

Automating clinical assessments of memory deficits: Deep Learning based scoring of the Rey-Osterrieth Complex Figure

Reviewed Preprint

v1 • June 21, 2024

Not revised

Nicolas Langer , Maurice Weber, Bruno Hebling Vieira, Dawid Strzelczyk, Lukas Wolf, Andreas Pedroni, Jonathan Heitz, Stephan Müller, Christoph Schultheiss, Marius Tröndle, Juan Carlos Arango-Lasprilla, Diego Rivera, Federica Scarpina, Qianhua Zhao, Rico Leuthold, Flavia Wehrle, Oskar G. Jenni, Peter Brugger, Tino Zaehle, Romy Lorenz, Ce Zhang

Methods of Plasticity Research, Department of Psychology, University of Zurich, Zurich, Switzerland • University Research Priority Program (URPP) Dynamics of Healthy Aging, Zurich, Switzerland • Neuroscience Center Zurich (ZNZ), University of Zurich and ETH Zurich, Zurich, Switzerland • Department of Computer Science, ETH, Zurich, Switzerland • Virginia Commonwealth University, Richmond, Virginia • Department of Health Science, Public University of Navarre, Pamplona, Spain • Instituto de Investigación Sanitaria de Navarra (IdiSNA), Pamplona, Spain • "Rita Levi Montalcini" Department of Neurosciences, University of Turin, Italy • I.R.C.C.S. Istituto Auxologico Italiano, U.O. di Neurologia e Neuroriabilitazione, Ospedale San Giuseppe, Piancavallo (VCO), Italy • Huashan Hospital, Shanghai, China • Smartcode, Zurich, Switzerland • University Children's Hospital Zurich, Child Development Center, Zurich, Switzerland • Rehabilitation Center, Valens, Switzerland • University Hospital Magdeburg University Department of Neurology, Magdeburg, Germany • MRC Cognition and Brain Sciences Unit, University of Cambridge, Cambridge, United Kingdom • Department of Neurophysics, Max Planck Institute for Human Cognitive and Brain Sciences, Leipzig, Germany • Stanford University, Stanford, CA, United States

 https://en.wikipedia.org/wiki/Open_access

 Copyright information

Abstract

Background

Memory deficits are a hallmark of many different neurological and psychiatric conditions. The Rey-Osterrieth complex figure (ROCF) is the state-of-the-art assessment tool for neuropsychologists across the globe to assess the degree of non-verbal visual memory deterioration. To obtain a score, a trained clinician inspects a patient's ROCF drawing and quantifies deviations from the original figure. This manual procedure is time-consuming, slow and scores vary depending on the clinician's experience, motivation and tiredness.

Methods

Here, we leverage novel deep learning architectures to automatize the rating of memory deficits. For this, we collected more than 20k hand-drawn ROCF drawings from patients with various neurological and psychiatric disorders as well as healthy participants. Unbiased ground truth ROCF scores were obtained from crowdsourced human intelligence. This dataset was used to train and evaluate a multi-head convolutional neural network.

Results

The model performs highly unbiased as it yielded predictions very close to the ground truth and the error was similarly distributed around zero. The neural network outperforms both online raters and clinicians. The scoring system can reliably identify and accurately score individual figure elements in previously unseen ROCF drawings, which facilitates explainability of the AI-scoring system. To ensure generalizability and clinical utility, the model performance was successfully replicated in a large independent prospective validation study that was pre-registered prior to data collection.

Conclusions

Our AI-powered scoring system provides healthcare institutions worldwide with a digital tool to assess objectively, reliably and time-efficiently the performance in the ROCF test from hand-drawn images.

eLife assessment

This study presented a **valuable** inventory in scoring a neuropsychological test, ROCFT. The level of evidence is **compelling**. The authors constructed large samples from multi-center international researchers and tested the model using internet data with excellent performance. Their deep learning method could potentially apply to neuropsychological tests as well as other related fields.

<https://doi.org/10.7554/eLife.96017.1.sa3>

Introduction

Neurological and psychiatric disorders are among the most common and debilitating illnesses across the lifespan. In addition, the aging of our population, with the increasing prevalence of physical and cognitive disorders, poses a major burden on our society with an estimated economic cost of 2.5 trillion US\$ per year (Trautmann, Rehm, and Wittchen 2016). Currently, neuropsychologists typically use paper-pencil tests to assess individual neuropsychological functions and brain dysfunctions, including memory, attention, reasoning and problem-solving. Most neuropsychologists around the world use the Rey-Osterrieth complex figure (ROCF) in their daily clinical practice (Rabin, Barr, and Burton 2005; Rabin, Paolillo, and Barr 2016), which provides insights into a person's nonverbal visuo-spatial memory capacity in healthy and clinical populations of all ages, from childhood to old age (Shin et al. 2006).

Our estimation revealed that a single neuropsychological division (e.g. at the University Hospital Zurich) scores up to 6000 ROCF drawing per year. The ROCF test has several advantages as it does not depend on auditory processing and differences in language skills that are omnipresent in cosmopolitan societies (Osterrieth 1944; Somerville, Tremont, and Stern 2000). The test exhibits adequate psychometric properties (e.g. sufficient test-retest reliability (John E. Meyers and Volbrecht 1999; Levine et al. 2004) and internal consistency (Berry, Allen, and Schmitt 1991;

Fastenau, Bennett, and Denburg 1996)). Furthermore, the ROCF test has demonstrated to be sensitive to discriminate between various clinical populations (Alladi et al. 2006) and the progression of Alzheimer's disease (Trojano and Gainotti 2016).

The ROCF test consists of three test conditions: First, in the *Copy condition* subjects are presented with the ROCF and are asked to draw a copy of the same figure. Subsequently, the ROCF figure and the drawn copy are removed and the subject is instructed to reproduce the figure from memory immediately (Shin et al., 2006) or 3 min after the Copy condition (Meyers J, Meyers K, 1996) (*Immediate Recall condition*). After a delay of 30 minutes, the subject is required to draw the same figure once again (*Delayed Recall condition*). For further description of the clinical usefulness of the ROCF please refer to Shin and colleagues (Shin et al. 2006).

The current quantitative scoring system (Meyers, Bayless, and Meyers 1996) splits the ROCF into 18 identifiable elements (see **Figure 1A** [↗](#)), each of which is considered separately and marked on the accuracy in both distortion and placement according to a set of rules and scored between 0 and 2 points (see supplementary Figure S1). Thus, the scoring is currently undertaken manually by a trained clinician, who inspects the reproduced ROCF drawing and tracks deviations from the original figure, which can take up to 15 minutes per figure (individually *Copy*, *Immediate Recall* and *Delayed Recall* conditions).

