

1. Record Nr.	UNISA996495570803316
Titolo	Interpretability of machine intelligence in medical image computing : 5th international workshop, iMIMIC 2022, held in conjunction with MICCAI 2022, Singapore, Singapore, September 22, 2022, proceedings // edited by Mauricio Reyes, Pedro Henriques Abreu, and Jaime Cardoso
Pubbl/distr/stampa	Singapore : , : Springer, , [2022] ©2022
ISBN	3-031-17976-5
Descrizione fisica	1 online resource (134 pages)
Collana	Lecture Notes in Computer Science ; ; v.13611
Disciplina	616.0754
Soggetti	Computer-assisted surgery Diagnostic imaging - Data processing
Lingua di pubblicazione	Inglese
Formato	Materiale a stampa
Livello bibliografico	Monografia
Note generali	Includes index.
Nota di contenuto	Intro -- Preface -- Organization -- Contents -- Interpretable Lung Cancer Diagnosis with Nodule Attribute Guidance and Online Model Debugging -- 1 Introduction -- 2 Materials -- 3 Methodology -- 3.1 Collaborative Model Architecture with Attribute-Guidance -- 3.2 Debugging Model with Semantic Interpretation -- 3.3 Explanation by Attribute-Based Nodule Retrieval -- 4 Experiments and Results -- 4.1 Implementation -- 4.2 Quantitative Evaluation -- 4.3 Trustworthiness Check and Interpretable Diagnosis -- 5 Conclusions -- References -- Do Pre-processing and Augmentation Help Explainability? A Multi-seed Analysis for Brain Age Estimation -- 1 Introduction -- 2 Related Work -- 3 Methods -- 4 Results -- 4.1 Performance -- 4.2 Voxel Agreement -- 4.3 Atlas-Based Analyses -- 4.4 Region Validation -- 5 Conclusion -- References -- Towards Self-explainable Transformers for Cell Classification in Flow Cytometry Data -- 1 Introduction -- 2 Related Work -- 3 Methods -- 3.1 Architecture -- 3.2 Preprocessing -- 3.3 Loss Function -- 3.4 Data Augmentation -- 4 Experiments -- 4.1 Data -- 4.2 Results -- 5 Conclusion -- References -- Reducing Annotation Need in Self-explanatory Models for Lung Nodule Diagnosis -- 1 Introduction -- 2 Method -- 3 Experimental Results -- 3.1 Prediction

Performance of Nodule Attributes and Malignancy -- 3.2 Analysis of
Extracted Features in Learned Space -- 3.3 Ablation Study -- 4
Conclusion -- References -- Attention-Based Interpretable Regression
of Gene Expression in Histology -- 1 Introduction -- 2 Methods -- 2.1
Datasets -- 2.2 Multiple Instance Regression of Gene Expression -- 2.3
Attention-Based Model Interpretability -- 2.4 Evaluation of
Performance and Interpretability -- 3 Experiments and Results -- 3.1
Network Training -- 3.2 Quantitative Model Evaluation -- 3.3
Attention-Based Identification of Hotspots and Patterns.
3.4 Quantitative Evaluation of the Attention -- 4 Discussion -- 5
Conclusion -- A Description of Selected Genes -- B Detailed Model
Evaluation -- C Additional Visualizations -- D Single-Cell Co-
expression -- References -- Beyond Voxel Prediction Uncertainty:
Identifying Brain Lesions You Can Trust -- 1 Introduction -- 2 Our
Framework: Graph Modelization for Lesion Uncertainty Quantification
-- 2.1 Monte Carlo Dropout Model and Voxel-Wise Uncertainty -- 2.2
Graph Dataset Generation -- 2.3 GCNN Architecture and Training -- 3
Material and Method -- 3.1 Data -- 3.2 Comparison with Known
Approaches -- 3.3 Evaluation Setting -- 3.4 Implementation Details --
4 Results and Discussion -- 5 Conclusion -- References --
Interpretable Vertebral Fracture Diagnosis -- 1 Introduction -- 1.1
Related Work -- 2 Methodology -- 2.1 Vertebral Fracture Detection --
2.2 Semantic Concept Extraction (Correlation) -- 2.3 Visualization of
Highly Correlating Concepts at Inference -- 3 Experimental Setup -- 4
Results and Discussion -- 4.1 Clinical Meaningfulness of Extracted
Semantic Concepts -- 4.2 Single-Inference Concept Visualization -- 5
Conclusion -- References -- Multi-modal Volumetric Concept
Activation to Explain Detection and Classification of Metastatic Prostate
Cancer on PSMA-PET/CT -- 1 Introduction -- 2 Data -- 3 Method --
3.1 Preprocessing -- 3.2 Detection -- 3.3 Classification -- 3.4
Explainable AI -- 4 Results -- 4.1 Detection -- 4.2 Classification -- 4.3
Explainable AI -- 5 Discussion -- 6 Conclusion -- References -- KAM -
A Kernel Attention Module for Emotion Classification with EEG Data --
1 Introduction -- 2 Related Work -- 3 Kernel Attention Module -- 4
Experiments -- 5 Conclusion -- References -- Explainable Artificial
Intelligence for Breast Tumour Classification: Helpful or Harmful -- 1
Introduction -- 2 Related Work -- 2.1 XAI in Medicine.
3 Model Setup -- 3.1 Data Pre-Processing -- 3.2 Model Architecture --
4 Explanations -- 4.1 LIME -- 4.2 RISE -- 4.3 SHAP -- 5 Evaluating
Explanations -- 5.1 One-Way ANOVA -- 5.2 Kendall's Tau -- 5.3
Radiologist Evaluation -- 5.4 Threats to Validity -- 6 Observations and
Discussion -- 6.1 Discussion -- A Appendix -- A.1 Model Training
Results -- A.2 Choosing L Parameter for LIME -- A.3 One-Way ANOVA
Results -- A.4 Pixel Agreement Statistics -- A.5 Ranked Biased Overlap
(RBO) Results -- A.6 Kendall's Tau Results -- A.7 Radiologist Opinions
-- A.8 Explanation Examples -- References -- Author Index.
