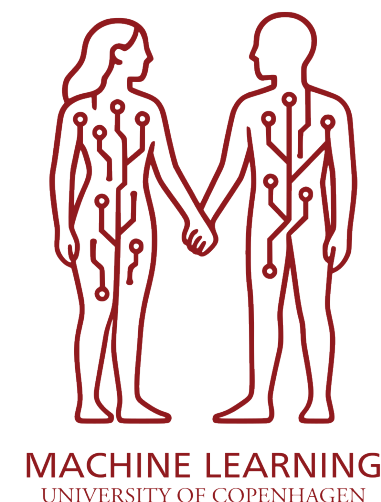
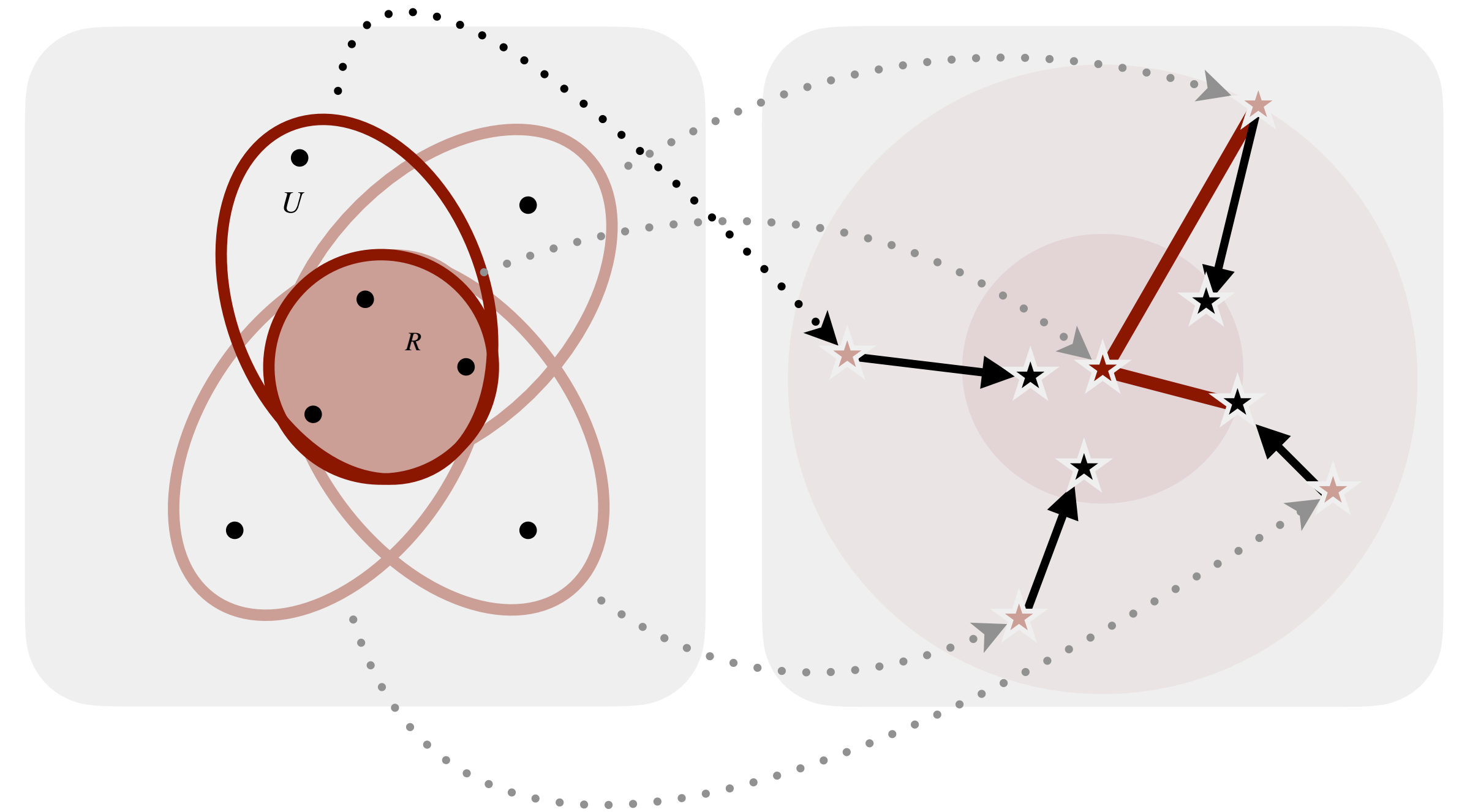


Less Noise, Same Certificate: Retain Sensitivity for Unlearning

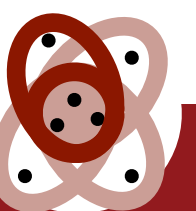
Carolin Heinzler
Kasra Malihi
Amartya Sanyal



