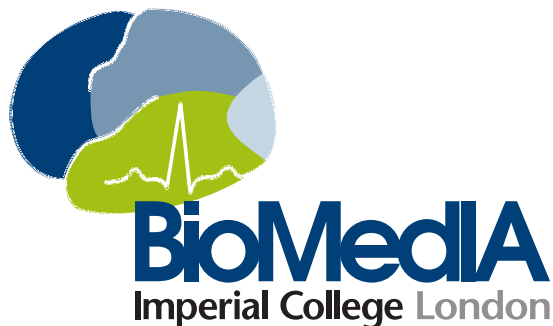


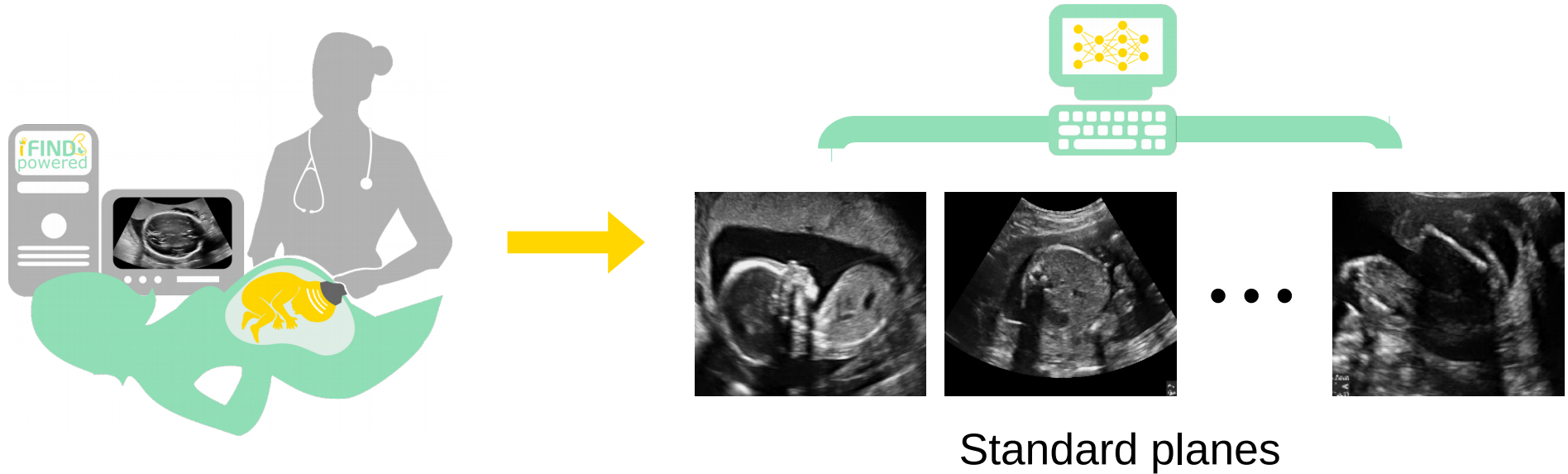
Representation Disentanglement for Multi-task Learning with application to Fetal Ultrasound

Qingjie Meng¹, Nick Pawlowski¹, Daniel Rueckert¹, Bernhard Kainz¹

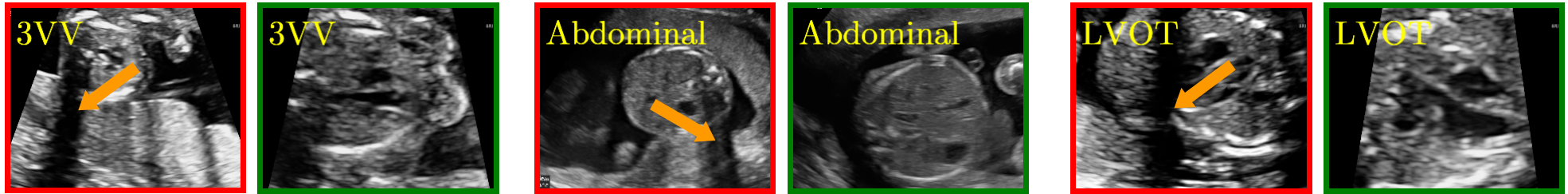
¹ Biomedical Image Analysis Group, Department of Computing,
Imperial College London

