Candes-23 etc.

but: unclear relation to true Γ^*

Previous paradigms that could be used to “flag”

- without RCT and using sensitivity bounds
 - quantification via critical value $\hat{\Gamma}_{ct}$ that changes causal conclusions
e.g. vanderWeele-Ding-17, Jin-Ren-Candes-23 etc.
but: unclear relation to true Γ^*
 - can test joint null hypothesis $ATE(\text{obs. study}) > 0$ and $MSM(\Gamma)$ holds
e.g. Yablowsky-Namkoong-Basu-Duchi-Tian-22, Jin-Ren-Candes-23

Previous paradigms that could be used to “flag”

- without RCT and using sensitivity bounds
 - quantification via critical value $\hat{\Gamma}_{ct}$ that changes causal conclusions
e.g. vanderWeele-Ding-17, Jin-Ren-Candes-23 etc.
but: unclear relation to true Γ^*
 - can test joint null hypothesis $ATE(\text{obs. study}) > 0$ and $MSM(\Gamma)$ holds
e.g. Yablowsky-Namkoong-Basu-Duchi-Tian-22, Jin-Ren-Candes-23
but: rejection only means either $MSM(\Gamma)$ assumption wrong or $ATE \leq 0$

Previous paradigms that could be used to “flag”

- without RCT and using sensitivity bounds
 - quantification via critical value $\hat{\Gamma}_{ct}$ that changes causal conclusions
e.g. vanderWeele-Ding-17, Jin-Ren-Candes-23 etc.
but: unclear relation to true Γ^*
 - can test joint null hypothesis $ATE(\text{obs. study}) > 0$ and $MSM(\Gamma)$ holds
e.g. Yablowsky-Namkoong-Basu-Duchi-Tian-22, Jin-Ren-Candes-23
but: rejection only means either $MSM(\Gamma)$ assumption wrong or $ATE \leq 0$
- with RCT:
 - binary test for existence of confounding with $H_0: \Gamma^* > 1$
e.g. Viele et al '14, Hussein-Oberst-Shih-Sontag '22

Previous paradigms that can be used for detection

- without RCT and using sensitivity bounds

- quantification via critical gamma value $\hat{\Gamma}_{ct}$ that changes causal conclusions

e.g. vanderWeide et al. 2023 etc.

but: can't test statement about Γ^*

- can test statement about Γ^* not possible!

e.g. Yablowsky-Namkoong-Basu-Duchi-Tian-22, Jin et al. 2023

but: rejection only means either MSM(Γ) assumption

→ true statement
about Γ^* not possible!



our paradigm:
statement about Γ^*
& flag only if Γ^* large

- with RCT:

- binary test for confounding with $H_0: \Gamma^* > 1$

e.g. Viele et al. 2022, Sontag '22

→ flag even if Γ^* small

Evaluation on real-world data (WHI)

- Randomized trial and observational study run by the NHLBI (1993-2005)
- Treatment: hormone replacement therapy
- Outcomes: coronary heart disease

Evaluation on real-world data (WHI)

- ▶ Randomized trial and observational study run by the NHLBI (1993-2005)
- ▶ Treatment: hormone replacement therapy
- ▶ Outcomes: coronary heart disease
- ▶ hidden confounder (revealed later): start of treatment



Evaluation on real-world data (WHI)

- Randomized trial and observational study run by the NHLBI (1993-2005)
- Treatment: hormone replacement therapy
- Outcomes: coronary heart disease
- hidden confounder (revealed later): start of treatment



treated	Coronary heart disease	
	as trial started	before trial
$\hat{\Gamma}_{CT}$	1.017	1.164
$\hat{\Gamma}_{LB}$	1.009	1.224
ψ_{bin}	1	1
ψ_{sens}	0	1

Different paradigms for “flagging” confounding:

- Compute $\hat{\Gamma}_{CT}$ that changes ATE sign and compare let “expert” assess “likeliness”
- ψ_{bin} : tests for existence, e.g. check $\hat{\Gamma}_{LB} > 1$
- ψ_{sens} (ours): check whether too large $\hat{\Gamma}_{LB} > \hat{\Gamma}_{CT}$

Current and future work

Higher power using

- kernelized test as opposed to averaging
- non-"adversarial" sensitivity model

Current and future work

Higher power using

- kernelized test as opposed to averaging
- non-“adversarial” sensitivity model

Extended applicability:

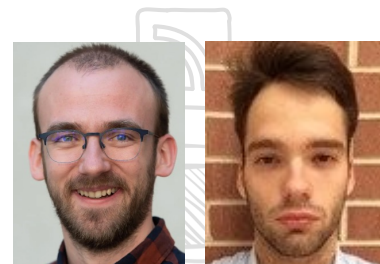
- multiple observational studies (no RCT)
- Automatic detection of hidden confounders from set of features



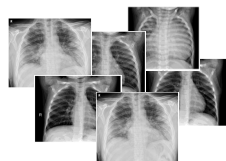
II. Semi-supervised novelty detection using ensembles with regularized disagreement

joint work with Alexandru Tifrea, Eric Stavarache

published at UAI '22



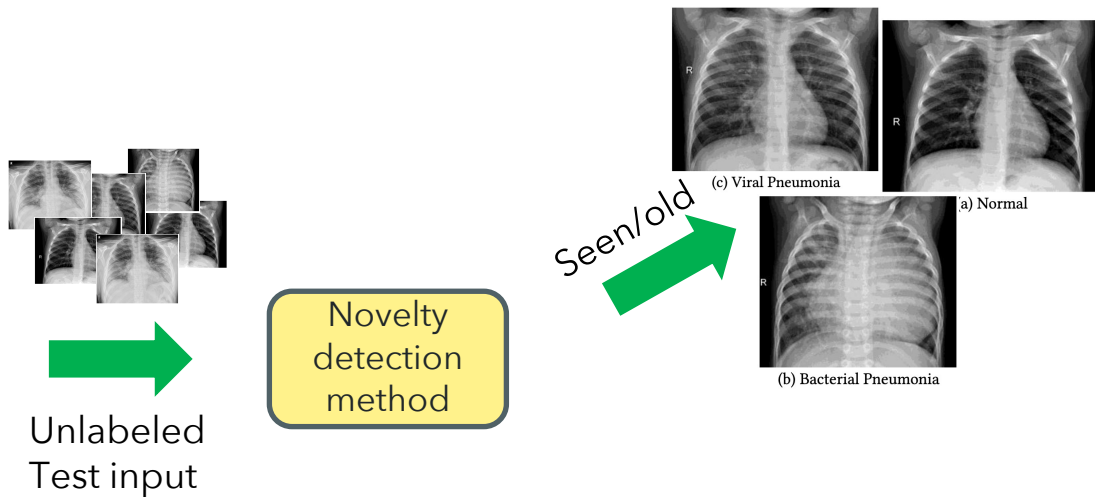
The novelty detection problem for classification



