

## SHORT PAPER SESSION I2

### I2.2 AI-assisted development of an educational atlas for acute CT brain imaging

*Dr Muath Ibrahim<sup>1,2</sup>, Dr David Rosewarne<sup>1,3</sup>, Dr Sahand Hooshmand<sup>1,2</sup>, Dr Ziba Gandomkar<sup>1,2</sup>, Dr Tan Hasdarngkul, Mr Zhufeng Shi, Mr Anirudh Singh Bhati, Ms Ella Grew, Dr Diogo J. Vidal Silva<sup>1</sup>, Prof Patrick Brennan<sup>1,2</sup>*

<sup>1</sup>DetectedX, Camperdown, Australia, <sup>2</sup>University of Sydney, Camperdown, Australia, <sup>3</sup>The Royal Wolverhampton NHS Trust, Wolverhampton, United Kingdom, <sup>4</sup>Norfolk and Norwich University Hospital NHS Foundation Trust, Norfolk, United Kingdom

#### Background

Radiology education faces pressure from increasing trainee numbers, limited educator time, and restricted access to high-quality case material from clinical archives. Scalable, structured educational resources are needed in diagnostic imaging. Advances in artificial intelligence (AI), particularly large language models (LLMs), offer opportunities to support efficient development of educational content. This project aimed to design and evaluate an AI-assisted pipeline to create a structured atlas for acute CT brain imaging education.

#### Method

A two-stage AI-driven workflow was developed. First, a hierarchical atlas structure was defined by chapters, pathologies, and topic-specific subheadings. A scripted Python pipeline generated targeted prompts to an LLM to produce draft teaching text for each subsection. Content underwent iterative human review by radiology trainees and consultant radiologists to ensure clinical accuracy and remove hallucinated or irrelevant material.

Second, illustrative imaging cases were identified from a large CT brain PACS archive. Rule-based natural language processing of radiology reports located examinations likely to demonstrate relevant pathologies. Selected anonymised cases were curated and integrated into the atlas to support case-based learning.

#### Results

This process produced a structured Atlas of Acute CT Brain Imaging combining expert-reviewed AI-generated text with clinical imaging cases. The atlas is undergoing pilot use and refinement in multiple training settings, with early feedback indicating educational value for radiology trainees and clinicians involved in acute neuroimaging.

#### Conclusion

An AI-assisted pipeline can support scalable development of structured radiology educational resources. The framework is adaptable across imaging subspecialties and may help address training capacity challenges.

### I2.3 Transforming Radiology Careers: An AI-Driven, Comprehensive Support System for Overseas Trainee Doctors and Consultants Joining the NHS

*Dr Rayya Wael Musa Naffa<sup>1</sup>, Ebinesh Arulnathan<sup>2</sup>, Sivakumar Manickam<sup>1</sup>*

<sup>1</sup>Mid And South Essex NHS Foundation, Essex, United Kingdom, <sup>2</sup>Alder Hey Children's NHS Foundation Trust, United Kingdom

#### Purpose

Internationally trained radiologists comprise 35% of the NHS workforce but face challenges adapting to NHS workflows, including PACS/RIS navigation, reporting standards, MDT preparation, audit/QI requirements, communication norms, and high-acuity service delivery. Existing induction processes are variable and often insufficient within the standard three-month probationary period. This pilot study evaluated an AI-enabled support system to facilitate professional integration.

#### Materials and Methods

A four-stage framework was implemented.

- (1) Challenge mapping used standardised questionnaires from 50 overseas radiologists across Southeast England to identify adaptation barriers and coping strategies.
- (2) AI capacity Building: These data were thematically analysed and structured into prompts to train a customised ChatGPT-based radiology support chatbot.
- (3) Testing: The system was tested by 20 radiologists within their first year in the NHS, supporting report drafting, urgency triage, MDT preparation, audit/QI workflows, communication templates, IT navigation, and clinical governance.
- (4) Post-implementation evaluation used structured questionnaires and interviews to assess usability, accuracy, relevance, and clinical value.

#### Results

Participants reported high perceived usefulness, describing the tool as accessible, intuitive, and non-judgemental. It enhanced confidence in reporting, MDT preparation, and system navigation. Suggested improvements included image-based examples, localisation to departmental protocols, and more scenario-driven responses.

#### Conclusions

This pilot demonstrates the feasibility of an AI-driven, personalised support system to address key adaptation challenges for overseas radiologists. With further validation, it could be implemented nationally to improve integration, efficiency, and professional confidence, and inform NHS induction frameworks and Royal College of Radiologists guidance.