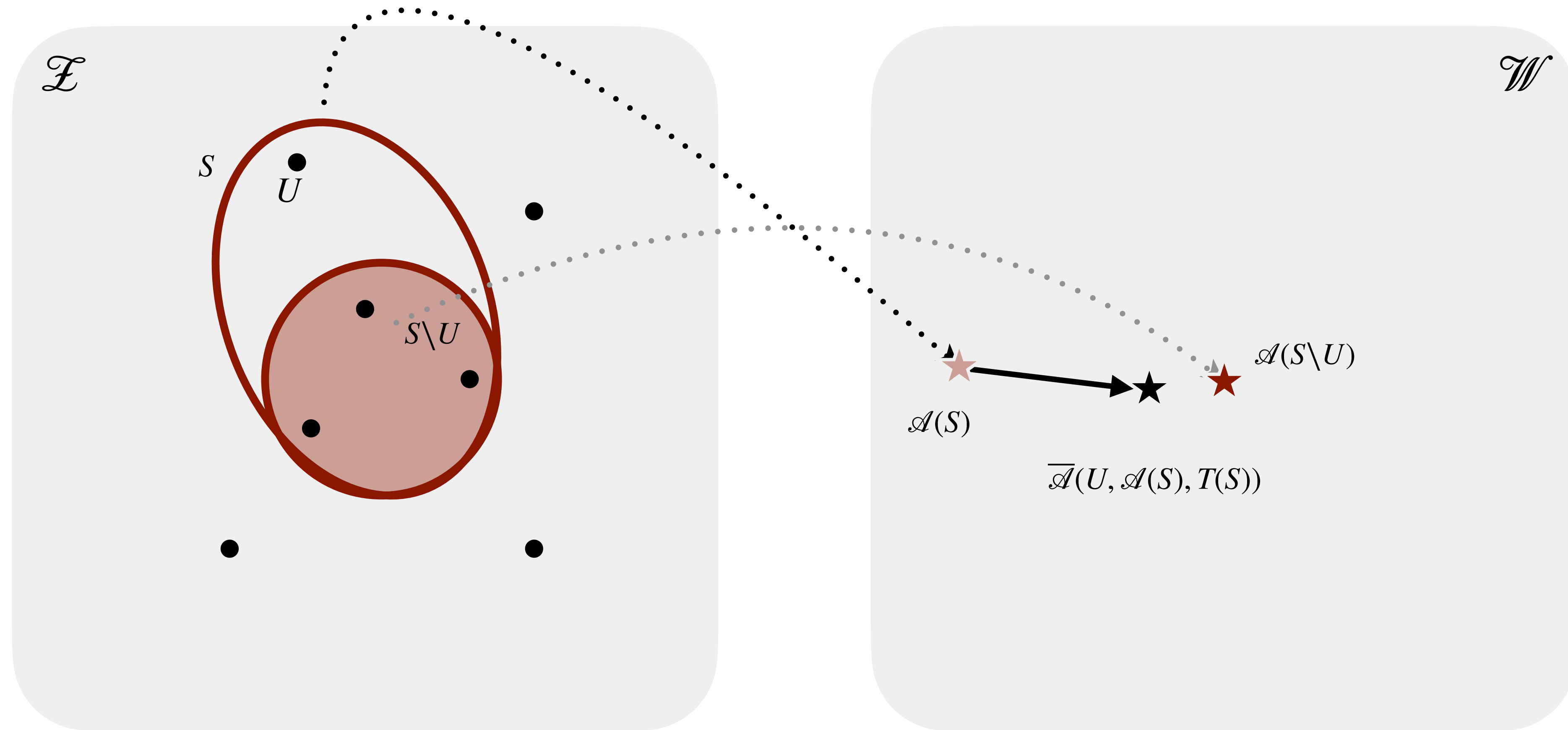
Foundations of Responsible Machine
Learning (FoRML)

University of Copenhagen,
Computer Science Department

FoRML
Foundations of Responsible Machine Learning



Unlearning



Certified Unlearning

Let $\varepsilon, \delta > 0$. A learning-unlearning algorithm pair $(\mathcal{A}, \overline{\mathcal{A}})$ satisfies (ε, δ) -Unlearning, if:

- for all **Datasets** $S \subseteq \mathcal{X}$, $|S| = n + m$
- and for all **Unlearning sets** $U \subseteq S$, $|U| \leq m$

$$\overline{\mathcal{A}}(U, \mathcal{A}(S), T(S))$$

Unlearning

$$\stackrel{\varepsilon, \delta}{\approx}$$

$$\overline{\mathcal{A}}(\emptyset, \mathcal{A}(S \setminus U), T(S \setminus U))$$

Retraining

Both depend on the same $S \setminus U$



Certified Unlearning

Retain
Perspective

Let $\varepsilon, \delta > 0$. A learning-unlearning algorithm pair $(\mathcal{A}, \overline{\mathcal{A}})$ satisfies (ε, δ) -Unlearning, if:

- for all **Retain Sets** $R \subseteq \mathcal{F}$
- and for all **Unlearning sets** $U \subseteq \mathcal{F}, |U| \leq m$,
 $|R \cup U| = n + m$

$$\overline{\mathcal{A}}(U, \mathcal{A}(R \cup U), T(R \cup U)) \stackrel{\varepsilon, \delta}{\approx}$$

Unlearning

$$\overline{\mathcal{A}}(\emptyset, \mathcal{A}(R), T(R))$$

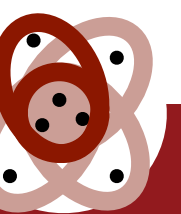
Retraining



Differential Privacy

Let $\epsilon, \delta > 0$. A randomised mechanism $\mathcal{M}$ satisfies (ϵ, δ) -differential privacy if for all **neighbouring datasets** S, S' (add/remove adjacency)

$$\mathcal{M}(S) \stackrel{\epsilon, \delta}{\approx} \mathcal{M}(S')$$



Assume

$$\overline{\mathcal{A}}(U, \mathcal{A}(S), T(S)) = \overline{\mathcal{A}}_0(U, \mathcal{A}(S), T(S)) + \nu$$

$\overline{\mathcal{A}}_0$ deterministic map

$\nu \sim \mathcal{N}(0, \sigma^2 I_d)$ i.i.d
 σ may depend on available
information $(U, T(S))$ or
 $(\emptyset, T(S \setminus U))$

What is minimum noise
necessary to guarantee
unlearning?

If $\overline{\mathcal{A}}_0(U, \mathcal{A}(S), T(S)) = \mathcal{A}(S) \longrightarrow$ Passive Unlearning

Else $\longrightarrow$ Active Unlearning

Full Side Information $T(S) = S$



Noise ~ Sensitivity: DP



DP

Global Sensitivity:

$$\text{GS}_f = \max_{S, S': d_{\pm}(S, S')=1} \|f(S) - f(S')\|$$



No DP

Local Sensitivity:

$$\text{LS}_f(S) = \max_{S': d_{\pm}(S, S')=1} \|f(S) - f(S')\|$$



Noise ~ Sensitivity: Unlearning

For $f = \overline{\mathcal{A}}_0$,
 $R = S \setminus U$,
 $\sigma \sim \text{Sensitivity}$



DP

Global Sensitivity:

$$\text{GS}_f^{(m)} = \max_{S, S': d_{\pm}(S, S')=m} \|f(S) - f(S')\|$$



Unlearning



No DP

Local Sensitivity:

$$\text{LS}_f^{(m)}(S) = \max_{S': d_{\pm}(S, S')=m} \|f(S) - f(S')\|$$



No Unlearning
for $\text{LS}_f^{(m)}(S)$



Unlearning
for $\text{LS}_f^{(m)}(R)$



No DP

Retain Sensitivity:

$$\text{RS}_f^{(m)}(S) = \max_{Z \subseteq \mathcal{L}: |Z| \leq m} \|f(S) - f(S \cup Z)\|$$



Ours



Unlearning
for $\text{RS}_f^{(m)}(R)$



No DP

Down Sensitivity:

$$\text{DS}_f^{(m)}(S) = \max_{Z \subseteq \mathcal{L}: |Z| \leq m} \|f(S) - f(S \setminus Z)\|$$