Unlabeled
Test input

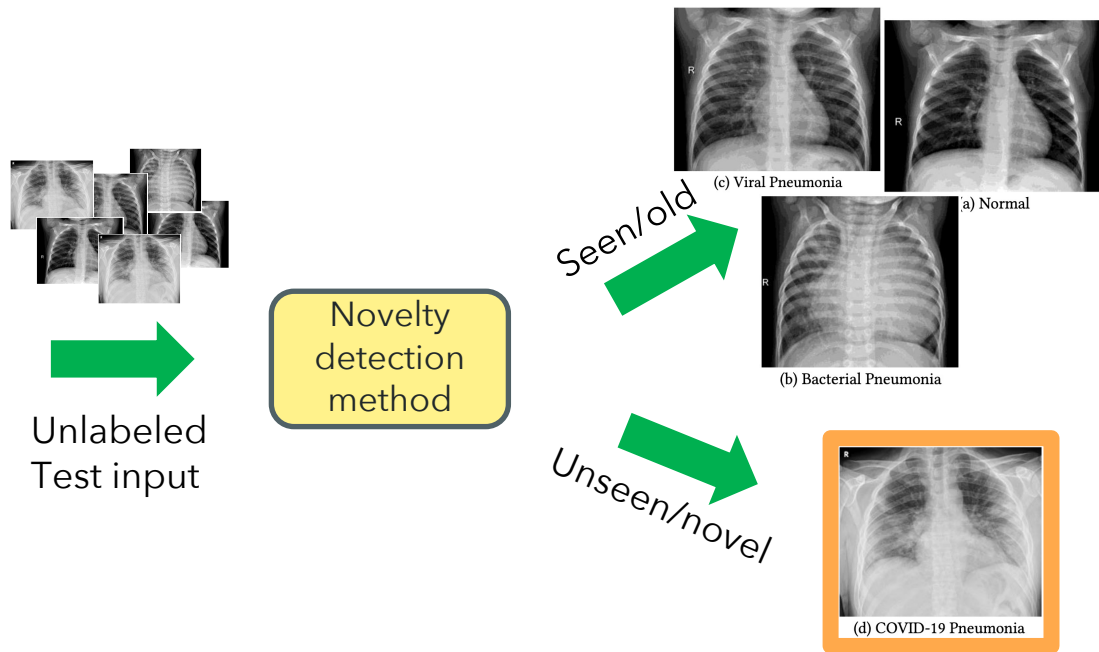
Novelty
detection
method

The novelty detection problem for classification



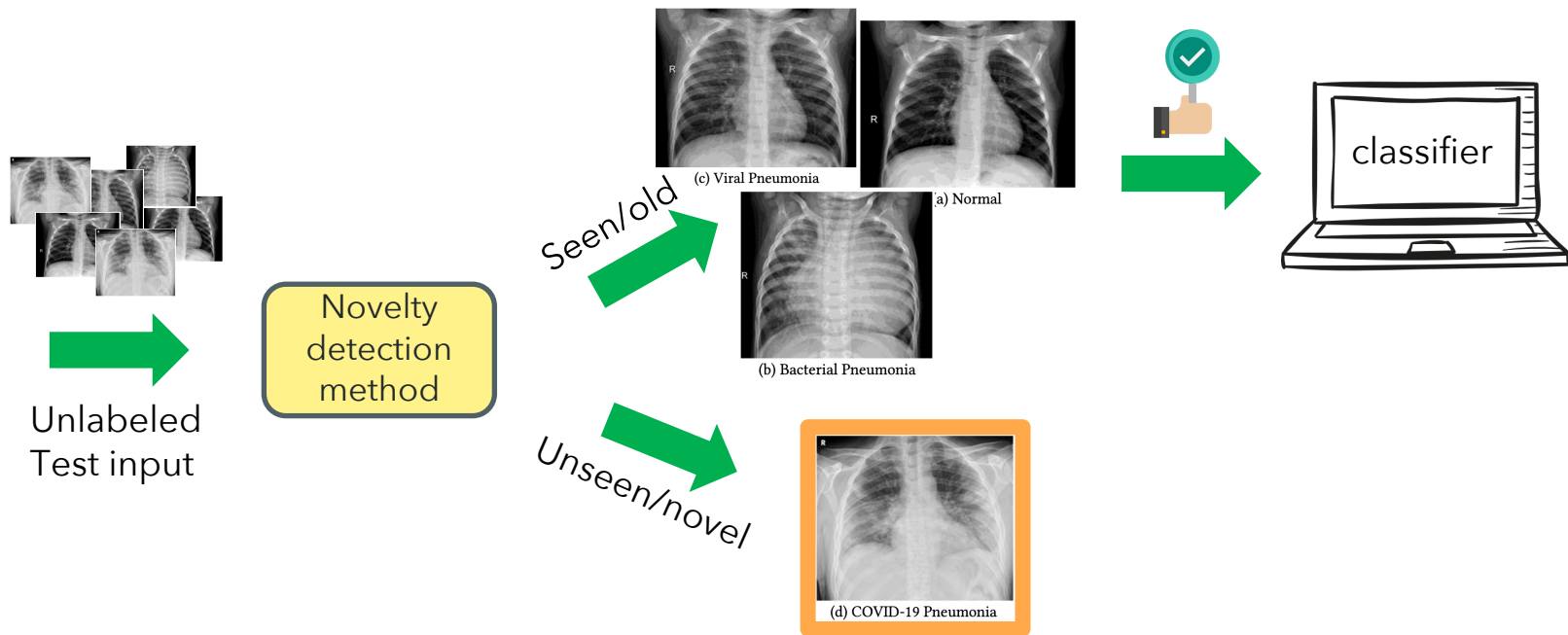
The novelty detection problem for classification

Novelty detection method tells user that software doesn't "know enough" to **predict** new point



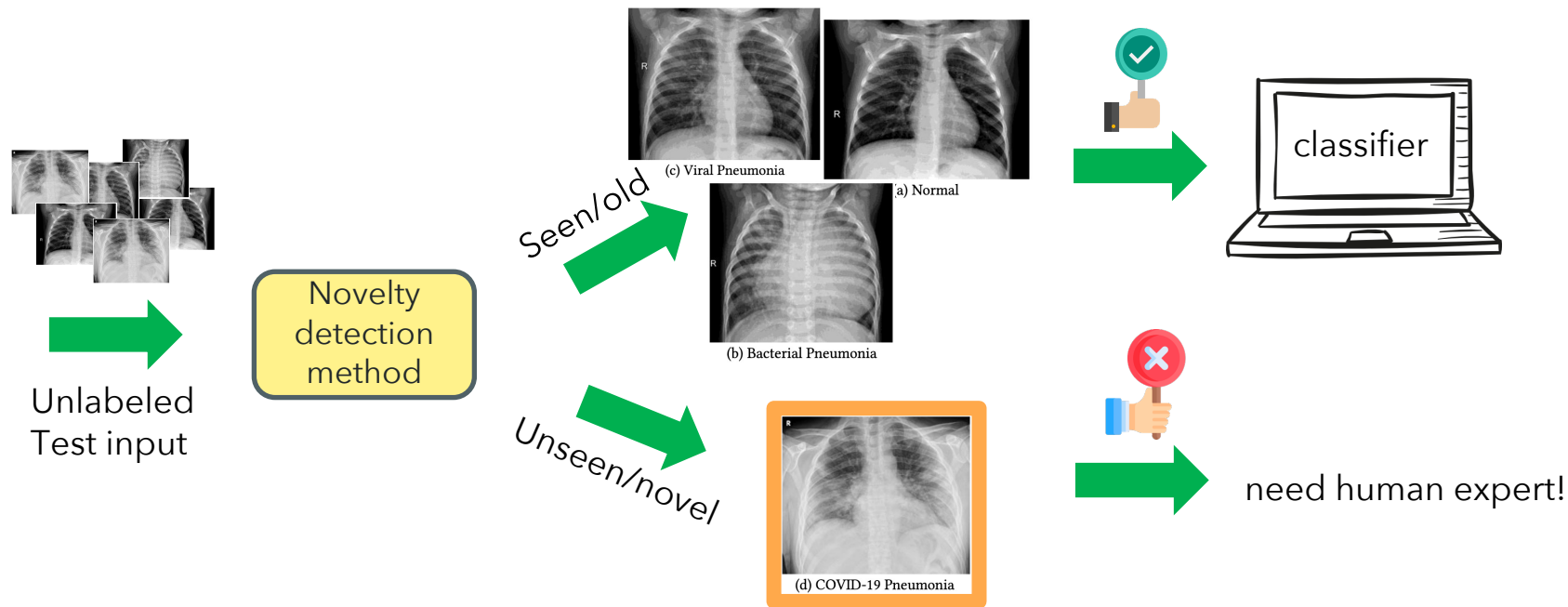
The novelty detection problem for classification