## 12.4 Evaluation of synthetic CT for paediatric bony facial asymmetry analysis

Ethan Pierce<sup>1</sup>, Alasdair Wilkinson<sup>1</sup>, Nina Joyce<sup>1</sup>, Abigail Crowther<sup>1</sup>, Tom Melichar<sup>1</sup>, Peter Sitch<sup>2</sup>, Lucy Siew Chen Davies<sup>1,3</sup>, Shermaine Pan<sup>3,4</sup>, Ed Smith<sup>1,2,3,4</sup>, Eliana Vasquez Osorio<sup>1</sup>, Marianne Aznar<sup>1</sup>, Angela Davey<sup>1</sup>, Chelsea Sargeant<sup>1</sup>

<sup>1</sup>University of Manchester, Manchester, United Kingdom, <sup>2</sup>The Christie Proton Beam Therapy Centre, The Christie NHS Foundation Trust, Manchester, United Kingdom, <sup>3</sup>Radiotherapy, The Christie NHS Foundation Trust, Manchester, United Kingdom, <sup>4</sup>Department of Clinical Oncology, The Christie NHS Foundation Trust, Manchester, United Kingdom

### Background

Bony facial asymmetry is a common late effect of radiotherapy in paediatric head and neck patients, leading to reduced post-treatment quality of life. Investigation of dose-response relationships is limited by the absence of standardised measurement procedures. Follow-up imaging is predominantly MR to minimise additional radiation exposure, restricting direct assessment of bony anatomy. Deep learning-based synthetic CT (sCT) may enable quantitative asymmetry analysis. This study evaluated the suitability of an open-source sCT model for bony facial asymmetry analysis.

### Method

Nine paediatric MRIs (aged 2-19), acquired during proton therapy treatment planning, were used to generate sCTs. These were produced using an open-source model developed for the SynthRAD2025 challenge and trained primarily on defaced adult data [1, 2]. sCTs were compared with paired planning CT scans using global and local image similarity metrics including mean absolute error (MAE). Anatomical landmarking assessed the impact of image modality on craniofacial growth-centre identification.

### Results

Visual inconsistencies between sCT and CT images were frequently observed, particularly in high-density regions such as teeth. Poor bony structure reproduction was reflected in local image metric performance, while global metrics were dominated by accurate soft-tissue agreement (Figure 1). No improvement in anatomical landmark identification was observed with sCT compared with MR (Figure 2).

### Conclusion

The evaluated model did not reliably reproduce bony facial structures and is unsuitable for facial asymmetry analysis in its current form. Limited generalisability likely reflects training on defaced adult datasets and application to paediatric anatomy. Targeted model adaptation may improve performance; further investigation is ongoing.

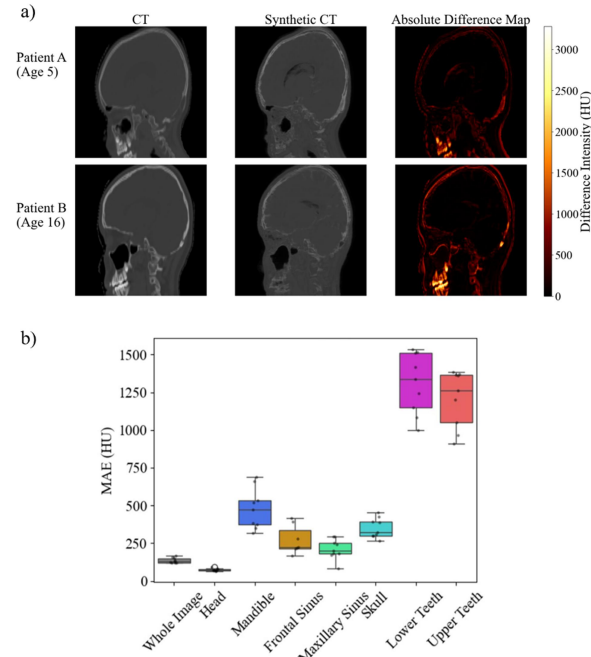


Figure 1: (a) Two example synthetic CT (sCT) images accompanied by paired CTs. Difference maps generated from these paired images are shown, with signal intensity measured in Hounsfield Units (HU). MR-related artifacts are present in the brain region, doubling is present at the top of the skull and signal intensity failures are observed in the dental region. (b) A comparison of the mean absolute error (MAE) calculated across the dataset for the whole head and local structures in Hounsfield Unites (HU). Local MAE values for bony tissue are consistently higher than those for the head, indicating superior soft tissue reproduction. Higher MAE values are also observed for the dental region.