One major limitation of this quantitative scoring system is that the criteria of what position and distortion is considered “accurate” or “inaccurate” may vary from clinician to clinician (Groth-Marnat 2000; Shin et al. 2006; Canham, Smith, and Tyrrell 2000). In addition, the scoring might vary as a function of motivation and tiredness or because the clinicians may be unwittingly influenced by interaction with the patient. Therefore, an automated system that offers reliable, objective, robust and standardized scoring, while saving clinicians' time, would be desirable from an economic perspective and more importantly leads to more accurate scoring and subsequently diagnosing.

In recent years, computerized image analysis and machine learning methods have entered the clinical neuropsychological diagnosis field providing the potential for establishing quantitative scoring methods. Using machine-learning methods, such as convolutional neural networks and support vector machines, studies have successfully recognized visual structures of interest produced by subjects in the Bender Gestalt Test (BGT) (Nazar et al. 2017), and the Clock Draw Task (CDT) (Kim, Cho, and Do 2010; Harbi, Hicks, and Setchi 2016) - both are less frequently applied screening tests for visuospatial and visuo-constructive (dis-)abilities (Webb et al. 2021). Given the wide application of the ROCF, it is not surprising that we are not the first to take steps towards a machine-based scoring system. Canham et al. (2000) have developed an algorithm to automatically identify a selection of parts of the ROCF (6 of 18 elements) with great precision. This approach provides first evidence for the feasibility of automated segmentation and feature recognition of the ROCF. More recently, Vogt et al. (2019) developed an automated scoring using a deep neural network. The authors reported a $r=0.88$ Pearson correlation with human ratings, but equivalence testing demonstrated that the automated scoring did not produce strictly equivalent total scores compared to the human ratings. Moreover, it is unclear if the reported correlation was observed in an independent test data set, or in the training set. Finally, Petilli et al. (2021) have proposed a novel tablet-based digital system for the assessment of the ROCF task, which provides the opportunity to extract a variety of parameters such as spatial, procedural, and kinematic scores. However, none of the studies described have been able to produce machine-based scores according to the original scoring system currently used in clinics (Meyers, Bayless, and Meyers 1996) that are equivalent or superior to human raters. A major challenge in developing automated scoring systems is to gather a rich enough set of data with instances of all possible mistakes that humans can make. In this study, we have trained convolutional neural networks with over 20'000 digitized ROCFs from different populations regarding age and diagnostic status (healthy individuals or individuals with neurological and psychiatric disorder (e.g. Alzheimer, Parkinson)).

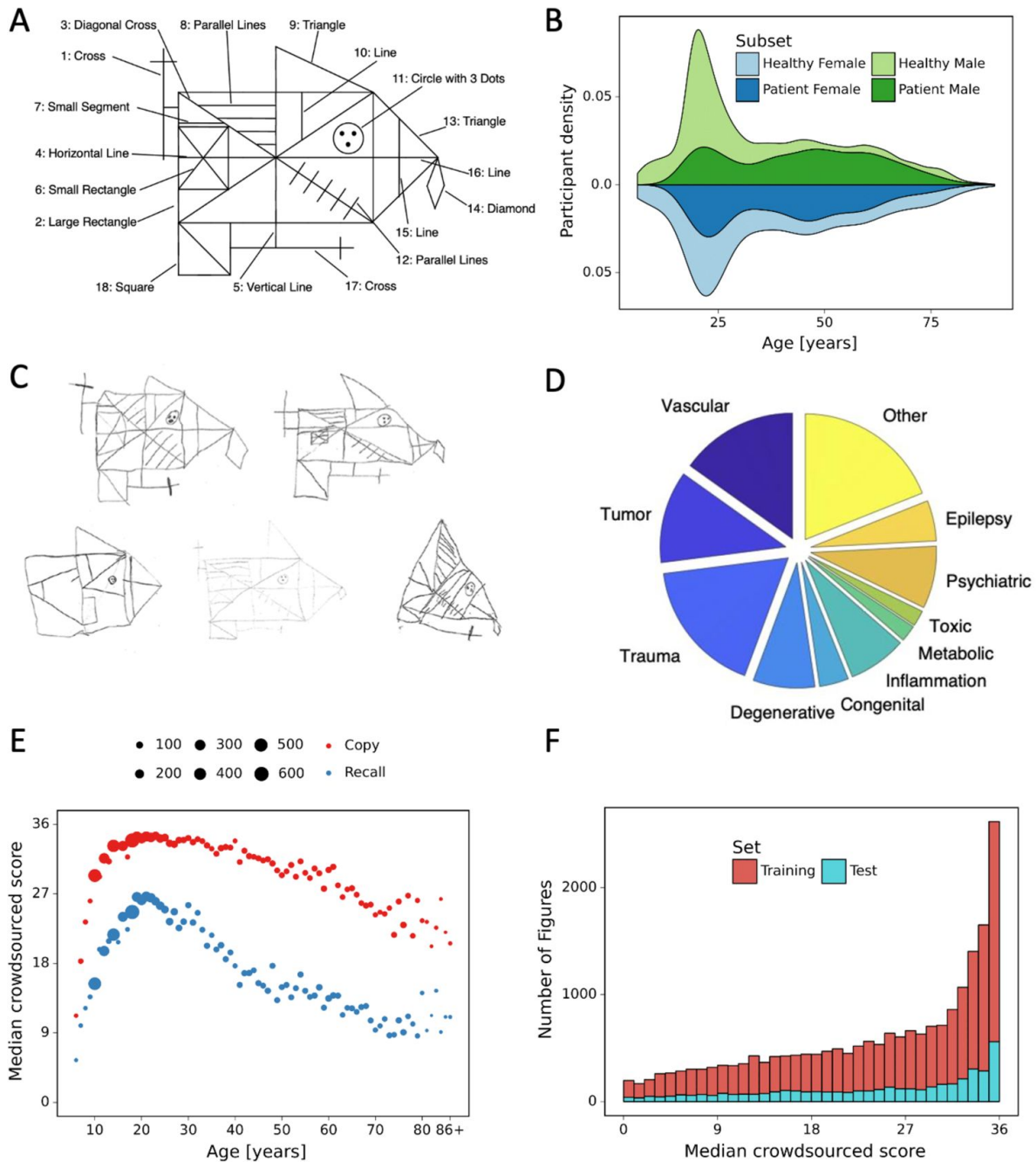


Figure 1.

A: ROCF figure with 18 elements. B: demographics of the participants and clinical population. C: examples of hand-drawn ROCF images. D: Pie chart illustrates the numerical proportion of the different clinical conditions. E: performance in the copy and (immediate) recall condition across the lifespan in the present data set. F: distribution of number of images for each total score (online raters).