No Unlearning
for $\text{DS}_f^{(m)}(S)$



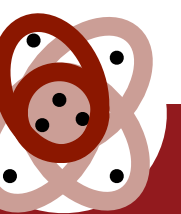
Result

Let $\varepsilon, \delta > 0$. For all $R \subseteq \mathcal{X}$, for all $U \subseteq \mathcal{X}$, $|U| \leq m$, $|R| + |U| = n + m$:

$$\begin{aligned} & \overline{\mathcal{A}}(U, \mathcal{A}(R \cup U), R \cup U) \\ &= \overline{\mathcal{A}}_0(U, \mathcal{A}(R \cup U), R \cup U) + \mathcal{N}(0, \sigma(R)^2 I_d) \end{aligned} \stackrel{\varepsilon, \delta}{\approx} \begin{aligned} & \overline{\mathcal{A}}(\emptyset, \mathcal{A}(R), R) \\ &= \overline{\mathcal{A}}_0(\emptyset, \mathcal{A}(R), R) + \mathcal{N}(0, \sigma(R)^2 I_d) \end{aligned}$$

with $\sigma(R) \sim \text{RS}_{\overline{\mathcal{A}}_0}^{(m)}(R)$.

Why? 1. Unlearning **and** Retraining can reconstruct R



Result

Let $\varepsilon, \delta > 0$. For all $R \subseteq \mathcal{X}$, for all $U \subseteq \mathcal{X}$, $|U| \leq m$, $|R| + |U| = n + m$:

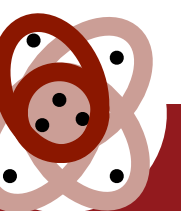
$$\begin{array}{ccc}
 \boxed{R \cup U \setminus U = R} & & \boxed{R \setminus \emptyset = R} \\
 \downarrow & & \downarrow \\
 \overline{\mathcal{A}}(U, \mathcal{A}(R \cup U), R \cup U) & \stackrel{\varepsilon, \delta}{\approx} & \overline{\mathcal{A}}(\emptyset, \mathcal{A}(R), R) \\
 = \overline{\mathcal{A}}_0(U, \mathcal{A}(R \cup U), R \cup U) + \mathcal{N}(0, \sigma(R)^2 I_d) & & = \overline{\mathcal{A}}_0(\emptyset, \mathcal{A}(R), R) + \mathcal{N}(0, \sigma(R)^2 I_d)
 \end{array}$$

with $\sigma(R) \sim \text{RS}_{\overline{\mathcal{A}}_0}^{(m)}(R)$.

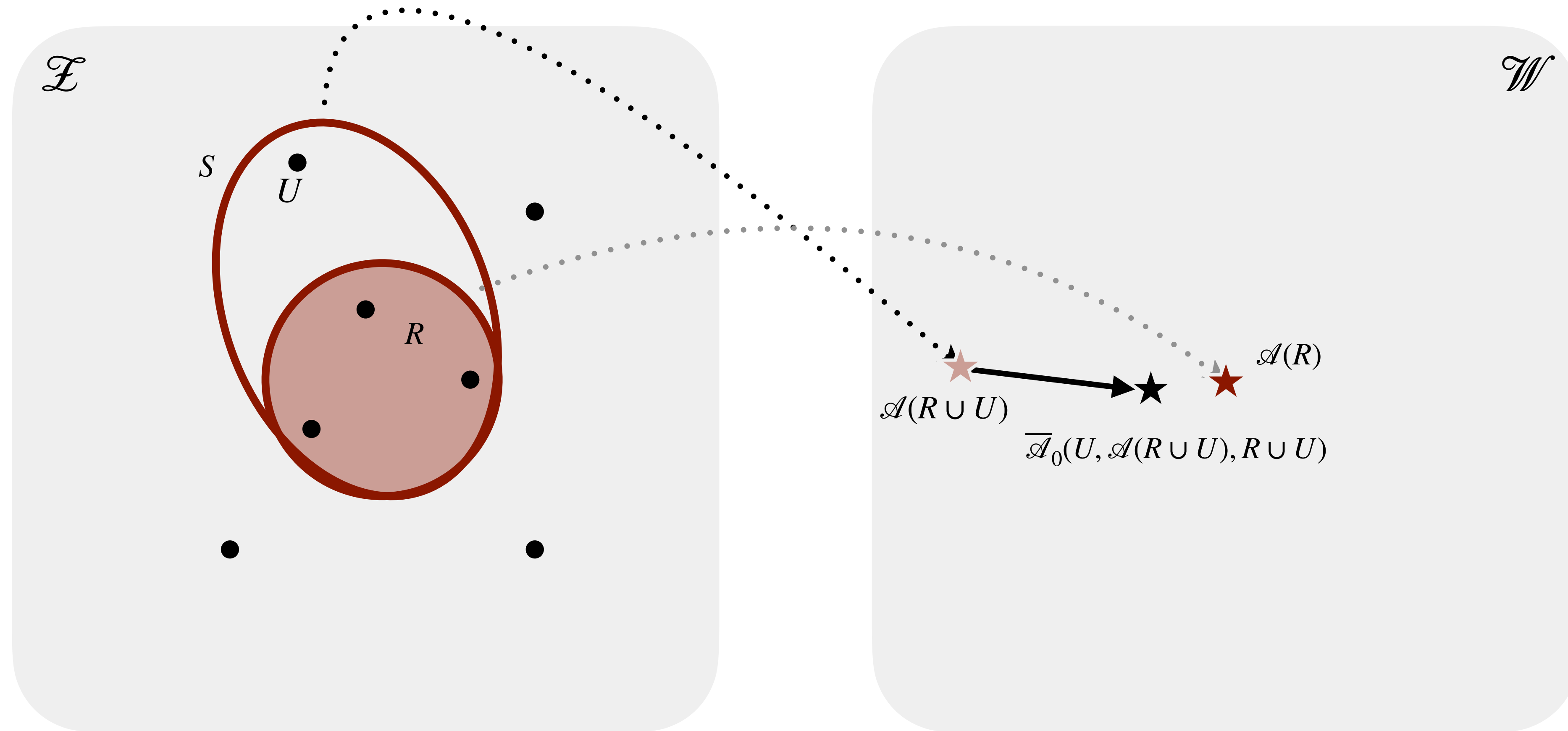
Why? 1. Unlearning **and** Retraining can reconstruct R

2. **Sufficient* noise** to hide unlearning of any Z from $R \cup Z$ (in particular U from $S = R \cup U$)