Novelty detection method tells user that software doesn't "know enough" to **predict** new point



The novelty detection problem for classification

Novelty detection method tells user that software doesn't "know enough" to **predict** new point



1. Definition: Points we can't make inference on
2. Approach: How to detect those samples?

What's "novel" to a trained model?

"novel" / o.o.d. points: test points $x \in X$ the model cannot reliably predict.

What's "novel" to a trained model?

"novel" / o.o.d. points: test points $x \in X$ the model cannot reliably predict.

First: which points $x \in X$ can a model predict "reliably" in an unseen test set?

- i.d. generalization from finite samples (traditional learning theory) and

What's "novel" to a trained model?

"novel" / o.o.d. points: test points $x \in X$ the model cannot reliably predict.

First: which points $x \in X$ can a model predict "reliably" in an unseen test set?

- i.d. generalization from finite samples (traditional learning theory) and
- o.o.d. generalization (**extrapolatable** from training distribution) -
depends on *test shift & model complexity*

Illustration: Extrapolatable vs. novel samples

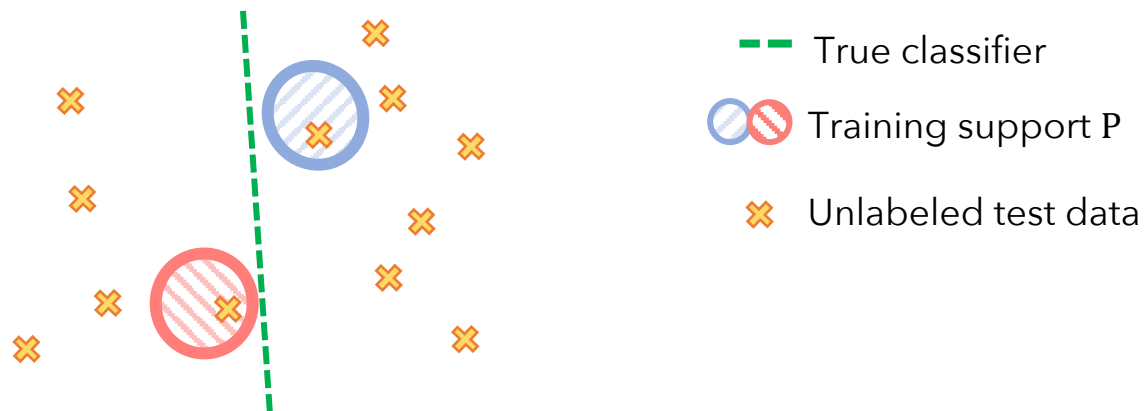


Illustration: Extrapolatable vs. novel samples

Extrapolatable given training distribution + linear ground truth:

Points $x \in X$ where the set of all linear Bayes optimal classifiers agree on

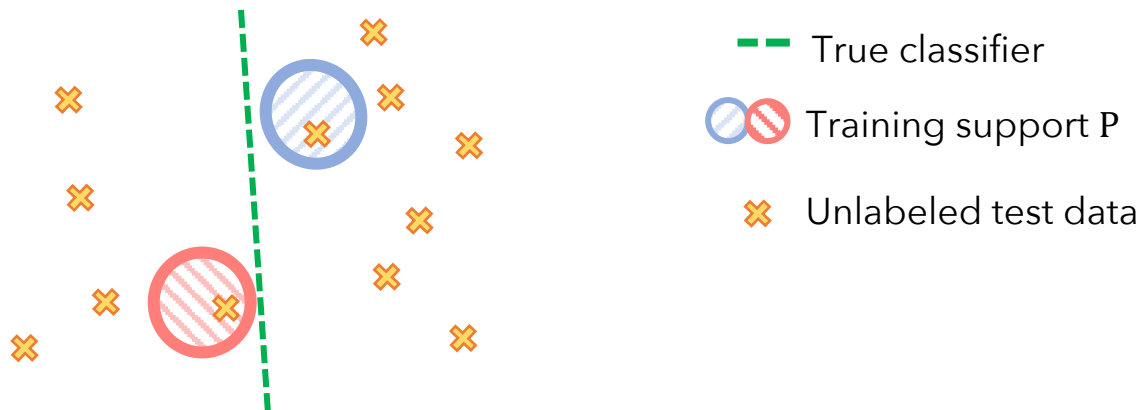
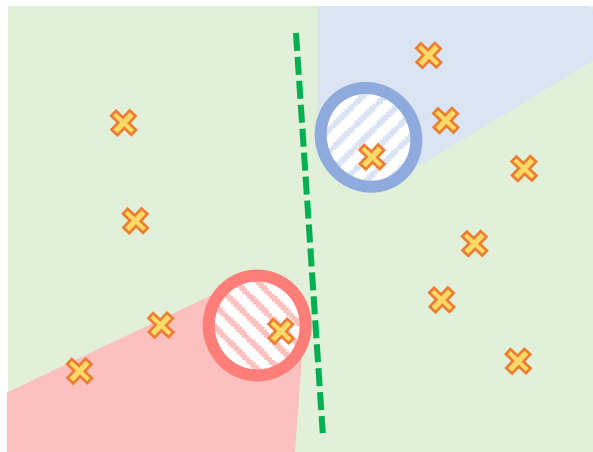


Illustration: Extrapolatable vs. novel samples

Extrapolatable given training distribution + linear ground truth:

Points $x \in X$ where the set of all linear Bayes optimal classifiers agree on

intersecting all
optimal classifiers
yields



--- True classifier

 Training support P

 Unlabeled test data

  Correctly extrapolatable

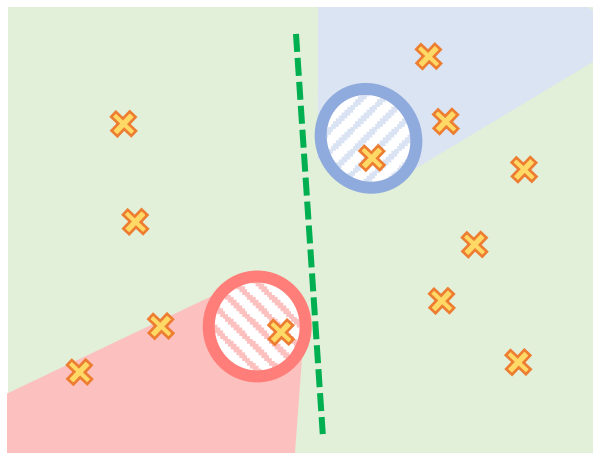
 Not extrapolatable (OOD)