Methods

Data

For our experiments, we used a data set of 20225 hand-drawn ROCF images collected from 90 different countries (see supplementary Figure S2) as well as various research and clinical settings. This large dataset spans the complete range of ROCF scores (**Figure 1F**) which allows the development of a more robust and generalizable automated scoring system. Our convergence analysis suggests that the error converges when approaching 10'000 digitized ROCFs. The data is collected from various populations regarding age and diagnostic status (healthy or with neurological and/or psychiatric disorder), shown respectively in **Figure 1E** and **Figure 1D**. The demographics of the participants and some example images are shown respectively in **Figure 1B** and **Figure 1C**. For a subset of the figures (4030), the clinician's scores were available. For each figure only one clinician rating is available. The clinicians ratings derive from six different international clinics (University Hospital Zurich, Switzerland; University Children's Hospital Zurich, Switzerland; BioCruces Health Research Institute, Spain; I.R.C.C.S. Istituto Auxologico Italiano, Ospedale San Giuseppe, Italy; Huashan Hospital, China; University Hospital Magdeburg, Germany).

The study was approved by the Institutional Ethics Review Board of the "Kantonale Ethikkommission" (BASEC-Nr. 2020-00206). All collaborators have written informed consent and/or data usage agreements for the recorded drawings from the participants. The authors assert that all procedures contributing to this work comply with the ethical standards of the relevant national and institutional committees on human experimentation and with the Helsinki Declaration of 1975, as revised in 2008. To ensure generalizability and clinical utility, the model performance was replicated in a large independent prospective validation study that was pre-registered prior to data collection (<https://osf.io/82796>). For the independent prospective validation study, an additional data set was collected and contained 2498 ROCF images from 961 healthy adults from and 288 patients with various neurological and psychiatric disorders. Further information about the participants demographics and performance in the copy and recall condition can be found in the supplementary Figure S3.

Convergence Analysis

To get an estimate on the number of ROCF needed to train our models we conducted a convergence analysis. That is, for a fixed test set with 4045 samples, we used different fractions of the remaining dataset to train the multilabel classification model, ending up with training set sizes of 1000, 2000, 3000, 4000, 6000, 8000, 12000, and 16000 samples. By evaluating the performance (i.e. MAE) of the resulting models on the fixed test set, we determined what amount of data is required for the model performance to converge. With ~3000 images, we obtain diminishing mean MAEs. After ~10000 images, no substantial improvements could be achieved by including more data, as can be seen from the progression plot in **Figure 2D**. From this, we conclude that the acquired dataset is rich and diverse enough to obtain a powerful machine learning model for the task at hand. In addition, this opens up an avenue for future research in that improvements to the model performance are likely to be achieved via algorithmic improvements, rather than via data-centric efforts.

Crowdsourced human intelligence for ground truth score

To reach high accuracy in predicting individual sample scores of the ROCFs, it is imperative that the scores of the training set are based on a systematic scheme with as little human bias as possible influencing the score. However, our analysis (see results section) and previous work

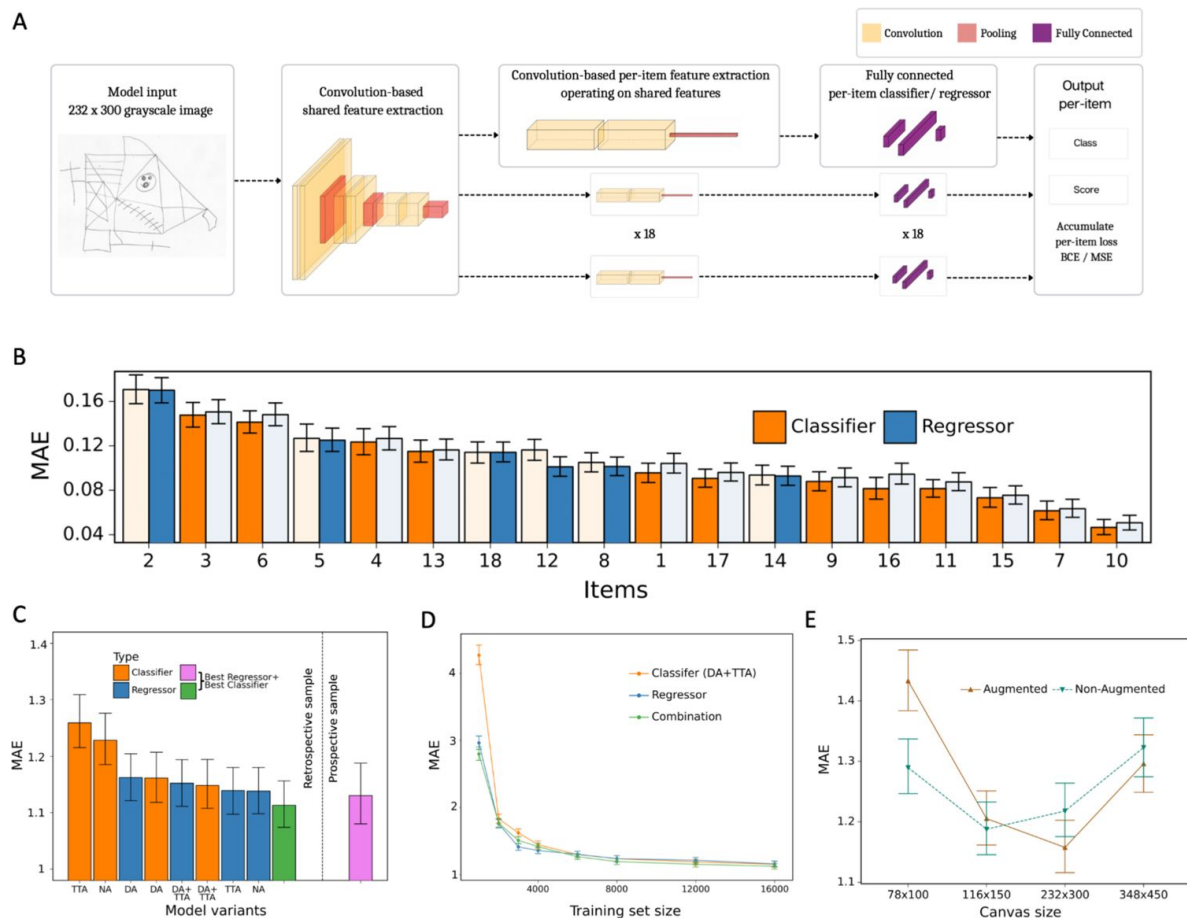


Figure 2.

A: network architecture, constituted of a shared feature extractor and 18 item-specific feature extractors and output blocks. The shared feature extractor consists of three convolutional blocks, whereas item-specific feature extractors have one convolutional block with global max-pooling. Convolutional blocks consist of two convolution and batch-normalization pairs, followed by max-pooling. Output blocks consist of two fully connected layers. ReLU activation is applied after batch normalization. After pooling, dropout is applied. B: item-specific MAE for the regression-based network (blue) and multilabel classification network (orange). In the final model, we determine whether to use the regressor or classifier network based on its performance in the validation data set, indicated by an opaque color in the bar chart. In case of identical performance, the model resulting in least variance was selected. C: Model variants were compared and the performance of the best model in the original, retrospectively collected (green) and the independent, prospectively collected (purple) test set is displayed; Clf: multilabel classification network; Reg: regression-based network; NA: no augmentation; DA: data augmentation; TTA: test time augmentation. D: Convergence analysis revealed that after ~8000 images, no substantial improvements could be achieved by adding more data. E: The effect of image size on the model performance measured in terms of MAE.