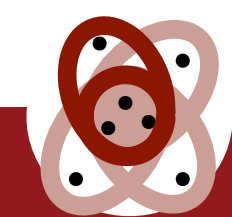
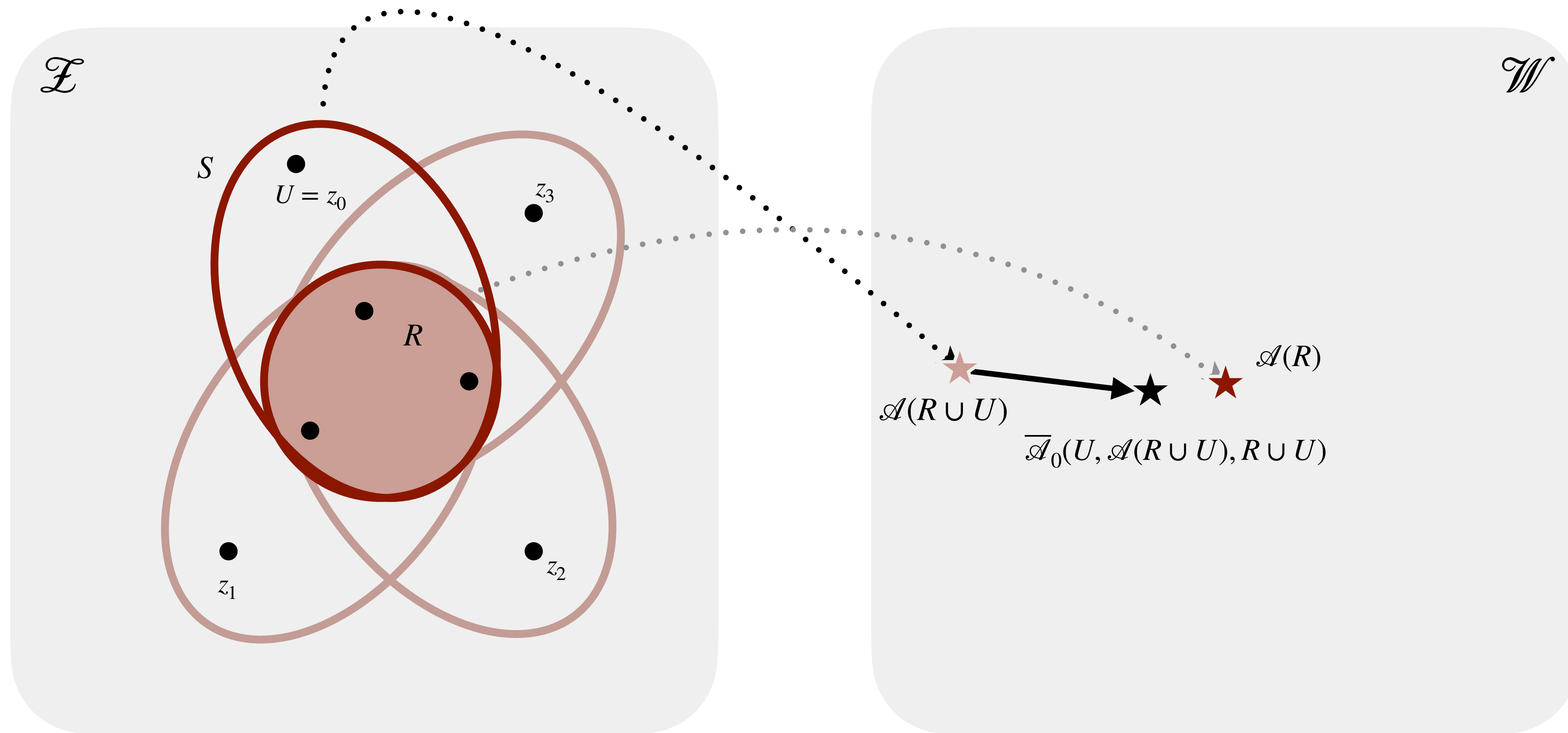
*and **necessary** for additive mechanisms with Gaussian noise by mean shift Lemma



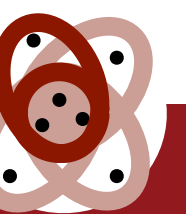
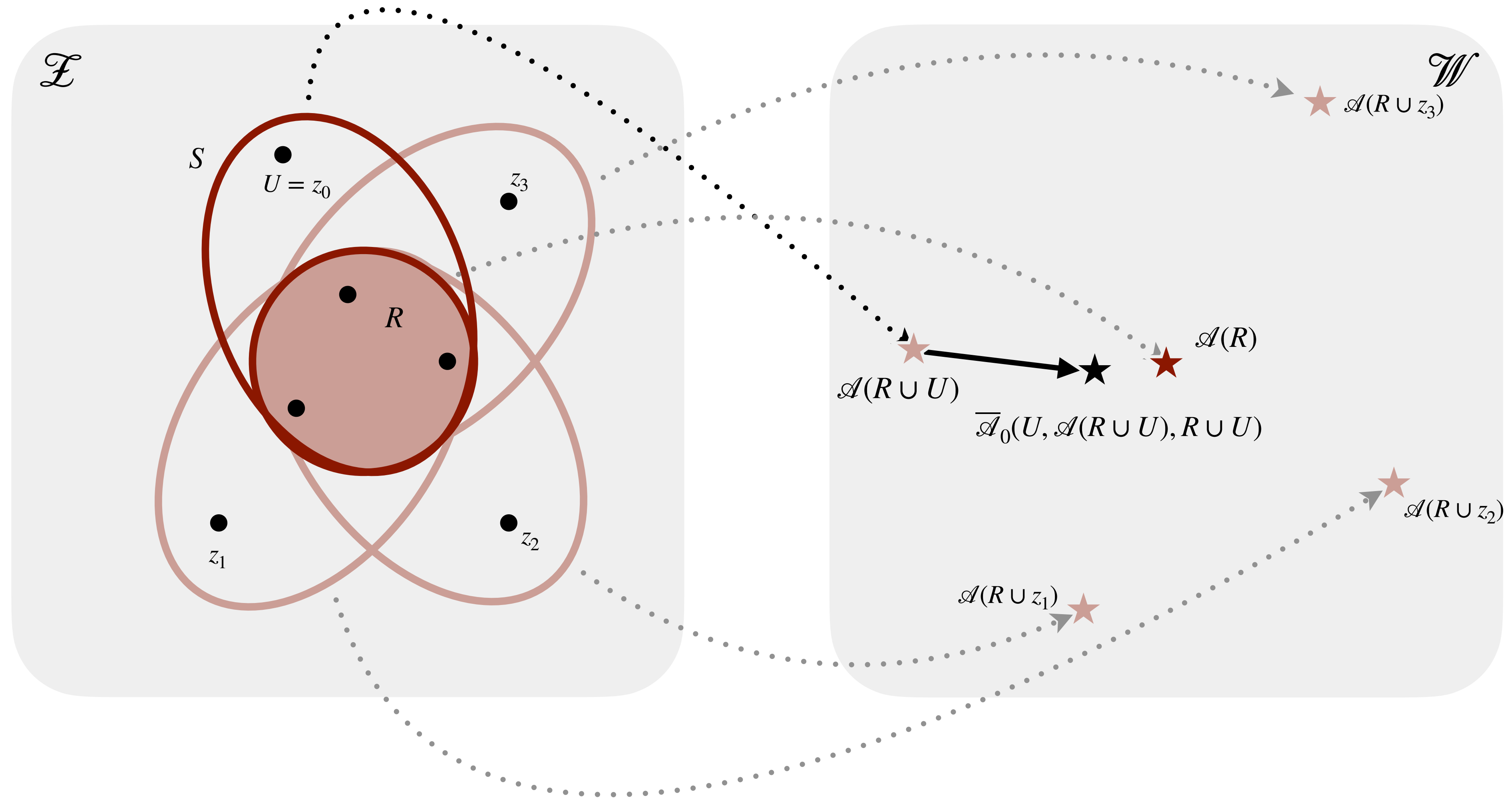
Unlearning