Qingjie Meng

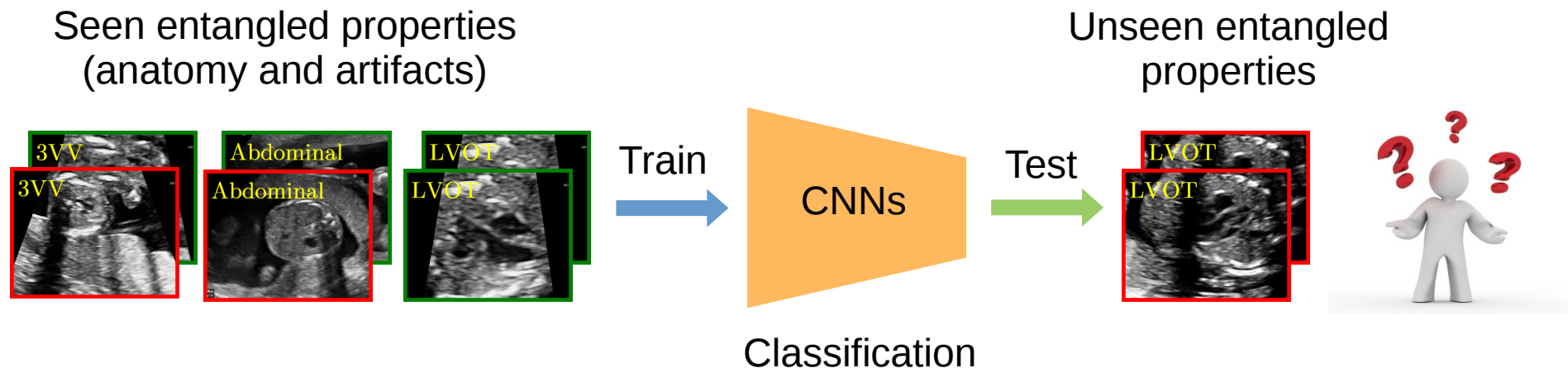




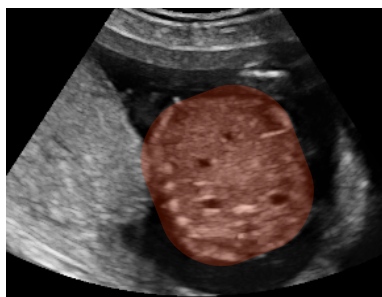
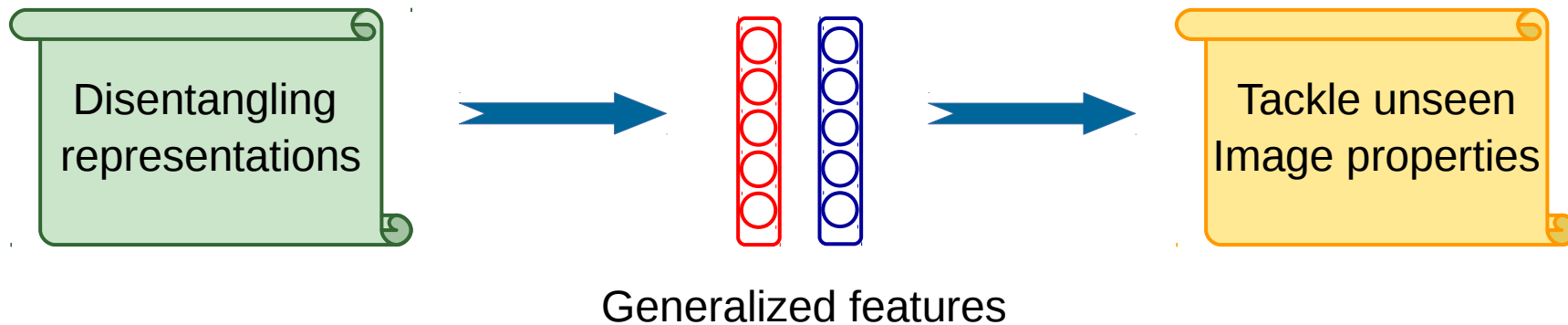
- ◆ Important for abnormality detection in early pregnancy