(Canham, Smith, and Tyrrell 2000) suggested that the scoring conducted by clinicians may not be consistent, because the clinicians may be unwittingly influenced by the interaction with the patient/participant or by the clinicians factor (e.g. motivation and fatigue). For this reason, we have harnessed a large pool (~5000) of human workers (crowdsourced human intelligence) who scored ROCFs, guided by our self-developed interactive web application (see supplementary Fig. S4). In this web application, the worker was first trained by guided instructions and examples. Subsequently, the worker rated each area of real ROCF drawings separately to guarantee a systematic approach. For each of the 18 elements in the figure, participants were asked to answer three binary questions: Is the item visible & recognizable? Is the item in the right place? Is the item drawn correctly? This corresponds to the original Osterrieth scoring system (Meyers, Bayless, and Meyers 1996) (see supplementary Figure S3). The final assessment consisted of answers to these questions for each of the 18 items and enabled calculation of item-wise scores (possible scores: 0, 0.5, 1, 2), which enabled to compute the total score for each image (i.e. sum over all 18 item-wise scores: range total score: 0-36).

Since participants are paid according to how many images they score, there is a danger of collecting unreliable scores when participants rush through too quickly. In order to avoid taking such assessments into account, 600 images have also been carefully scored manually by trained clinicians at the neuropsychological unit of the University Hospital Zurich (in the same interactive web application). We ensured that each crowdsourcing worker rated at least two of these 600 images, which results in 108 ratings (2 figures * 18 items * 3 questions) to compute a level of disagreement. The assessments of crowdsourcing participants that have a high level of disagreement (>20% disagreement) with this rating from clinicians are considered cheaters and are excluded from the data set. After this data cleansing, there remained an average of 13.38 crowdsource-based scores per figure. In order to use this information for training an algorithm, we require one ground truth score for each image. We assume that the scores approximately follow a normal distribution centered around the true score. With this in mind, we have obtained the ground truth for each drawing by computing the median for each item in the figure, and then summed up the medians to get the total score for the drawing in question. Taking the median has the advantage of being robust to outliers. The identical crowdsourcing approach has also used to obtain ground truth scores for the independent prospective validation study (an average of 10,36 crowdsource-based scores per figure; minimum number of ratings = 10).

Human MSE

As described in the previous section there are multiple independent scores (from different crowdsourcing workers) available for each image. It frequently happens that two people scoring the same image produce different scores. In order to quantify the disagreement among the crowdsourcing participants, we calculated the empirical standard deviation of the scores for each image. With $s_1, \dots, s_k$ referring to the k crowdsourcing scores for a given image and $\bar{s}$ to their mean, the empirical standard deviation is calculated as

$$SD_{\text{empirical}} = \sqrt{\frac{1}{k-1} \sum_{i=1}^k (s_i - \bar{s})^2}$$

The mean empirical standard deviation is 3.25. Using as ground truth the sum of item score medians, we can also calculate a *human MSE*. Assuming that the sum of median scores of all items of an image is the ground truth, the average rating performance of human raters can be estimated by computing the mean squared error (MSE) and the mean absolute error (MAE) of the human ratings by first computing the mean error for each human rater, and then computing the mean of all individual MSEs. These metrics describe how close one assessment is to this ground truth on

average and lets us make a statement on the difficulty of the task. Reusing above notation, we denote the k crowdsourcing scores for a given image by $s_1, \dots, s_k$. Let S be their median. For an image or item, we define the *human MSE* as

$$MSE_{human} = \frac{1}{k} \sum_{i=1}^k (s_i - \tilde{s})^2.$$

Similarly, we define the human MAE as

$$MAE_{human} = \frac{1}{k} \sum_{i=1}^k |s_i - \tilde{s}|.$$

Based on clinician ratings, we define both $MSE_{clinician}$ and $MAE_{clinician}$ in similar fashion. These metrics are used to compare individual crowdsourced scores, clinician scores and AI-based scores.

Convolutional Neural Network (CNN)

For the automated scoring of ROCF drawings, two variations of deep learning system were implemented: a regression-based network and a multilabel classification network. The multilabel classification problem is based on the assumption that one drawing can contain multiple classes. That is, the scores for a specific element correspond to four mutually exclusive classes (i.e. 0, 0.5, 1, 2), which corresponds to the original Osterrieth scoring system (Meyers, Bayless, and Meyers 1996) (see also supplementary Figure S1). Each class can appear simultaneously with score classes corresponding to other elements in the figure. While the scores for different elements share common features such as e.g., texture or drawing style, which are relevant for the element score, one can view the scoring of different elements as independent processes, given these shared features. These considerations are the starting point for the architectural design of our scoring system, shown in **Figure 2A**. The architecture consists of a shared feature extractor network and eighteen parallel decision heads, one for each item score. The shared feature extractor is composed of three blocks, introducing several model hyperparameters: Each block consists of a 3×3 convolutional layer, batch normalization, 3×3 convolutional layer, batch normalization and max pooling with 2×2 stride. The ReLU activation function is applied after each batch normalization and dropout is applied after max pooling. The output of each block is a representation with 64, 128 and 256 channels, respectively. Each of the eighteen per-item networks is applied on this representation. These consist of a similar feature extractor block, but with global max pooling instead, which outputs a 512-dimensional representation that is subsequently fed into a fully connected layer, that outputs a 1024-dimensional vector, followed by batch normalization, ReLU, dropout and an output fully connected layer.

Two variations of this network were implemented: a regression-based network and a multilabel classification network. They differ in the output layer and in the loss function used during training. The regression-based network outputs a single number, while the multilabel classification network outputs class probabilities, each of which corresponds to one of the four different item scores {0, 0.5, 1.0, 2.0}. Secondly, the choice of loss functions also incurs different dynamics during optimization. The MSE loss used for the regression model penalizes smaller errors less than big errors, e.g., for an item with score 0.5, predicting the score 1.0 is penalized less than when the model predicts 2.0. In other words the MSE loss naturally respects the ordering in the output space. This is in contrast to the multilabel classification model for which a cross entropy loss was used for each item, which does not differentiate between different classes in terms of “how wrong” the model is: in the example above, 1.0 is considered equally wrong as 2.0.