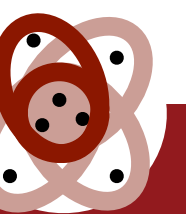
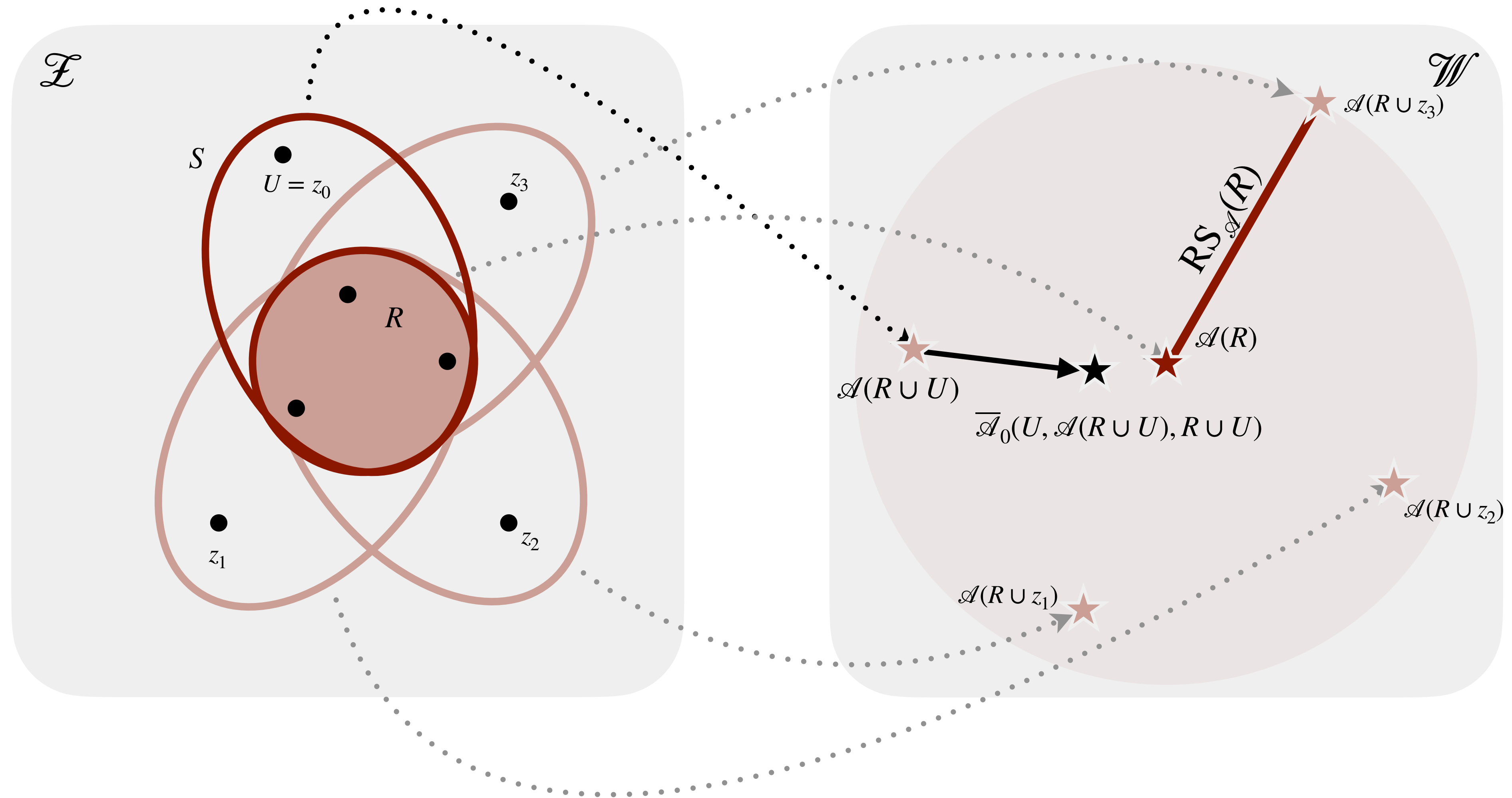
Unlearning