Illustration: Extrapolatable vs. novel samples

Extrapolatable given training distribution + linear ground truth:

Points $x \in X$ where the set of all linear Bayes optimal classifiers agree on

intersecting all
optimal classifiers
yields



--- True classifier

 Training support P

 Unlabeled test data

 Correctly extrapolatable

 Not extrapolatable (OOD)

Goal now: how to output green area

1. Definition: Points we can't make inference on
2. Approach: How to detect those samples?

Semi-supervised novelty detection using ensembles

OOD definition suggests following procedure: given K models

- with good validation accuracy on old classes
- but different predictions outside of training distribution

Semi-supervised novelty detection using ensembles

OOD definition suggests following procedure: given K models

- with good validation accuracy on old classes
- but different predictions outside of training distribution

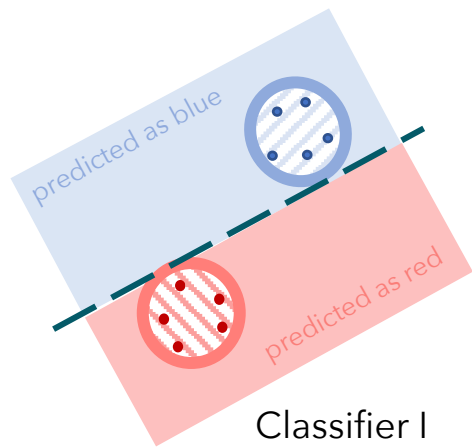
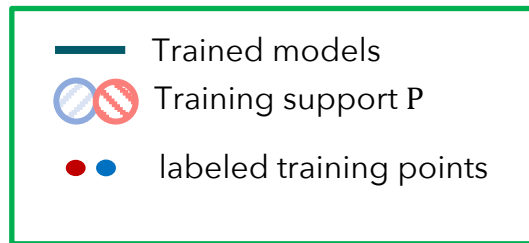
→ flag all points where the models disagree as “novel”

Semi-supervised novelty detection using ensembles

OOD definition suggests following procedure: given K models

- with good validation accuracy on old classes
- but different predictions outside of training distribution

→ flag all points where the models disagree as “novel”



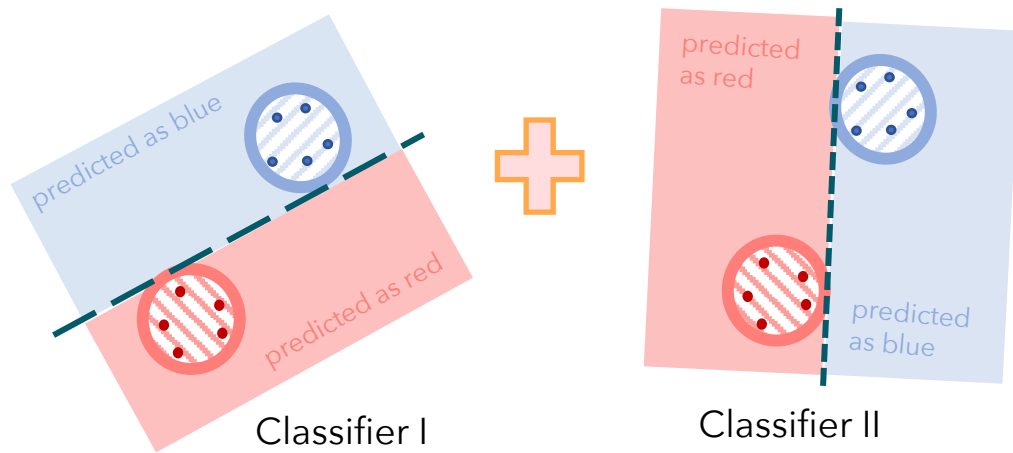
Semi-supervised novelty detection using ensembles

OOD definition suggests following procedure: given K models

- with good validation accuracy on old classes
- but different predictions outside of training distribution

→ flag all points where the models disagree as “novel”