Data augmentation

In a further study, we investigated the effect of data augmentation during training from two perspectives. Firstly, data augmentation is a standard technique to prevent machine learning models from overfitting to training data. The intuition is that enriching the dataset with perturbed versions of the original images leads to better generalization by enriching the dataset with

additional and more diverse training set. Secondly, data augmentation can also help in making models more robust against semantic transformations like rotations or changes in perspective. These transformations are particularly relevant for the present application since users in real-life are likely to take pictures of drawings which might be slightly rotated or with a slightly tilted perspective. With these intuitions in mind, we randomly transformed drawings during training. Each transformation was a combination of Gaussian blur, a random perspective change and a rotation with angles chosen randomly between -10° and 10° .

Training & Validation

To evaluate our model, we set aside 4045 of the 20225 ROCF drawings as a testing set, corresponding to 20% of the dataset. The remaining 16180 images were used for training and validation. Specifically, we used 12944 (80%) drawings for training and 3236 (20%) as a validation set for early stopping.

The networks and training routines were implemented in PyTorch 1.8.1 (Paszke et al. 2019). The training procedure introduces additional hyperparameters: Our main model was trained on drawings of size 232×300 (see **Figure 2D** and supplementary on details on the image resolution analysis), using batches of 16 images. Model parameters were optimized by the Adam optimizer (Kingma and Ba 2014) with the default parameters $\beta_1 = 0.9$, $\beta_2 = 0.999$ and $\epsilon = 1e-8$. The initial learning rate was set to 0.01 and decayed every epoch exponentially with decay rate $\gamma = 0.95$. To prevent overfitting dropout rates were set to 0.3 and 0.5 in the convolutional and fully-connected layers, respectively. Additionally, we used the validation split (20%) of the training data for early stopping. Since our dataset is imbalanced and contains a disproportionately high number of high score drawings, we sampled the drawings in a way that results in evenly distributed scores in a single batch. We trained the model for 75 epochs and picked the weights corresponding to the epoch with the smallest validation loss. The multilabel classification model was trained with cross entropy loss applied to each item classifier independently so that the total loss is obtained by summing the individual loss terms. The regression model was trained in an analogous manner except that we used the MSE for each item score, instead of the cross entropy loss.

Performance evaluation

We evaluated the overall performance of the final model on the test set based on MSE, MAE and R-squared. Even though item-specific scores are discrete and we also implement a classification-based approach, we chose not to rely on accuracy to assess performance. Because Accuracy penalizes errors equally, minor differences in scores can lead to arbitrarily small Accuracy, or even worse scores than a model that leads to bigger differences on average. MAE gives a better picture of how accurate the model is.

Performance metrics (MSE, MAE, R-squared) were examined for the total score and additionally for each of the 18 items individually. To probe the performance of the model across a wide range of scores, we also obtained MSE and MAE for each total score. Given predictions $\hat{y}$, true scores y and the average of true scores $\bar{y}$, MSE, MAE and R-squared are defined, respectively, as

$$MSE = \frac{1}{N} \sum_i^N (y_i - \hat{y}_i)^2,$$

$$MAE = \frac{1}{N} \sum_i^N |y_i - \hat{y}_i|,$$

and

$$R^2 = 1 - \frac{\sum_i^N (y_i - \hat{y}_i)^2}{\sum_i^N (y_i - \bar{y})^2}.$$

In addition, we evaluated the performance in a pre-registered (<https://osf.io/82796>) independent prospective validation study. Importantly, both data sets (original test set (i.e. 4045 images) and the independent prospective data set (i.e. 2498 images) were never seen by the neural networks.

Test-time augmentation

Test-time augmentation (TTA) was performed to promote robustness at test-time to slight variations in input data. First, we experimented with the same augmentations that we used in our data augmentation training policy, namely randomly applying perspective change, resizing the image and applying rotations. In addition to the original image, we fed four to seven randomly transformed drawings (see description of data augmentation for the possible transformations) to the model and averaged the output logits to obtain a prediction. In addition, we experimented with approaches exclusively applying either random rotations or random perspective change. None of these random procedures improved the model performance. Nonetheless, performance improvement was achieved by applying deterministic rotations along positive and negative angles on the unit circle. We performed five rotations, uniformly distributed between -2° and 2° . For this choice, computational restrictions for the model's intended use-case were considered. Specifically, time complexity increases linearly with the number of augmentations. In keeping a small set of augmentations, we believe that small perturbations to rotation present an obvious correspondence to the variations the model will be exposed to in a real application.

Final model selection

Both model and training procedure contain various hyperparameters and even optional building blocks such as using data augmentation (DA) during training, as well as applying test-time-augmentation (TTA) during inference. After optimizing the hyperparameters of each model independently, all of the possible variants that emerge from applying DA and TTA were explored on both classification and regression models. Therefore, the space of possible model variants is spanned by the non-augmented model, the model trained with data augmentation, the model applying test-time-augmentation during inference, as well as the model variant applying both aforementioned building blocks.

To take advantage of particularities in both models, a mixture of the best performing regression and classification models was obtained. In this model, the per-item performances of both models are compared on the held-out validation set. During inference, for each item, we output the prediction of the best performing model. Therefore, on a figure-wide scale, the prediction of the combined model uses both classification and regression models concurrently, while not combining the models' per-item predictions.

Robustness Analysis

We investigated the robustness of our model against different semantic transformations which are likely to occur in real life scenarios. These were rotations, perspective changes, and changes to brightness and contrast. To assess the robustness against these transformations, we transformed images from our test set with different transformation parameters. For rotations, we rotated drawings with angles chosen randomly from increasingly higher orders, both clockwise and counterclockwise. That is, we sampled angles between 0° and 5° , between 5° and 10° and so on, up to 40° to 45° . The second transform we investigated were changes in perspective as these are also likely to occur, for example when photos of drawings are taken with a camera in a tilted position. The degree of perspective change was guided by a parameter between 0 and 1, where 0 corresponds to the original image and 1 corresponds to an extreme change in perspective, making the image almost completely unrecognizable. Furthermore, due to changes in light conditions, an image might appear brighter or with a different contrast. We used brightness parameters between

0.1 and 2.0, where values below 1.0 correspond to a darkening of the image and values above 1.0 correspond to a brighter image. Similarly, for contrast, we varied the contrast parameter between 0.1 and 2.0, with values above (below) 1.0 corresponding to high (low) contrast images.

Results

A human MSE and clinicians' scoring

We have harnessed a large pool (~5000) of human workers (crowdsourced human intelligence) to obtain ground truth scores and compute the *human MSE*. An average of 13.38 (sd = 2.23) crowdsourced-based scores per figure were gathered. The *average human MSE* over all images is 16.3, and the *average human MAE* is 2.41.

For a subset (4030) of the 20225 ROCF images, we had access to scores conducted by different professional clinicians. This enabled us to analyze and compare the scoring of professionals to the crowdsourcing results. The clinician MSE over all images is 9.25 and the clinician MAE is 2.15, indicating a better performance of the clinicians compared to the average human rater.