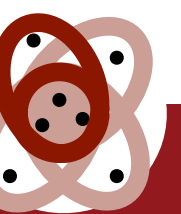
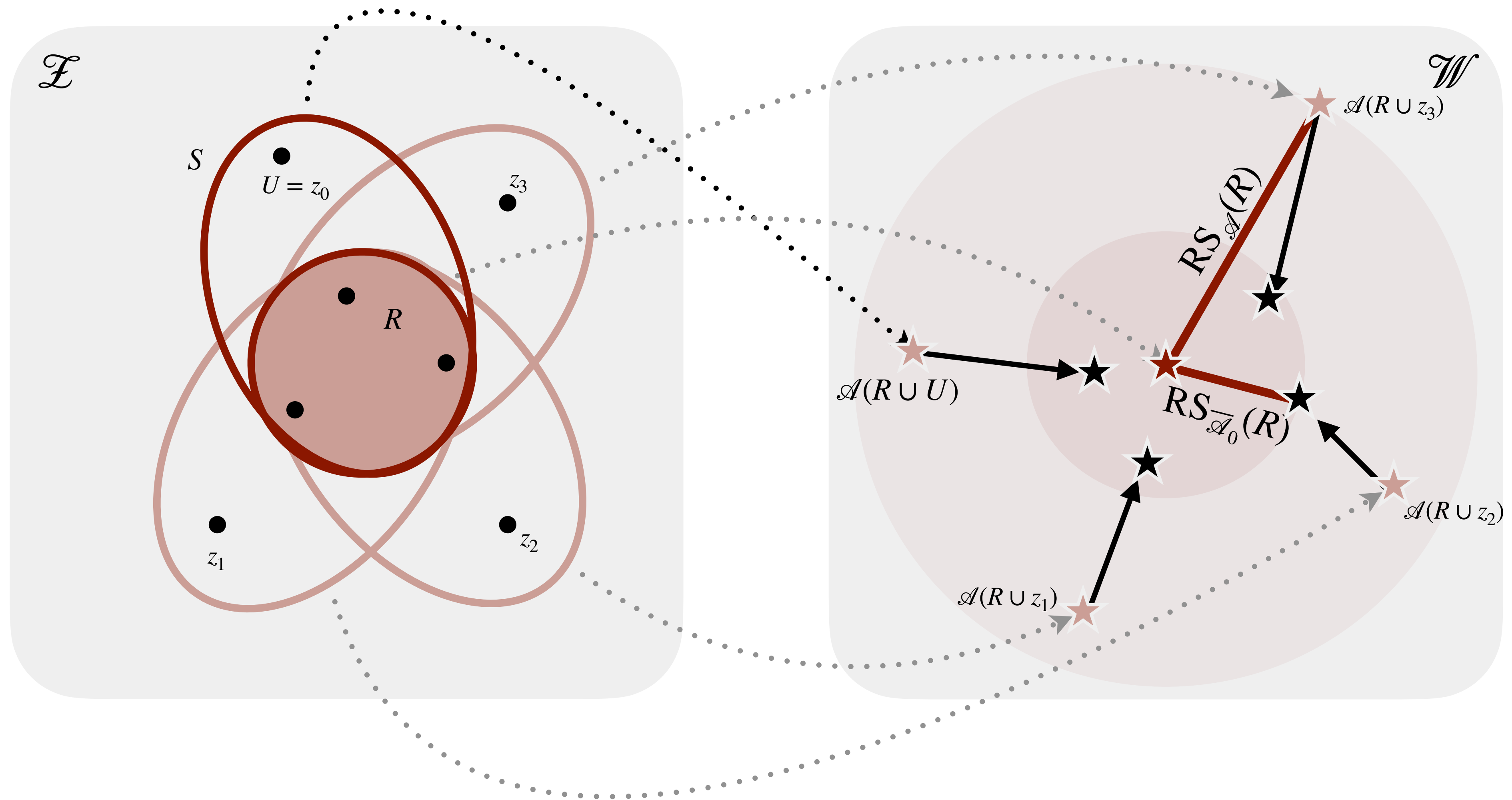
Unlearning



Unlearning



Unlearning

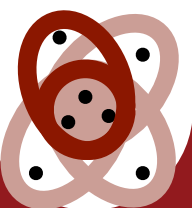


$$|R| = n, \\ \|x\| \leq B$$

Noise ~ Properties of R

Better retain-set ‘conditioning’ → less noise for same certificate → higher utility after unlearning

Problem	Global Sensitivity	Retain Sensitivity
Median	$GS_{\text{median}}^{(m)} = B$	<i>local spacings</i> around median(R)
weight of MST	$GS_{\text{MST}}^{(m)} = mB$	maximum total weight of m <i>replaceable MST edges</i> in R
PCA (rank- k projector)	$GS_{\text{PCA}}^{(m)} \leq \sqrt{2k}$	<i>eigengap</i> of covariance of R
SVM	$GS_{\text{SVM}}^{(m)} = \frac{1}{\gamma}$	<i>empirical margin</i> γ_R
ERM	$GS_{\text{ERM}}^{(m)} \leq \frac{mL}{n\lambda}$	<i>data-dependent strong convexity</i> λ_R of $\hat{F}_R$: $RS_{\text{ERM}}^{(m)}(R) \leq \frac{mL}{n\lambda_R}$



Noise ~ Properties of R : Active Methods

$$|R| = n, \\ \|x\| \leq B$$

$$\text{ERM} \quad \text{GS}_{\text{ERM}}^{(m)} \leq \frac{mL}{n\lambda} \quad \text{data-dependent strong convexity } \lambda_R \text{ of } \hat{F}_R \quad \text{RS}_{\text{ERM}}^{(m)}(R) \leq \frac{mL}{n\lambda_R}$$

Descent-to-Delete [Neel'20]: (gradient descent on retain data)

$$\overline{\mathcal{A}}(U, w_{R \cup U}, R \cup U) = \text{Proj}_{\mathcal{W}}(w_{R \cup U} - \eta_R \nabla \hat{F}_R(w_{R \cup U}))^I + \nu$$

Newton Step Update [Sekhari'21]: (Hessian update)