- Trained models
- ⊗ Training support P
- • labeled training points

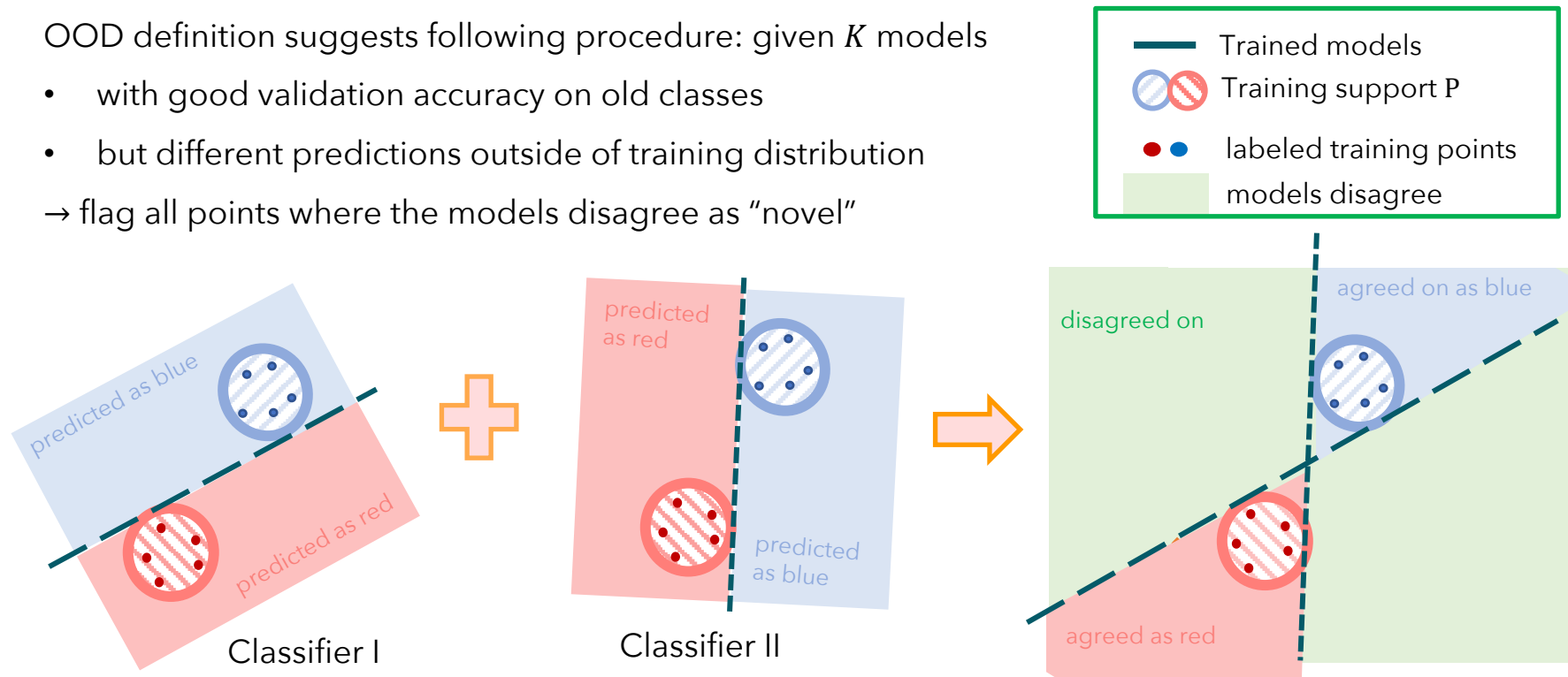


Semi-supervised novelty detection using ensembles

OOD definition suggests following procedure: given K models

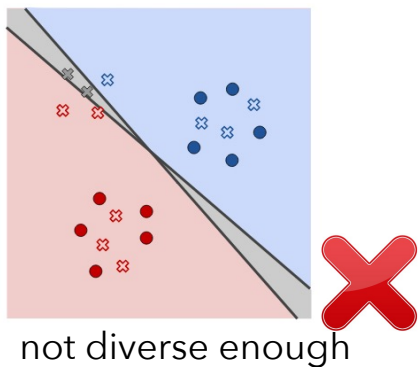
- with good validation accuracy on old classes
- but different predictions outside of training distribution

→ flag all points where the models disagree as “novel”



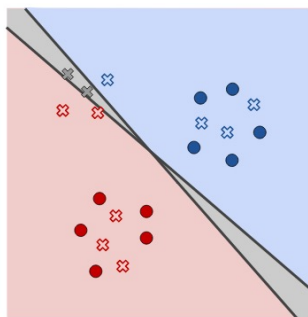
Key for our improvement: Regularized disagreement

Key for “good performance”: Complexity of ensemble models being only as large as needed

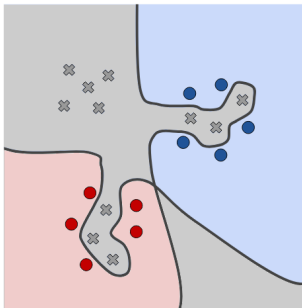


Key for our improvement: Regularized disagreement

Key for “good performance”: Complexity of ensemble models being only as large as needed



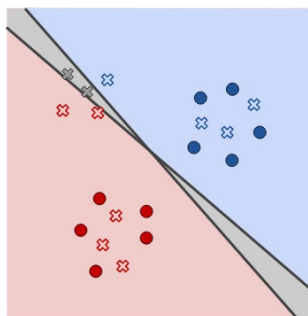
not diverse enough



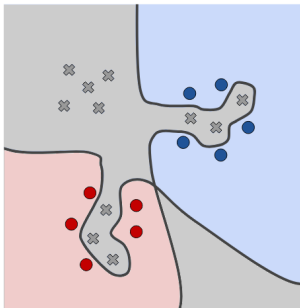
too diverse

Key for our improvement: Regularized disagreement

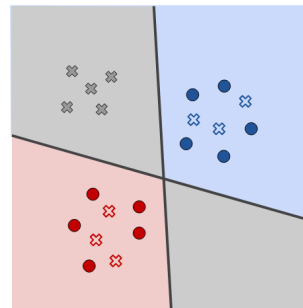
Key for “good performance”: Complexity of ensemble models being only as large as needed



not diverse enough



too diverse

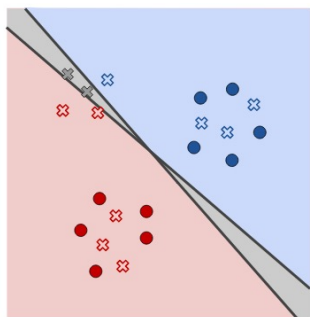


right amount of
diversity

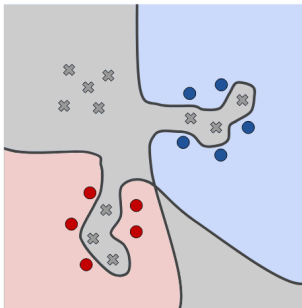


Key for our improvement: Regularized disagreement

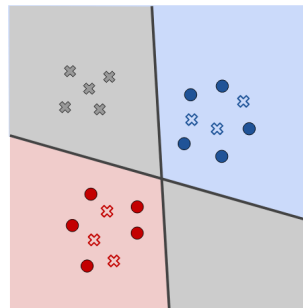
Key for “good performance”: Complexity of ensemble models being only as large as needed



not diverse enough



too diverse

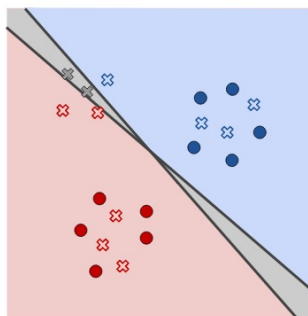


right amount of
diversity

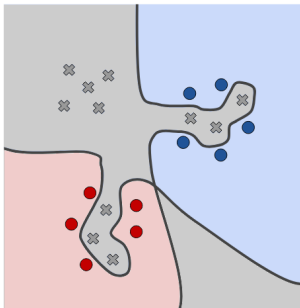
Idea for right amount of disagreement: maximize disagreement s.t. $\underbrace{\text{validation error of all models small}}_{\text{“regularization”}}$

Key for our improvement: Regularized disagreement

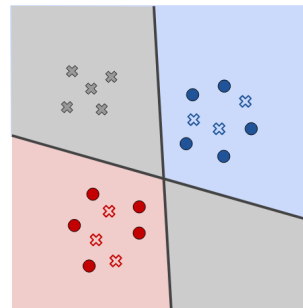
Key for “good performance”: Complexity of ensemble models being only as large as needed



not diverse enough



too diverse



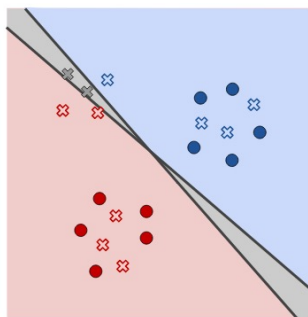
right amount of
diversity

Idea for right amount of disagreement: maximize disagreement s.t. $\underbrace{\text{validation error of all models small}}_{\text{“regularization”}}$