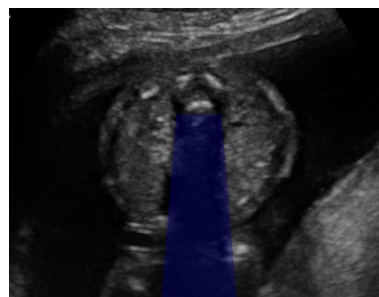
- ◆ CNNs learn *anatomical features* and *shadow features* simultaneously.



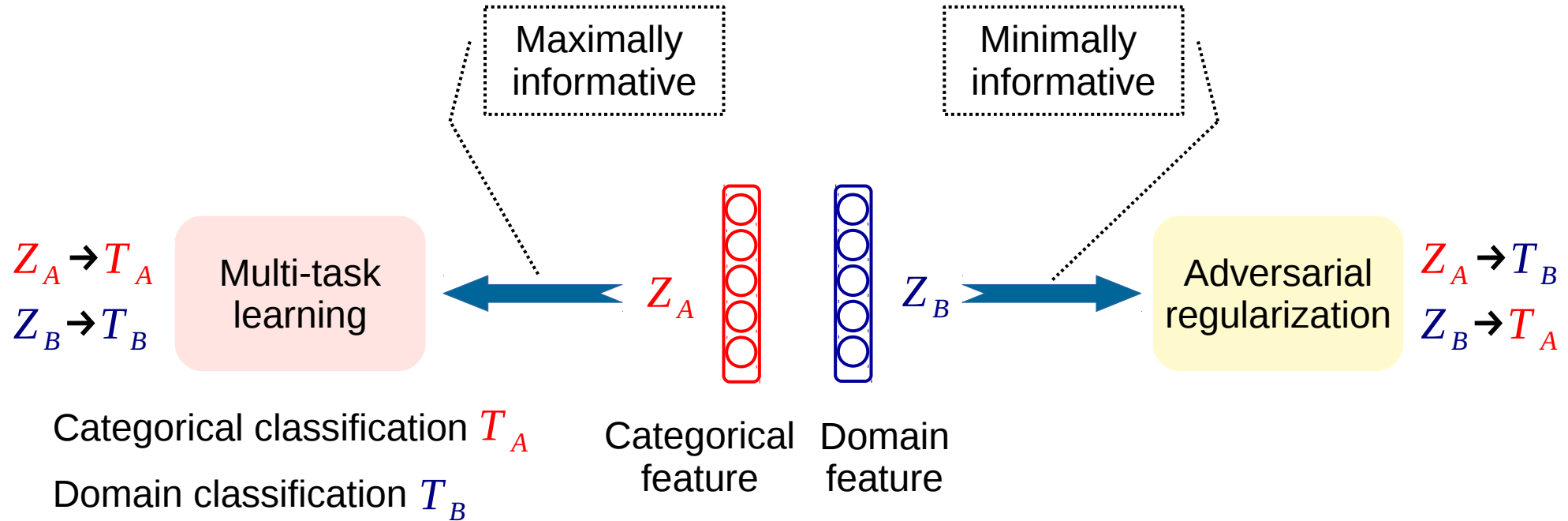
- ◆ CNNs exhibits weak generalizability
- ◆ Our goal is to disentangle *anatomical features* from *shadow features*