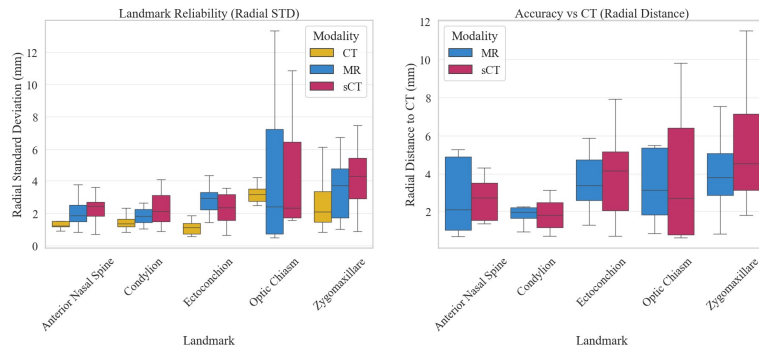


Figure 2: Landmarking performance across modalities for each landmark. Left: landmark reliability measured by radial standard deviation across repeated annotations. Right: landmark accuracy measured as radial distance to CT ground truth for MR and synthetic CT (sCT).

1. Boussot, V., Hémon, C., Nunes, J.C. and Dillenseger, J.L., 2025. Why Registration Quality Matters: Enhancing sCT Synthesis with IMPACT-Based Registration. arXiv preprint arXiv:2510.21358.
2. Thummerer, A., van der Bijl, E., Galapon, A.J., Kamp, F., Savenije, M., Muijs, C., Aluwini, S., Steenbakkers, R.J., Beuel, S., Intven, M.P. and Langendijk, J.A., 2025. SynthRAD2025 Grand Challenge dataset: Generating synthetic CTs for radiotherapy from head to abdomen. Medical physics, 52(7), p.e17981.

## 12.5 Woman vs Machine: Can AI be used to aid or replace title and abstract screening for literature reviews in healthcare?

[Sophie O'Riordan<sup>1</sup>](#), [Danielle Cartwright<sup>1</sup>](#)

<sup>1</sup>The Clatterbridge Cancer Centre NHS Foundation Trust, Liverpool, United Kingdom

Screening papers for literature reviews is time-intensive. This project aimed to evaluate the use of Copilot 2025<sup>1</sup> to screen titles and abstracts for a study on adaptive radiotherapy (ART) and to develop a prompt template that could be used for further healthcare research.

Prompt variants from Syriani et. al (2024)<sup>2</sup> were adapted to screen 130 and 92 papers related to gynaecological and bladder ART respectively, with PRISMA manual screening as the comparator. Prompts were favoured that achieved minimal false negatives, to minimise the exclusion of relevant papers. The repeatability of the best performing prompts was investigated, and final template was validated on an independent dataset of 376 papers.

Table 1 shows that the 'Simple' prompt produced results with least false negatives across the datasets (0%-0.769%) and was the most repeatable (83.7% unanimous decisions across five repeats). Overall agreement was lower than some prompts, as was the Cohen's kappa score. This being due to the higher number of false positives (18.4%-39.1%). Yet, fair-moderate agreement ( $\kappa=0.252-0.512$ ) was still demonstrated. The 'Simple' prompt was estimated to save 34.8 hours for the validation set, extrapolated from the number of true negatives and an estimated 7.8 minutes to review each paper manually.

The trade-off between sensitivity and specificity was significant. Reducing false negatives resulted in increased false positives. However, utilising the 'Simple' prompt template before manual screening removes many true negatives and offers significant time saving. This study concludes AI must be seen as a preliminary tool to aid, rather than replace, manual screening.

Table 1:

|                                                   | Prompt Variants |         |          |           |          |              |
|---------------------------------------------------|-----------------|---------|----------|-----------|----------|--------------|
|                                                   | Simple          | SimpleX | Positive | PositiveX | Balanced | BalancedX    |
| <i>130 gynaecological papers</i>                  |                 |         |          |           |          |              |
| <b>TP</b>                                         | <b>52</b>       | 47      | 50       | 48        | 45       | 46           |
| <b>FP</b>                                         | 32              | 22      | 34       | 28        | 26       | <b>19</b>    |
| <b>TN</b>                                         | 45              | 55      | 43       | 49        | 51       | <b>58</b>    |
| <b>FN</b>                                         | <b>1</b>        | 6       | 3        | 5         | 8        | 7            |
| <b>Sensitivity</b>                                | <b>0.981</b>    | 0.887   | 0.943    | 0.906     | 0.849    | 0.868        |
| <b>Specificity</b>                                | 0.584           | 0.714   | 0.558    | 0.636     | 0.662    | <b>0.753</b> |
| <b>% Agree</b>                                    | 74.6            | 78.5    | 71.5     | 74.6      | 73.9     | <b>80.0</b>  |
| <b>Cohen's kappa</b>                              | 0.518           | 0.574   | 0.460    | 0.508     | 0.486    | <b>0.600</b> |
| <i>92 bladder papers</i>                          |                 |         |          |           |          |              |
| <b>TP</b>                                         | <b>43</b>       | 38      |          | 27        |          | 39           |
| <b>FP</b>                                         | 36              | 25      |          | 11        |          | 29           |
| <b>TN</b>                                         | 13              | 24      |          | 38        |          | <b>10</b>    |
| <b>FN</b>                                         | <b>0</b>        | 5       |          | 16        |          | <b>39</b>    |
| <b>Sensitivity</b>                                | <b>1.00</b>     | 0.884   |          | 0.628     |          | 0.674        |
| <b>Specificity</b>                                | 0.265           | 0.490   |          | 0.776     |          | <b>0.796</b> |
| <b>% Agree</b>                                    | 60.9            | 67.4    |          | 70.7      |          | <b>73.9</b>  |
| <b>Cohen's kappa</b>                              | 0.252           | 0.363   |          | 0.406     |          | <b>0.473</b> |
| <i>376 gynaecological papers (validation set)</i> |                 |         |          |           |          |              |
| <b>TP</b>                                         | 46              |         |          |           |          |              |
| <b>FP</b>                                         | 69              |         |          |           |          |              |
| <b>TN</b>                                         | 268             |         |          |           |          |              |
| <b>FN</b>                                         | 2               |         |          |           |          |              |
| <b>Sensitivity</b>                                | 0.958           |         |          |           |          |              |
| <b>Specificity</b>                                | 0.817           |         |          |           |          |              |
| <b>% Agree</b>                                    | 83.5            |         |          |           |          |              |
| <b>Cohen's kappa</b>                              | 0.512           |         |          |           |          |              |

\***TP** = True positive, **FP** = False positive, **TN** = True negative, **FN** = False negative, **Sensitivity** = Proportion of papers that were correctly included out of all that should have been included, **Specificity** = Proportion of papers that were correctly excluded out of all that should have been excluded, **% Agree** – How many of the include/exclude decisions agreed with the manual decisions.

1. Copilot (2025), Available at: <https://copilot.microsoft.com>, Accessed: 28/07/25 – 28/08/25, Think Deeper settings enabled.

2. Syriani, E., David, I. and Kumar, G. (2024) Screening articles for systematic reviews with ChatGPT, Journal of Computer Languages, 80, p. 101287. Available at: <https://doi.org/10.1016/j.cola.2024.101287>.

## 12.6 The RAPID-LC study: Perceptions of Artificial Intelligence in the Lung Cancer Diagnostic Pathway: A Cross Sectional Survey of Patients and Clinicians in Northern Ireland

[Dr Laura McLaughlin<sup>1</sup>](#), [Dr Sonyia McFadden<sup>2</sup>](#), [Dr Clare Rainey<sup>1</sup>](#)

<sup>1</sup>University College Cork, Cork, Ireland, <sup>2</sup>Ulster University, Belfast, United Kingdom

Lung cancer remains the leading cause of cancer mortality in the UK, with the majority of cases diagnosed at an advanced stage. Artificial intelligence (AI) enabled imaging tools have been proposed to support radiology and radiography practice by improving detection, triage and workflow efficiency.

A cross-sectional survey study was conducted to explore perceptions of AI use within the lung cancer imaging pathway. Two e-questionnaires were distributed: one to members of the public with recent experience of diagnostic imaging (n = 169) and one to clinicians involved in lung cancer care (n = 45). Surveys assessed attitudes towards AI-assisted imaging, trust, perceived impact on time-to-diagnosis, workforce implications and preferences for AI as a decision-support tool.

Public respondents were supportive of AI integration. Many perceived AI as having the potential to reduce time to diagnosis. Clinicians demonstrated cautious optimism, with 51% expressing positive views on AI use. Support was strongest for AI applications in (i) prioritisation (ii) detection of incidental findings and (iii) provision of a second opinion. Both groups emphasised the importance of human oversight, patient–clinician communication, data governance and validation.

Findings indicate broad support among UK service users and imaging professionals for AI as an assistive tool within the lung cancer diagnostic pathway. For radiographers and radiologists, AI is viewed as a means to support workflow efficiency and patient safety rather than replace clinical judgement. These insights are directly relevant to imaging practice, service planning and the safe, patient-centred implementation of AI within UK cancer pathways.

---