

experts. **METHODS/STUDY POPULATION:** Common to many medical imaging projects, we have a small number of expert-marked patient photos ($N = 36$, $n = 360$), and many unmarked photos ($N = 337$, $n = 25,842$). Dark skin (Fitzpatrick type 4+) is underrepresented in both sets; 11% of patients in the marked set and 9% in the unmarked set. In addition, a set of 20 expert-marked photos from 20 patients were withheld from training to assess model performance, with 20% dark skin type. Our gold standard markings were manual contours around affected skin by a trained expert. Three AI training methods were tested. Our established baseline uses only the small number of marked photos (supervised method). The semi-supervised method uses a mix of marked and unmarked photos with human feedback. The self-supervised method uses only unmarked photos without any human feedback. **RESULTS/ANTICIPATED RESULTS:** We evaluated performance by comparing predicted skin areas with expert markings. The error was given by the absolute difference between the percentage areas marked by the AI model and expert, where lower is better. Across all test patients, the median error was 19% (interquartile range 6 – 34) for the supervised method and 10% (5 – 23) for the semi-supervised method, which incorporated unmarked photos from 83 patients. On dark skin types, the median error was 36% (18 – 62) for supervised and 28% (14 – 52) for semi-supervised, compared to a median error on light skin of 18% (5 – 26) for supervised and 7% (4 – 17) for semi-supervised. Self-supervised, using all 337 unmarked patients, is expected to further improve performance and consistency due to increased data diversity. Full results will be presented at the meeting. **DISCUSSION/SIGNIFICANCE OF IMPACT:** By automating skin assessment for cGVHD, AI could improve accuracy and consistency compared to manual methods. If translated to clinical use, this would ease clinical burden and scale to large patient cohorts. Future work will focus on ensuring equitable performance across all skin types, providing fair and accurate assessments for every patient.

27

Automated IRB compliance and secure data delivery in i2b2

Hakob Abajian, Venkat Vuppula Johnson Kwong, Praveen Angyan, Shakeel Usmani, Saurav Limbu Brittney Arvizu Hira Sherazi, Swaroop Samek, Maryam Abdallah, Amy Chuang, Daniella Garofalo and Neil Bahroos
University of Southern California

OBJECTIVES/GOALS: To address the manual, time-consuming processes of validating IRB compliance and ensuring the secure delivery of i2b2 data, this project automates compliance checks, streamlines Protected Health Information (PHI) access, and provides timely, secure data availability while reducing administrative burdens and non-compliance risks. **METHODS/STUDY POPULATION:** This project enhances the i2b2 application to automate compliance processes and facilitate secure data delivery through integration with REDCap. By linking i2b2 with the IRB system, the application performs automatic compliance checks for project requests, verifying GCP and HIPAA certifications, only allowing the release of IRB-approved PHI variables, safeguarding against unauthorized data access. Manual signatures confirm non-automated compliance processes. Once verified, the application

automatically creates a REDCap project, assigns user access, and securely delivers data, ensuring compliance with HIPAA regulations. **RESULTS/ANTICIPATED RESULTS:** The automated system successfully streamlined IRB compliance checks and data delivery for i2b2 requests. Validation of certifications like GCP and HIPAA, now occurs automatically, significantly reducing the risk of non-compliance. Personnel access to data is limited to IRB-approved PHI, ensuring data security and adherence to institutional standards. The integration with REDCap has reduced manual processes, cutting data request processing time to approximately 30 minutes. Researchers and administrative staff experienced a notable decrease in administrative burden, with faster, more efficient access to approved data while maintaining full compliance with IRB and HIPAA regulations. **DISCUSSION/SIGNIFICANCE OF IMPACT:** The lessons learned can be adapted by institutions to improve compliance efficiency and reduce administrative overhead. Implementing similar automation of certification checks and data delivery, sites can enhance data security, minimize errors, and ensure faster, compliant access to research data.

28

Using AI to predict molecular subtype from histopathology slides in endometrial cancer[†]

Vincent Wagner¹, Jesus Gonzalez, Bosquet², Michael Goodheart², Xiaodong Wu³ and Megan Samuelson⁴

¹University of Iowa; ²Division of Gynecologic Oncology, University of Iowa; ³Department of Electrical and Computer Engineering, University of Iowa and ⁴Department of Surgical Pathology, University of Iowa

OBJECTIVES/GOALS: Endometrial cancer is one of the few cancers that has both a rising incidence and mortality rate. Molecular classification is becoming more important for the management of endometrial cancer but the ability to translate this into clinical practice remains constrained. Our goal is to use AI to predict the molecular subtype from histopathology slides. **METHODS/STUDY POPULATION:** We utilized the open source endometrial cancer datasets from The Cancer Genome Atlas (TCGA) ($N = 387$) and Cancer Proteomics Transcriptomic Tumor Analysis Consortium (CPTAC) ($N = 135$) to develop and train a vision transformer AI model. We used a proprietary cohort of patients ($N = 548$) for external validation. Whole slide images (WSI) and molecular subtype data were collected. Subtypes include POLE ultramutated (POLE), microsatellite instability (MSI-H), copy-number low (CNV-L), and copy-number high (CNV-H). WSI were preprocessed, and features were extracted. Modified STAMP protocol was used in training, utilizing a pretrained foundation transformer model (Virchow2). Cross-validation of the TCGA was used for initial training, followed by testing on the CPTAC dataset and validation on our proprietary cohort. **RESULTS/ANTICIPATED RESULTS:** Fivefold cross-validation of the TCGA database (60% training, 20% testing, and 20% validation) developed a best overall model with a mean AUC of 0.74 (POLE 0.78, MSI-H 0.76, CNV-H 0.86, CNV-L 0.77). Overall precision 0.58, recall 0.55. CNV-H was the subtype with the most accurate prediction. CPTAC holdout testing revealed moderately high AUC (POLE 0.63, MSI-H 0.62, CNV-H 0.98, and CNV-L 0.76). Overall precision 0.54 and recall 0.58. Again, CNV-H was the most accurate