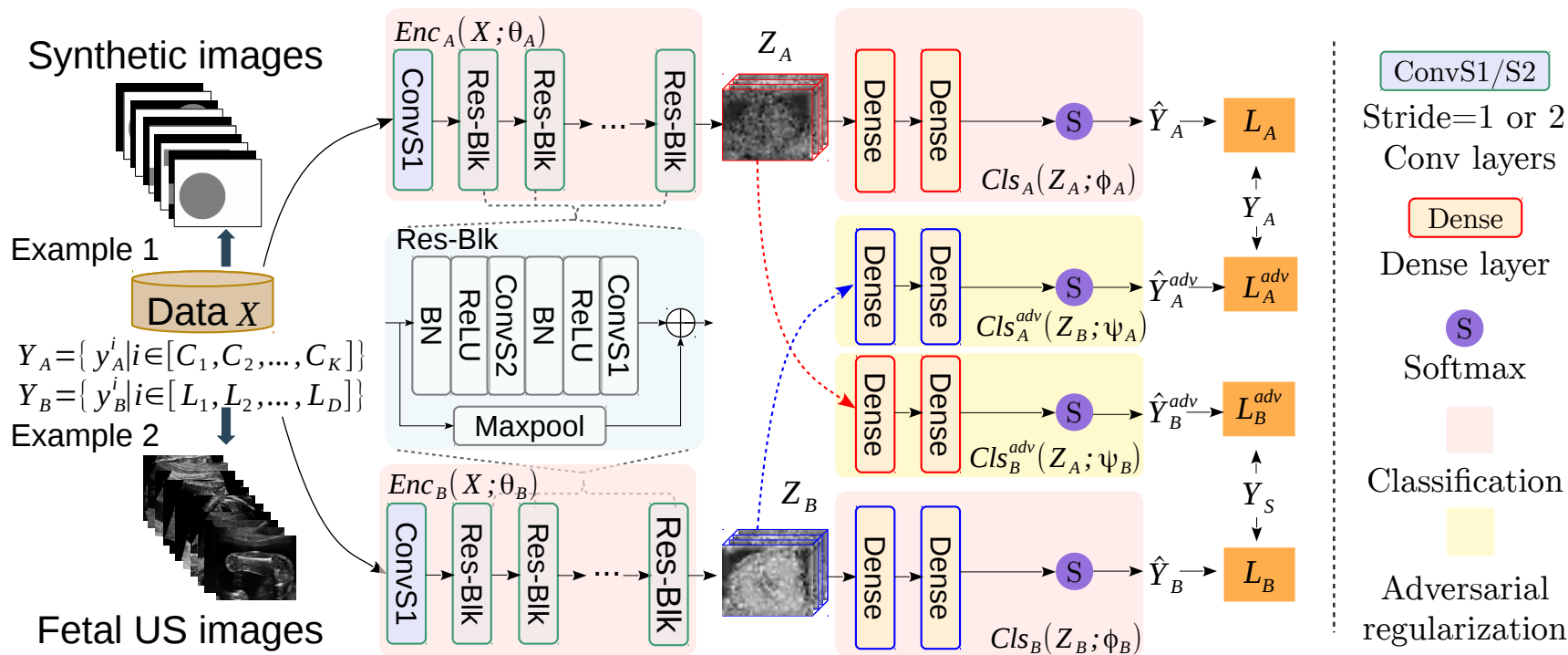
Anatomical feature

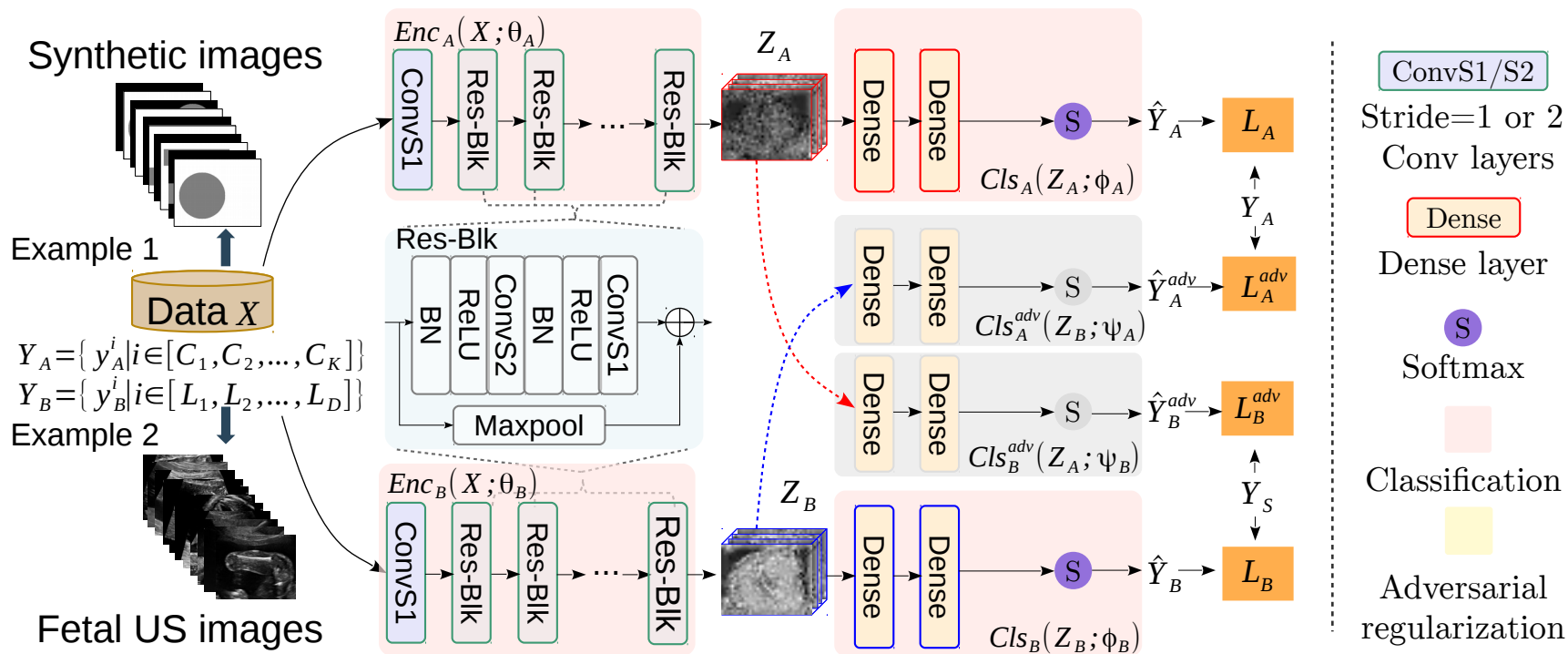


Shadow feature