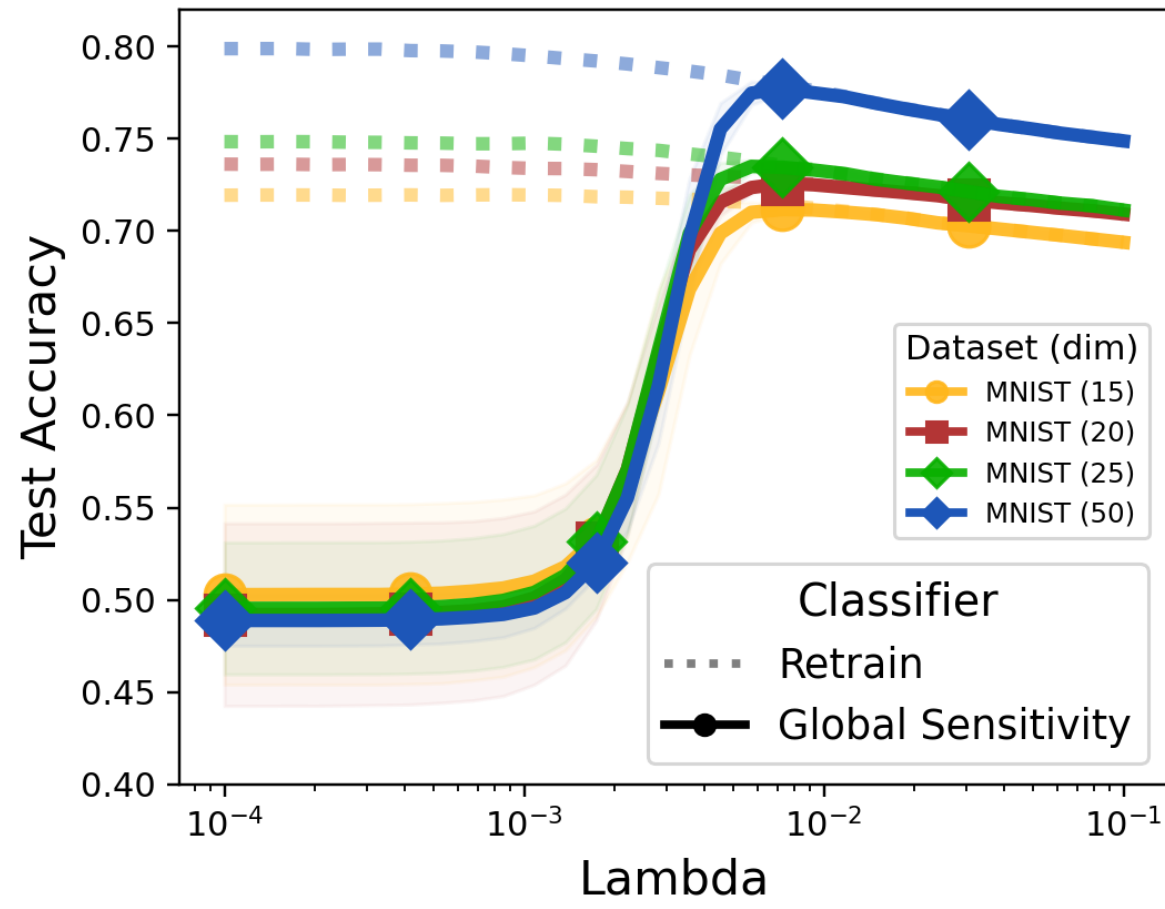
$$\overline{\mathcal{A}}(U, w_{R \cup U}, R \cup U) = w_{R \cup U} + \frac{1}{n} (\nabla^2 \hat{F}_R(w_{R \cup U}))^{-1} \sum_{z \in U} \nabla f(w_{R \cup U}, z)$$



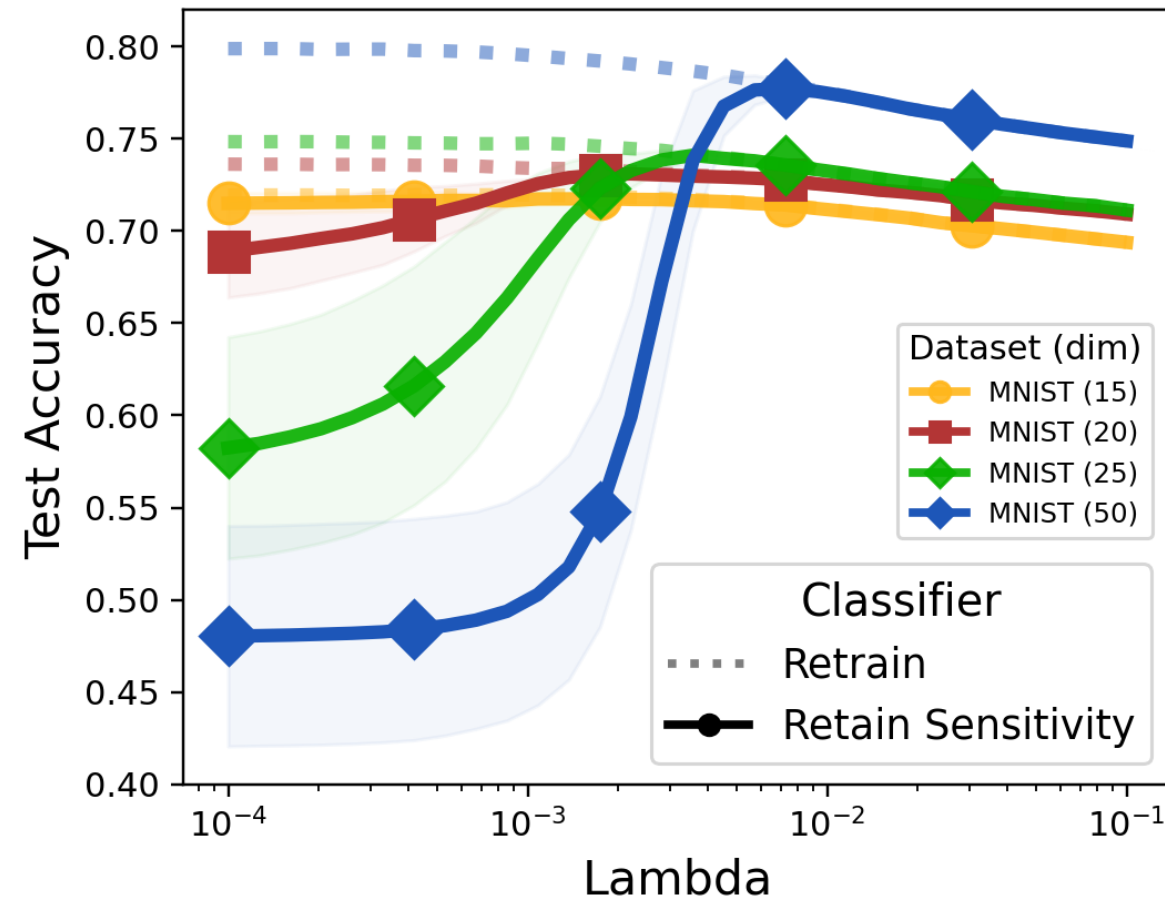
Noise ~ Properties of R : Active Methods