Machine-based scoring system

The multilabel classification network achieved a total score MSE of 3.56 and MAE of 1.22, which is already considerably lower than the corresponding human performance. Implementing additional data augmentation has led to a further improvement in accuracy with an MSE of 3.37 and MAE of 1.16. Furthermore, when combining data augmentation with test time augmentation leads to a further improvement of performance with an MSE of 3.32 and MAE of 1.15.

The regression based network variant led to a further slight improvement in performance, reducing the total score MSE to 3.07 and the MAE to 1.14. Finally, our best model results from combining the regression model with the multilabel classification model in the following way: For each item of the figure we determine whether to use the regressor or classifier based on its performance on the validation set (**Figure 2B**). Aggregating the two models in this way, leads to a MSE of 3.00 and a MAE of 1.11. **Figure 2C** presents the error for all combinations of data and test time augmentation with the multilabel classification and regression model separately and in combination. During the experiments it became apparent that for the multilabel classification network applying both data augmentation as well as test-time-augmentation jointly improves the model's performance (see **Figure 2C**). Somewhat surprisingly, applying these elements to the multihead regression model did not improve the performance compared to the non-augmented version of the model. The exact performance metrics (MSE, MAE and R-squared) of all model variants are reported in the supplementary table S1 and for each figure element in table S2. In addition, the model performance was replicated in the independent prospective validation study (i.e. MSE = 3.32; MAE = 1.13). The performance metrics of each figure element for the independent prospective validation study are reported in the supplementary table S4 and S5.

We further conducted a more fine-grained analysis of our best model. **Figure 3A** shows the score for each figure in the data set contrasted with the ground truth score (i.e. median of online raters). In addition, we computed the difference between ground truth score and predicted score revealing that errors made by our model are concentrated closely around 0 (**Figure 3B**), while the distribution of clinician's errors is much more spread out (**Figure 3D** and **3E**). It is interesting to note that, like the clinicians, our model exhibits a slight bias towards higher scores, although much less pronounced. Importantly, the model does not demonstrate any considerable bias towards specific figure elements. In contrast to the clinicians, the MAE is very balanced across each individual item of the figure (**Figure 3C** and **3F**, and supplementary table S2). Finally, a detailed breakdown of the expected performance across the entire range of total scores is displayed for the model (**Figure 3G** and supplementary table S3), the clinicians (3H) and

average online raters (3I). As can be seen, the model exhibits a comparable MAE for the entire range of total scores, although there is a trend that high scores exhibit lower MAEs. These results were confirmed in the independent prospective validation study (see **Figure 3G** [↗](#)-**3I** [↗](#), supplementary Figure S5 and supplementary table S5).

Robustness Analysis

Using data augmentation did not critically improve the accuracy, which is likely due to the fact that our dataset is already large and diverse enough. However, data augmentation does significantly improve robustness against semantic transformations, as can be seen from supplementary Figure S6. In particular, using our data augmentation pipeline, the model becomes much more robust against rotations and changes in perspective. On the other hand, for changes to brightness and contrast we could not find a clear trend. This is however not surprising as the data augmentation pipeline does not explicitly encourage the model to be more robust against these transformations. Overall, we demonstrate that our scoring system is highly robust for realistic changes in rotations, perspective, brightness, and contrast of the images (**Figure 4** [↗](#)). Performance is degraded only under unrealistic and extreme transformations.

Discussion

In this study, we developed an AI-based scoring system for a nonverbal visuo-spatial memory test that is being utilized in clinics around the world on a daily basis. For this, we trained two variations of deep learning systems and used a rich data set of more than 20k hand-drawn ROCF images covering the entire life span, different research and clinical environments as well as representing global diversity. By leveraging human crowdsourcing we obtained unbiased high-precision training labels. Our best model results from combining a multihead convolutional neural network for regression with a network for multilabel classification. Data and test time augmentation were used to improve the accuracy and made our model more robust against geometric transformations like rotations and changes in perspective.

Overall, our results provide evidence that our AI-scoring tool outperforms both amateur raters and clinicians. Our model performed highly unbiased as it yielded predictions very close to the ground truth and the error was similarly distributed around zero. Importantly, these results have been replicated in an independent prospectively collected data set. The scoring system reliably identified and accurately scored individual figure elements in previously unseen ROCF drawings, which facilitates explainability of the AI-scoring system. Comparable performance for each figure element indicates no considerable bias towards specific figure compositions. Furthermore, the error of our scoring system is rather balanced across the entire range of total scores. An exception are the highest scored images displaying the smallest error, which has two explanations: First, compared to low score images, for which the same score can be achieved by numerous combinations of figure elements, the high score images by nature do not have as many possible combinations. Secondly, the data set contains a proportional larger amount of available training data for the high score images.

While the ROCF is one of the most commonly employed neuropsychological test to evaluate nonverbal visuo-spatial memory in clinical setting (Shin et al. 2006), the current manual scoring in clinics is time-consuming (5 to 15 minutes for each drawing), requires training, and ultimately depends on subjective judgements, which inevitably introduces human scoring biases (Watkins 2017; Franzen 2000) and low inter-rater reliability (Franzen 2000; Huygelier et al. 2020) (see also supplementary Figure S7). Thus, an objective and reliable automated scoring system is in great demand as it can overcome the problem of intra- and inter-rater variability. More importantly, such an automated system can remove a time-consuming, tedious and repetitive task from