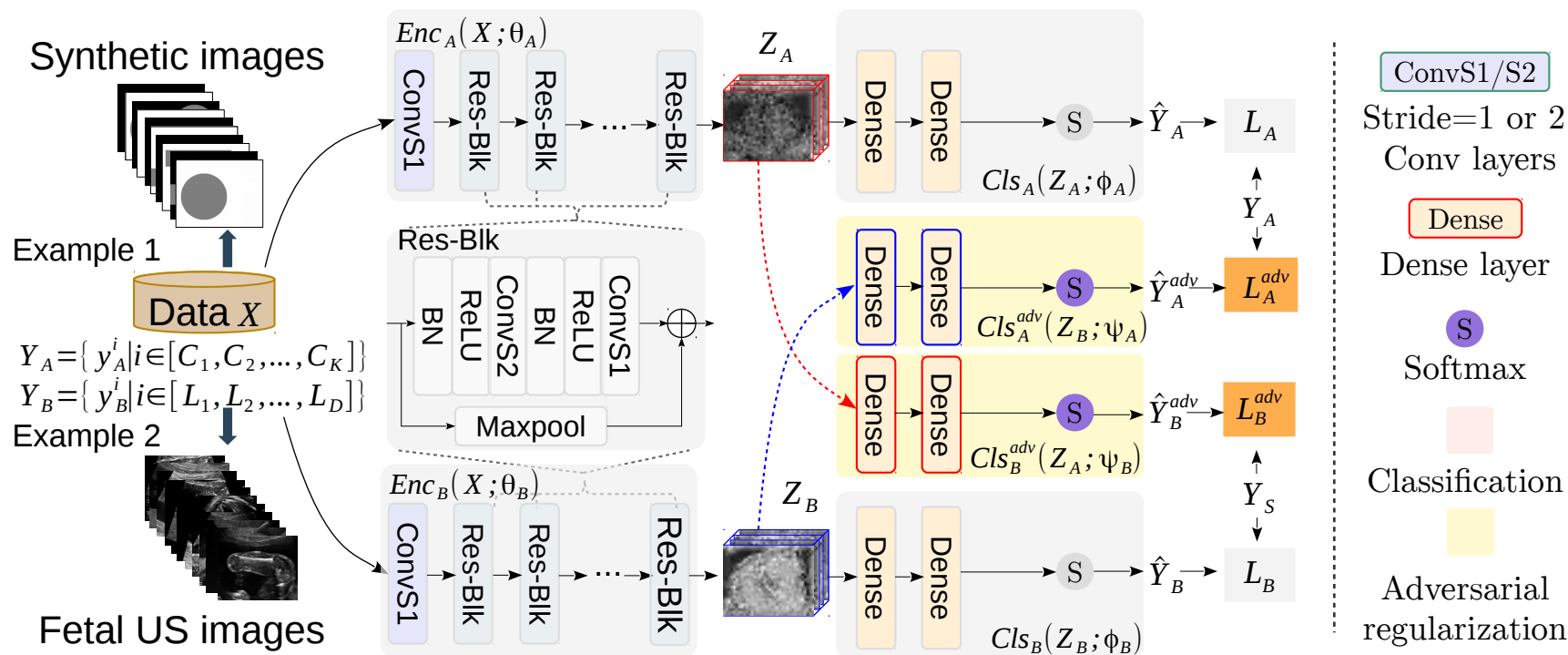
Method





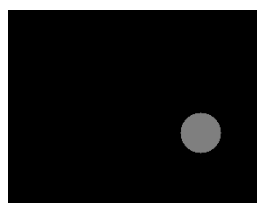
Training objective
$$\min_{\{\theta_A, \theta_B, \phi_A, \phi_B\}} \{ L_A + L_B - \lambda * (L_A^{adv} + L_B^{adv}) \}, \lambda > 0$$

