

Measuring interpretability in chest radiograph classification

ProtoCXR: every prediction grounded in a real training case

Tran Quoc Hung · Le Tuan Hung · Nguyen Ngoc Dung · Luong Thi Tra My · Tran Hoang Linh
Academic supervisors: Hoang Tuan Anh and Joshua Dwight

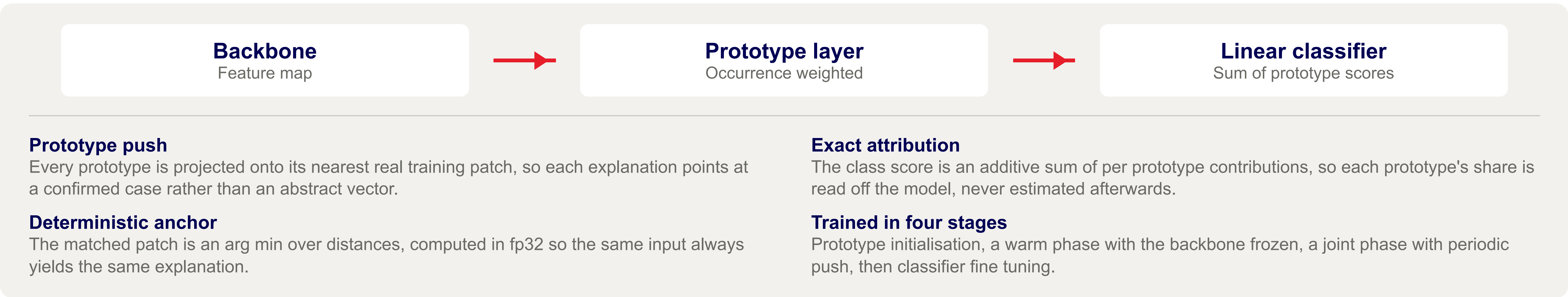
School of Science, Engineering & Technology | RMIT University Vietnam

The problem

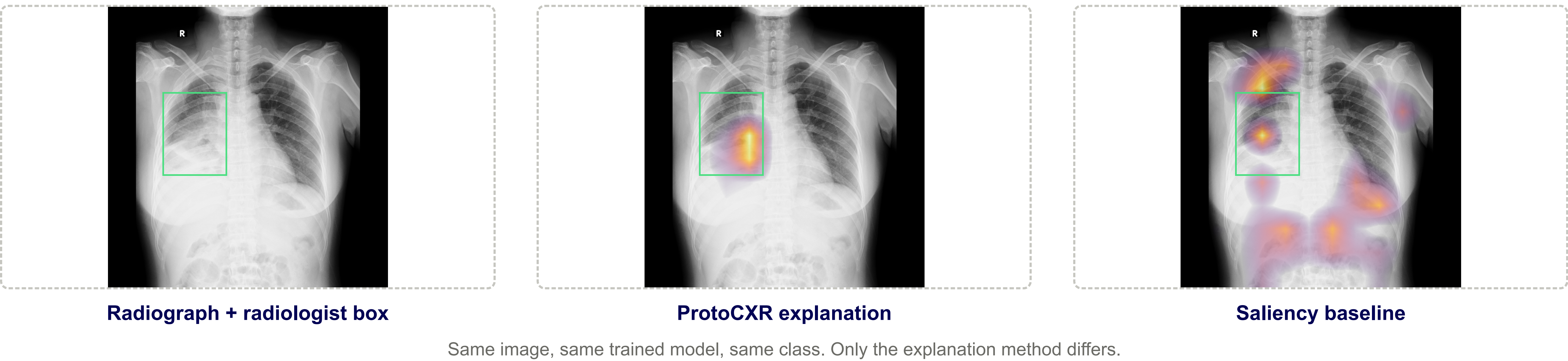
Deep learning models now match radiologists on many chest radiograph tasks, but a clinician cannot safely act on a diagnosis they cannot interrogate.

The common remedy is a saliency map drawn after training, colouring in where the model supposedly looked. Such maps barely change even when the model's weights are destroyed, so they may not be describing the model at all.

How it works



Where the model actually looked



What we measured

Prototype grounding
 Distance from each prototype to its nearest real training patch

Lesion localisation
 Agreement with boxes drawn by radiologists

Region coverage
 Share of activation falling inside the lesion

Model dependence
 Does the map change when the weights are destroyed

Where it falls short

Boundary agreement
 Localisation overlap is low in absolute terms for every method tested, including the saliency baselines. Models find roughly the right region without matching boundaries.

Accuracy
 Classification sits below black box models trained on a single dataset. The contribution here is measured interpretability, not peak predictive performance.

Evaluation
 Metrics are functionally grounded. A clinical reader study relating explanation quality to radiologist judgement is left for future work.

The explanation is the computation

Interpretability should be measured, not illustrated. ProtoCXR is a reproducible protocol and an inherently interpretable baseline for chest radiograph research.