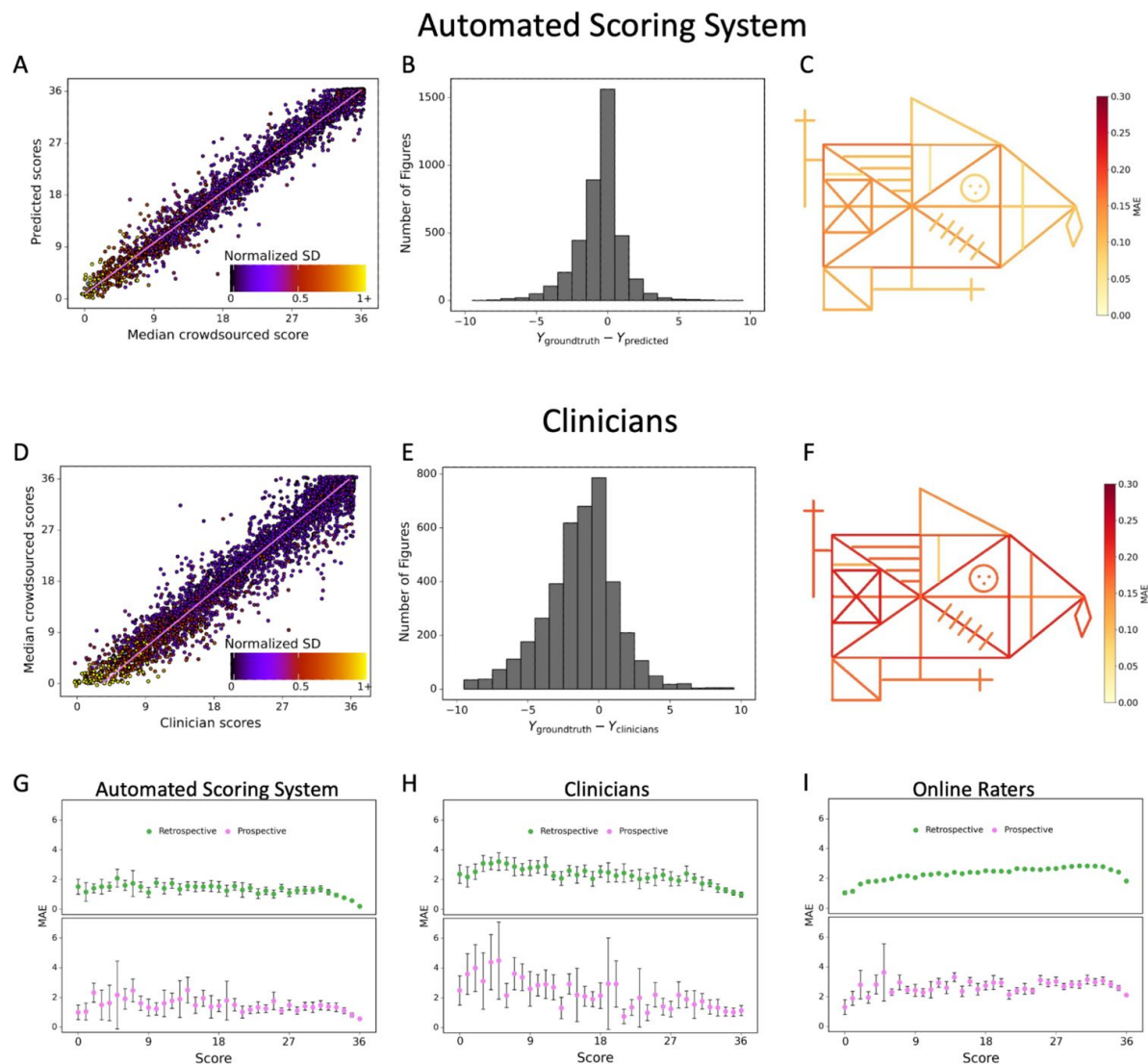


Figure 3.

Contrasting the ratings of our model (A) and clinicians (D) against the ground truth revealed a larger deviation from the regression line for the clinicians. A jitter is applied to better highlight the dot density. The distribution of errors for our model (B) and the clinicians ratings (E) are displayed. The MAE of our model (C) and the clinicians (F) is displayed for each individual item of the figure (see also supplementary table S2). The corresponding plots for the performance on the prospectively collected data is displayed in the supplementary Figure S5. The model performance for the retrospective (green) and prospective (purple) sample across the entire range of total scores for model (G), clinicians (H) and online raters (I) are presented.

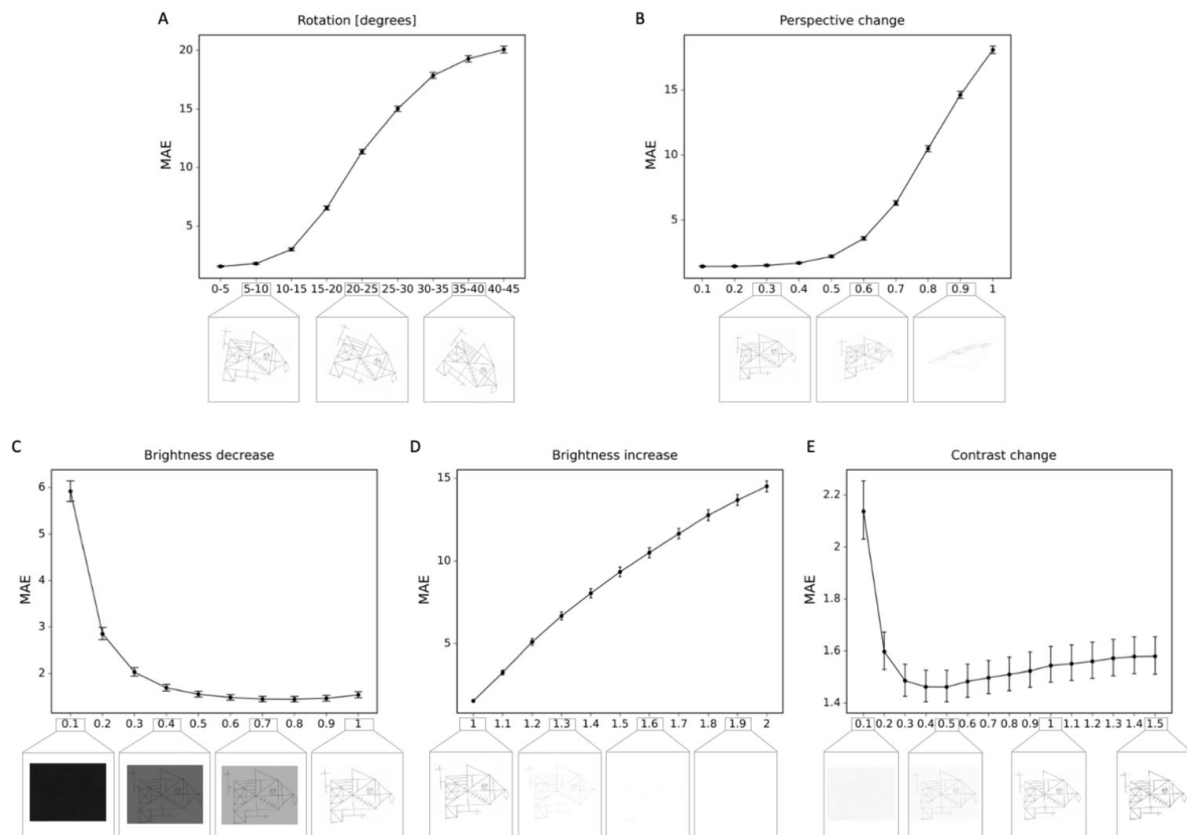


Figure 4.

Robustness to geometric, brightness and contrast variations. The MAE is depicted for different degrees of transformations. In addition examples of the transformed ROCF draw are provided.

clinicians and thereby helps to reduce workload of clinicians and/or allowing more time for more patient-centered care. Overall, automation can support clinicians to make healthcare decisions more accurate and timely.