Newton method

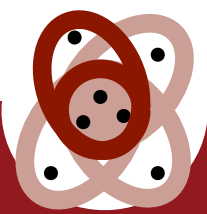
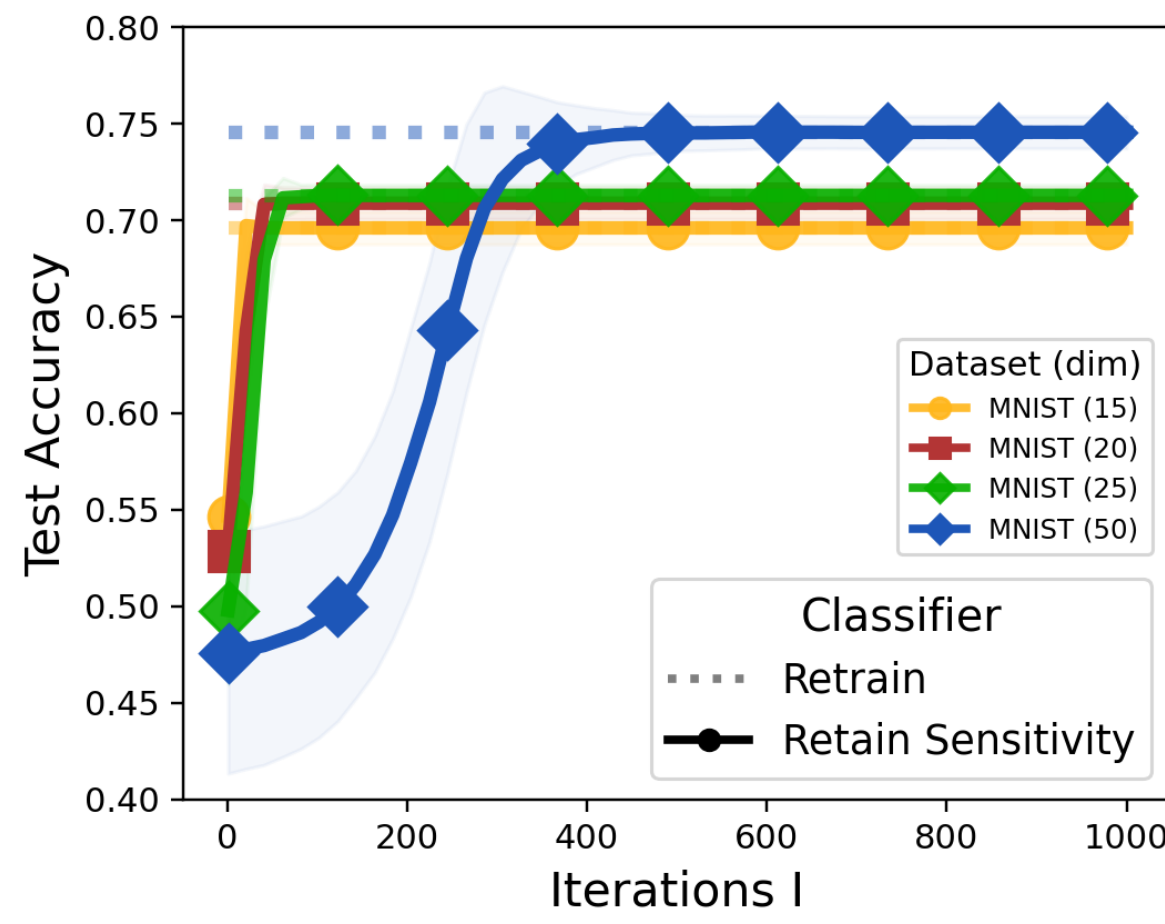
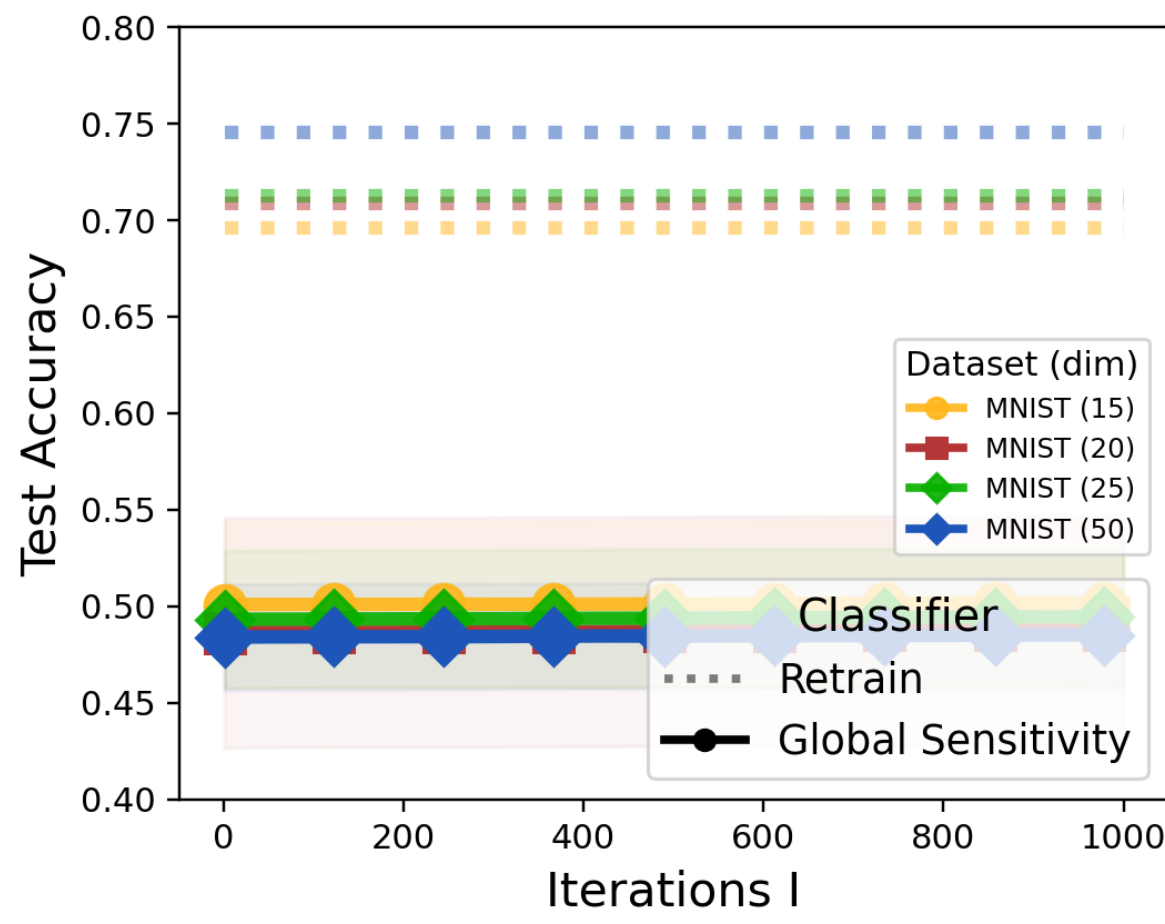
Noise ~ GS



Noise ~ RS



D2D



Discussion

Unlearning mechanisms can rely on noise scaling with Retain Sensitivity at R

Consequence: ϵ -DP-algorithm is ϵ' -Unlearning with **lower** ϵ' than by just group privacy

No **privacy requirement for** R in plain unlearning - desirable? See [Cohen'26]

Limitations: **Efficient computation** of necessary statistics at Unlearning/Training time

Assumption of **Full Side Information** to construct RS: sufficient — but necessary?