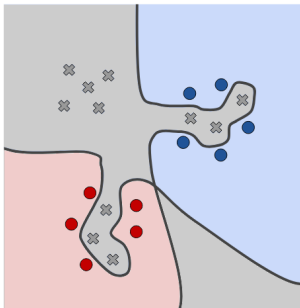
using unlabeled test data

Key for our improvement: Regularized disagreement

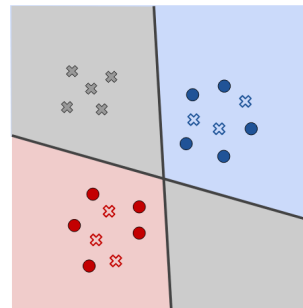
Key for “good performance”: Complexity of ensemble models being only as large as needed



not diverse enough



too diverse



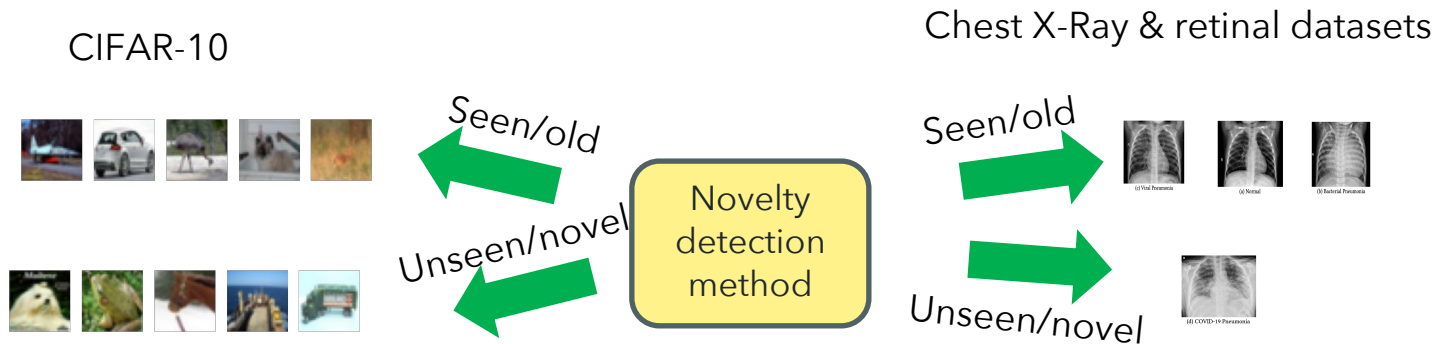
right amount of
diversity

Idea for right amount of disagreement: maximize disagreement s.t. $\underbrace{\text{validation error of all models small}}_{\text{“regularization”}}$

using unlabeled test data

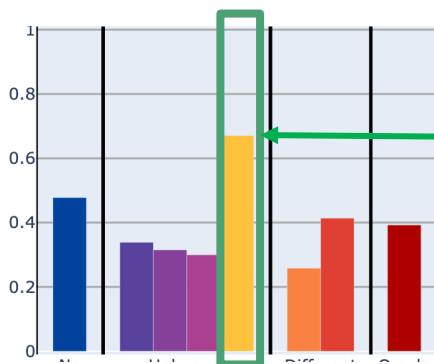
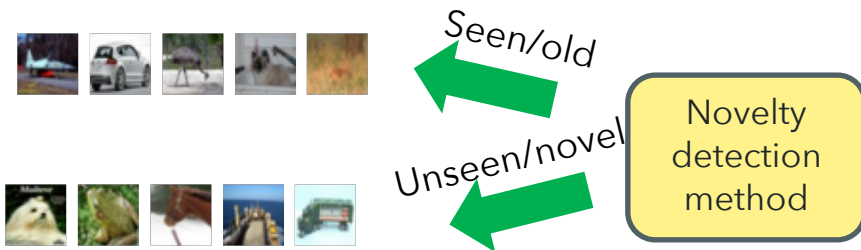
using labeled training data

The near OOD problem on images with DNN



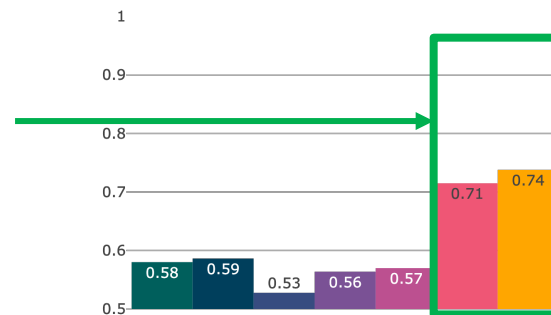
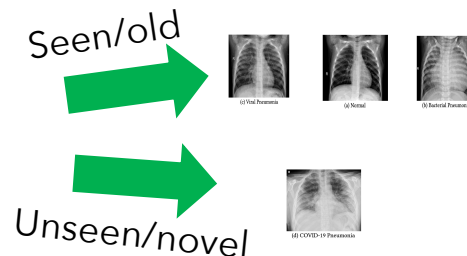
The near OOD problem on images with DNN

CIFAR-10



Ensembles with regularized disagreement

Chest X-Ray & retinal datasets

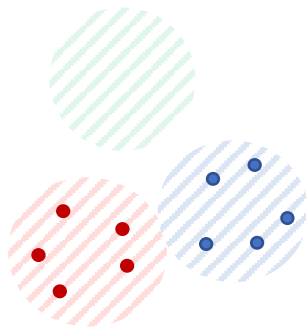
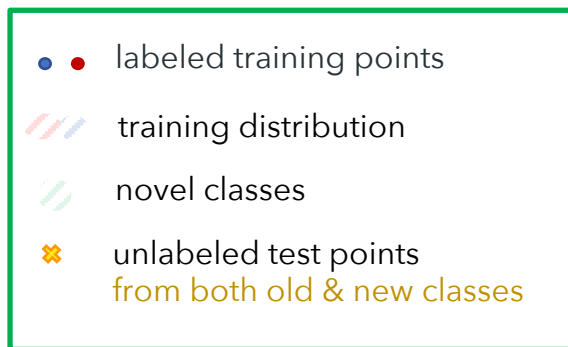


Thanks!

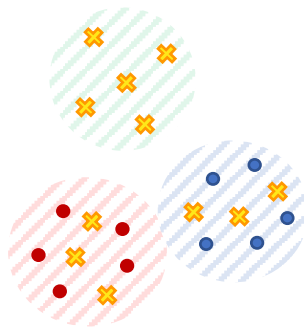
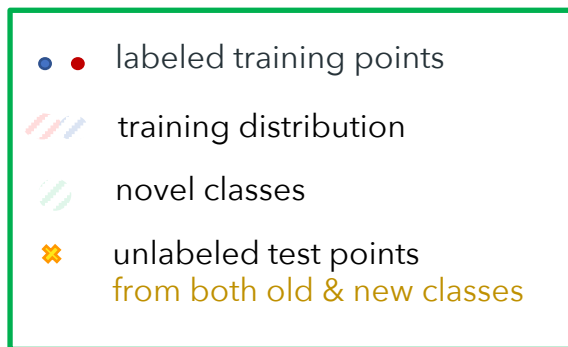


- *"Hidden yet quantifiable: A lower bound for confounding strength using randomized trials"* by Piersilvio De Bartolomeis*, Javier Abad*, Konstantin Donhauser, FY, arxiv preprint
- *"Semi-supervised novelty detection using ensembles with regularized disagreement"* by Alexandru Țifrea, Eric Stavarache, and FY, (UAI), 2022