Method





Train/Test_{seen}



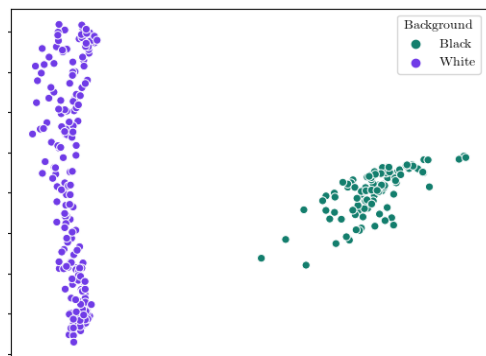
Unseen test

Background classification T_A

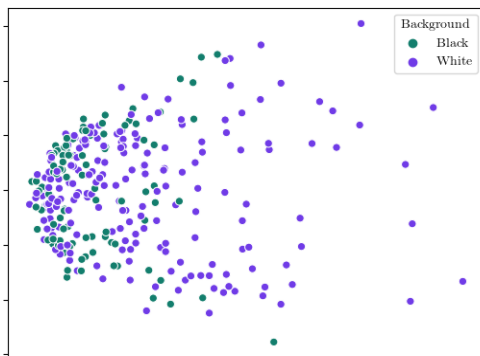
Background feature Z_A

Shape classification T_B

Shape feature Z_B



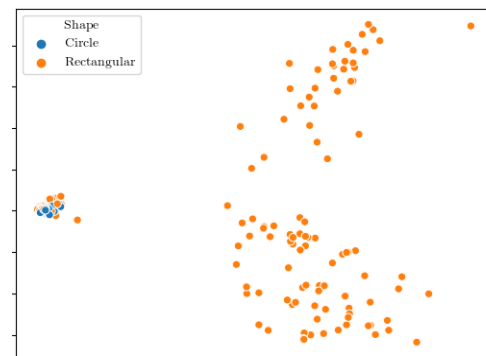
(a) $Z_A \rightarrow T_A$ 100%



(b) $Z_B \rightarrow T_A$ 59.67%



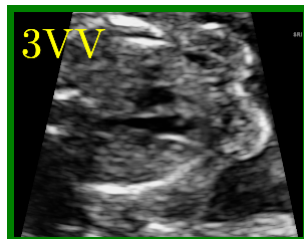
(c) $Z_B \rightarrow T_B$ 99.67%



(d) $Z_A \rightarrow T_B$ 62%

Our model achieves 99% accuracy on unseen test data (baseline is 10%)

Train/
Test_seen



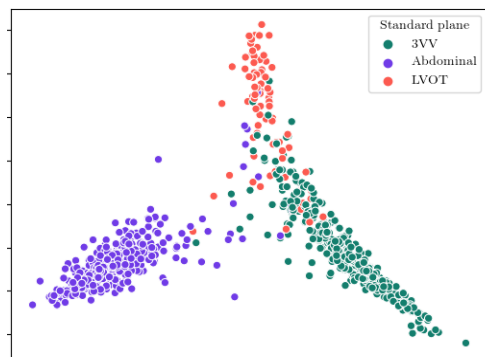
Unseen test for shadow artifacts classification

Unseen test for
anatomical classification

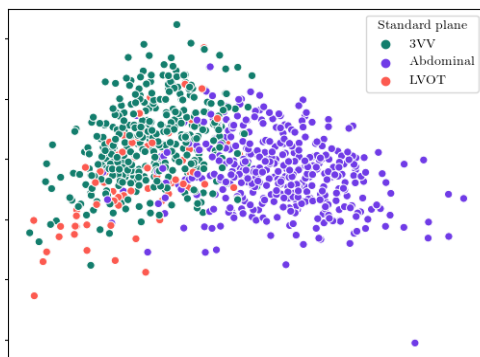
- ◆ Anatomical classification task T_A / feature Z_A
- ◆ Shadow artifacts classification task T_B / feature Z_B

► Classification accuracy of different methods for various tasks.

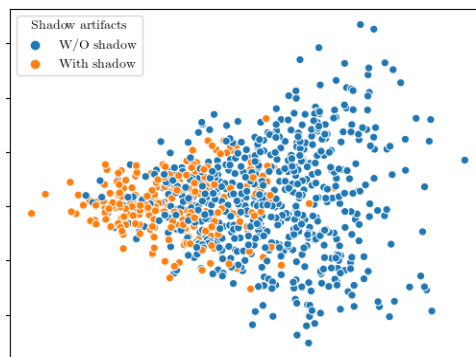
Methods			Std plane only	Artifacts only	Proposed _{W/O_adv}	Proposed	Proposed _{irr_task}
Test_seen	T_A		78.36	--	81.25	94.44	64.35
	T_B		--	78.7	78.74	79.05	72.57
Unseen test	LVOT(W_S)	T_A	34.93	--	37.56	73.68	--
	Artifacts(OTHS)	T_B	--	69.26	69.50	69.50	--



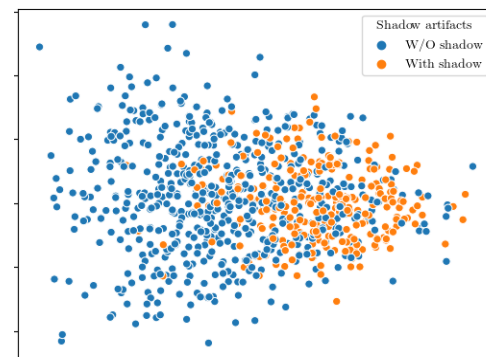
(a) $Z_A \rightarrow T_A$



(b) $Z_B \rightarrow T_A$



(c) $Z_B \rightarrow T_B$



(d) $Z_A \rightarrow T_B$

We require Less
data collection
compare with data
augmentation

Disentanglement
occurs in last layers
because of
adversarial training

Disentanglement is
able to provide
benefits for model
generalization

Thank you!