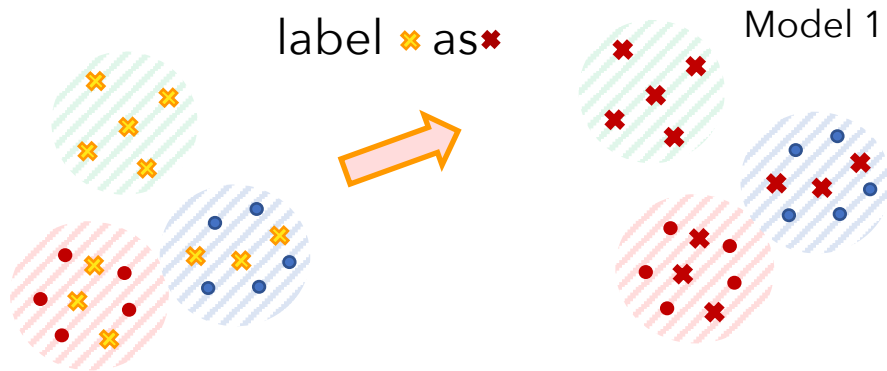
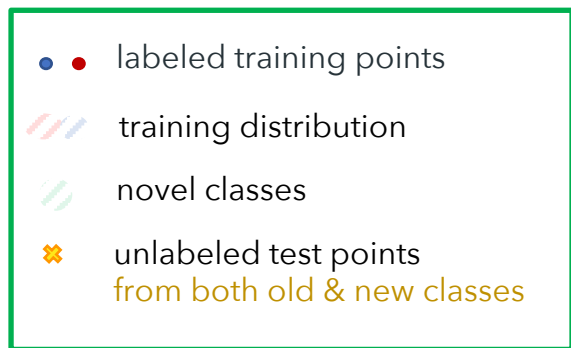
Maximizing disagreement using unlabeled data



Maximizing disagreement using unlabeled data

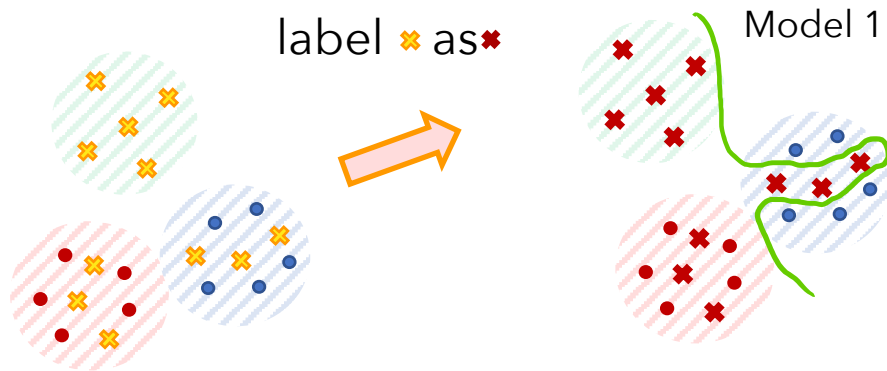
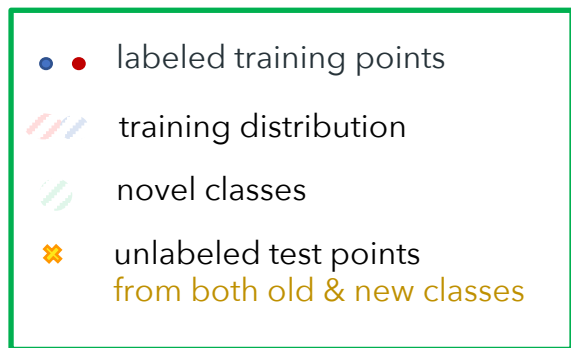


Maximizing disagreement using unlabeled data



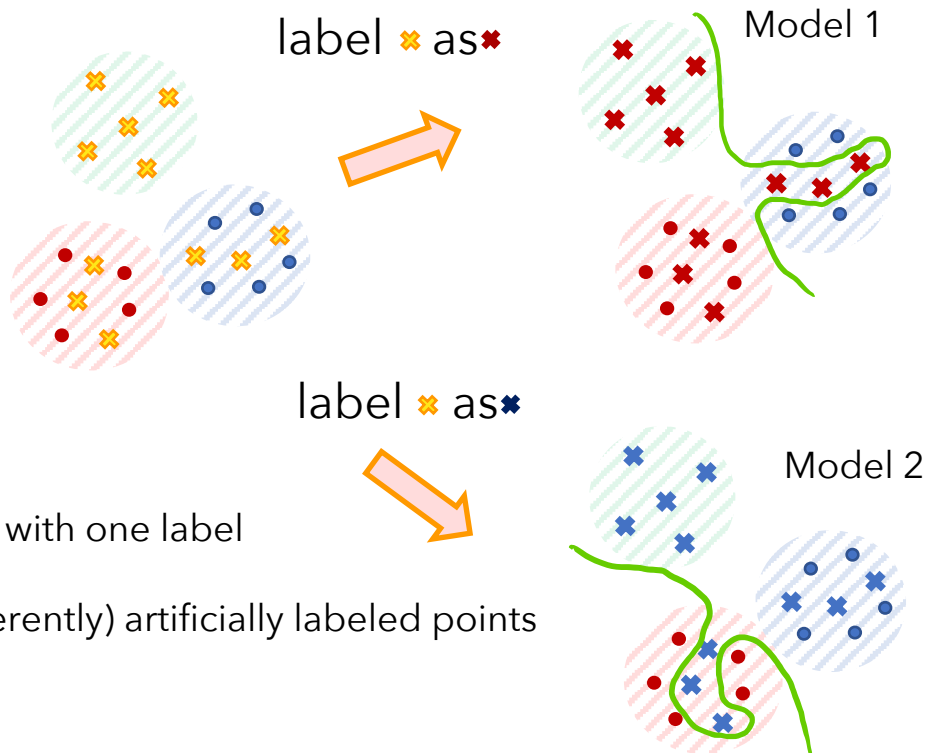
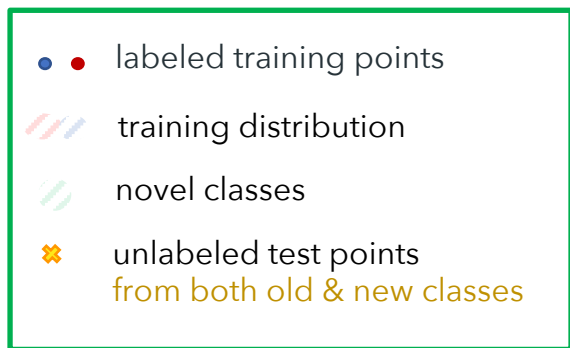
- Artificially label all unlabeled test data with one label

Maximizing disagreement using unlabeled data



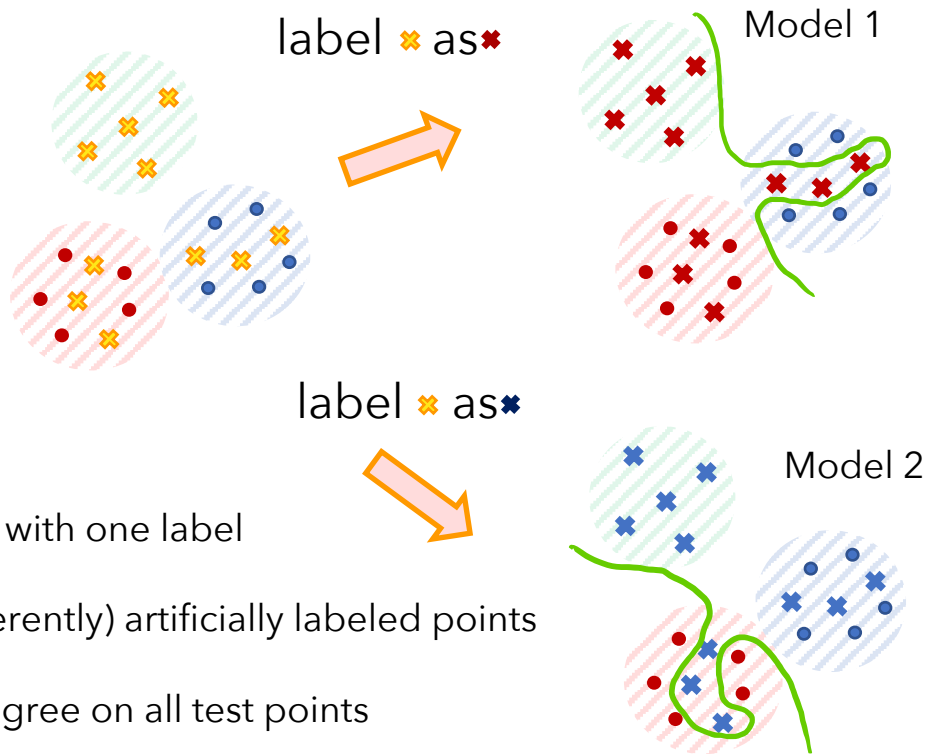
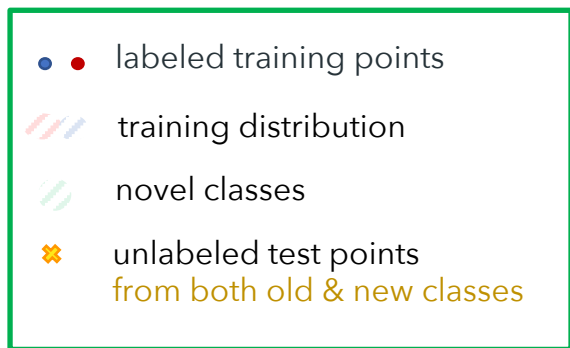
- Artificially label all unlabeled test data with one label

Maximizing disagreement using unlabeled data



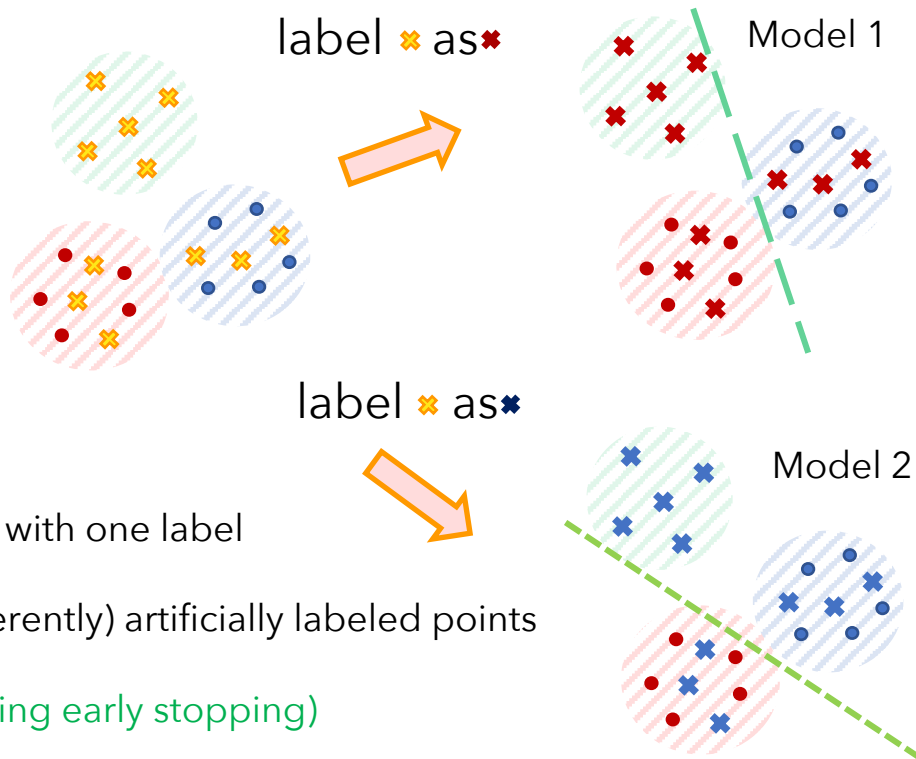
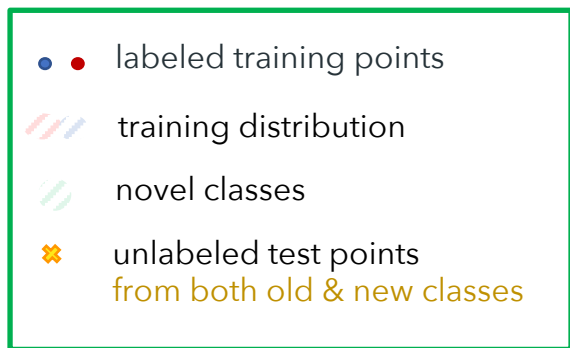
- Artificially label all unlabeled test data with one label
- Fit different models on labeled & (differently) artificially labeled points

Maximizing disagreement using unlabeled data



- Artificially label all unlabeled test data with one label
 - Fit different models on labeled & (differently) artificially labeled points
- ... NNs can fit every point perfectly → disagree on all test points

Regularizing disagreement using labeled data

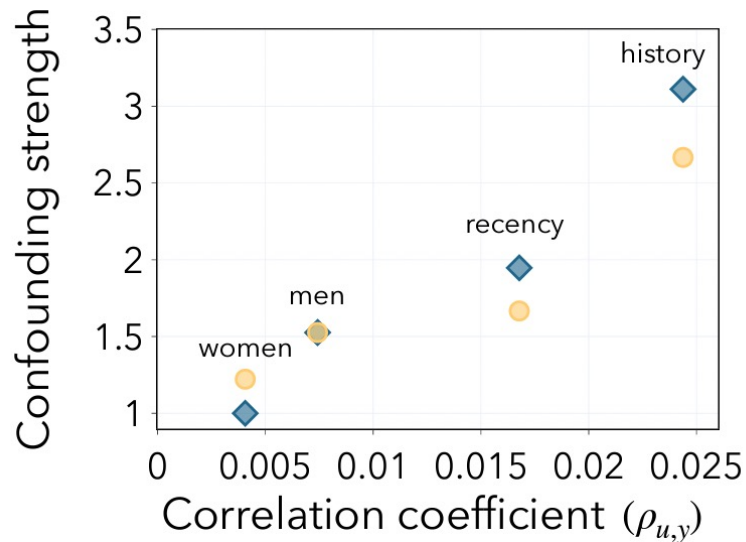


- Artificially label all unlabeled test data with one label
- Fit different models on labeled & (differently) artificially labeled points

... such that validation error is low (e.g. using early stopping)

Current and future work

Non-adversarial confounding



Discussion of the paradigm

- Propose two tests $\phi(\Gamma)$ based on (C)ATE sensitivity analysis intervals
 - obs: estimate μ with importance weighting rct, then ATE sensitivity valid when ATE bounds are asymptotically normal
 - rct: estimate μ on rct, then CATE sensitivity on obs $\rightarrow$ average on rct valid when CATE sensitivity bounds converge at a $1/\sqrt{n}$ rate and $n_{rct} \ll n_{os}$