Over the past years, deep learning has made significant headway in achieving above human-level performance on well-specified and isolated tasks in computer vision. However, unleashing the power of deep learning depends on the availability of large training sets with accurate training labels. We obtained over 20k hand-drawn ROCF images and invested immense time and labor to manually scan the ROCF drawings. This effort was only possible by orchestrating across multiple clinical sites and having funds available to hire the work force to accomplish the digitization. Our approach highlights the need to intensify national and international effort of the digitalization of health record data. Although electronic health records (EHR) are increasingly available (e.g. quadrupled in the US from 2007 to 2012 (Hsiao, Hing, and Ashman 2014)), challenges remain in terms of handling, processing and moving such big data. Improved data accessibility and better consensus in organization and sharing could simplify and foster similar approaches in neuropsychology and medicine.

Another important difficulty is that data typically comes from various sources and thus exhibits large heterogeneity in terms of quality, size, format of the images, and crucially the quality of the labeled dataset. High-quality labeled and annotated datasets are the foundation of a successful computer vision system. A novelty of our approach is that we have trained a large pool of human internet workers (crowdsourced human intelligence) to score ROCFs drawings guided by our self-developed interactive web application. Each element of the figure was scored by several human workers (13 workers on average per figure). To derive ground truth scores, we took the median score, as it has the advantage of being robust to outliers. To further ensure high-quality data annotation, we identified and excluded crowdsourcing participants that have a high level of disagreement (>20% disagreement) with this rating from trained clinicians, who carefully scored manually a subset of the data in the same interactive web application. Importantly our two-fold innovative approach that combines the digitization of neuropsychological tests and the high-quality scoring using crowdsourced intelligence can provide a roadmap for how AI can be leveraged in neuropsychological testing more generally as it can be easily adapted and applied to various other neuropsychological tests (e.g. Clock Drawing Test (Freedman et al. 1994), Taylor Complex Figure Test (Awad et al. 2004), Hamasch 5-point test (Regard, Strauss, and Knapp 1982)).

Another important prerequisite of using AI for the automated scoring of neuropsychological tests is the availability of training data that is sufficiently diverse and obtains sufficient examples from all possible scores. Recent examples in computer vision have demonstrated that insufficient diversity in training data can lead to biases and adverse consequences (Buolamwini and Gebre 23-Feb 2018; Ensign et al. 2017). Given that our training data set included data from children and a representative sample of patients (see **Figure 1B** [↗](#)), exhibiting a range of neurological disorders, drawings were frequently extremely dissimilar to the target figure. Importantly, our scoring system delivers accurate scores even in cases where drawings are distorted (e.g. caused by fine motor impairments) or omissions (e.g. due to visuospatial deficits), which are typical impairment patterns in neurological patients. In addition, our scoring system is robust against different semantic transformations (e.g. rotations, perspective change, brightness) which are likely to occur in real life scenarios, when the examiner takes a photo of the ROCF drawing, which naturally will lead to varying viewpoints and illumination conditions. Thus, this robustness is a pivotal prerequisite for any potential clinical utility.

To improve usability of our system and guarantee clinical utility, we have integrated our model into a tablet-(and smartphone-) based application, in which users can take a photo (or upload an image) of an ROCF drawing and the application instantly provides a total score. Importantly, the automated scoring system also maintains explainability by providing visualization of detailed score breakdowns (i.e. item-specific scores), which is a highly valued property of AI-assisted

medical decision making and key to help interpret the score and communicate them to the patient (see supplementary Figure S8). The scoring is completely reproducible and thus facilitates valid large-scale comparisons of ROCF data, which enable population-based cognitive screening and (assisted) self-assessments.

In summary, we have created a highly robust and automated AI-based scoring tool, which provides unbiased, reproducible and immediate scores for the ROCF test in research and clinical environments. Our current machine-based automated quantitative scoring system outperforms both amateur raters and clinicians.

The present findings demonstrate that automated scoring systems can be developed to advance the quality of neuropsychological assessments by diminishing reliance on subjective ratings and efficiently improving scoring in terms of time and costs. Our innovative approach can be translated to many different neuropsychological tests, thus providing a roadmap for paving the way to AI-guided neuropsychology.

Financial Support

This work was supported by the URPP “Dynamics of Healthy Aging” and BRIDGE [40B2-0_187132], which is a joint programme of the Swiss National Science Foundation SNSF and Innosuisse. Furthermore, BHV is funded by the Swiss National Science Foundation [10001C_197480]. Finally, CL is supported by the Swiss State Secretariat for Education, Research and Innovation (SERI) under contract number MB22.00036. None of the authors have been paid to write this article by a pharmaceutical company or other agency. Authors were not precluded from accessing data in the study, and they accept responsibility to submit for publication.

Declaration of Interests

The authors declare no conflict of interest.

Data sharing

The clinical dataset can not be shared publicly because the consent for data sharing was not obtained from study participants. The Prolific dataset is available on request after signing a data use agreement. All preprocessing and analysis scripts used for the analyses are uploaded on github (<https://github.com/methlabUZH/rey-figure> .

Authorship Contribution

Nicolas Langer: Conceptualization, Methodology, Software, Formal analysis, Writing - original draft, Writing - review & editing, Visualization, Supervision, Project administration, Resources, Funding acquisition.

Maurice Weber: Conceptualization, Methodology, Software, Formal analysis, Validation, Data curation, Writing - original draft, Writing - review & editing, Visualization.

Bruno Hebling Vieira: Conceptualization, Methodology, Software, Formal analysis, Validation, Data curation, Writing - original draft, Writing - review & editing, Visualization.

Dawid Strzelcy: Conceptualization, Methodology, Software, Formal analysis, Data curation, Validation, Resources, Writing - original draft, Writing - review & editing, Visualization.

Lukas Wolf: Conceptualization, Methodology, Software, Formal analysis, Data curation, Validation, Writing - original draft, Writing - review & editing, Visualization.

Andreas Pedroni: Conceptualization, Methodology, Writing - review & editing

Jonathan Heitz: Conceptualization, Methodology, Writing - review & editing

Stephan Müller: Conceptualization, Methodology, Writing - review & editing

Christoph Schultheiss: Conceptualization, Methodology, Writing - review & editing

Marius Tröndle: Conceptualization, Methodology, Data curation, Resources, Writing - review & editing

Juan Carlos Arango-Lasprilla: Data curation, Resources, Writing - review & editing

Diego Rivera: Data curation, Resources, Writing - review & editing

Federica Scarpina: Data curation, Resources, Writing - review & editing

Qianhua Zhao: Data curation, Resources, Writing - review & editing

Rico Leuthold: Conceptualization, Methodology, Software

Flavia Wehrle: Data curation, Resources, Writing - review & editing

Oskar G. Jenni: Data curation, Resources, Writing - review & editing

Peter Brugger: Data curation, Resources, Writing - review & editing

Tino Zaehle: Data curation, Resources, Writing - review & editing

Romy Lorenz: Conceptualization, Methodology, Writing - original draft, Writing - review & editing

Ce Zhang: Conceptualization, Methodology, Formal analysis, Writing - review & editing, Supervision, Resources, Funding acquisition.

References

1. Alladi Suvarna, Arnold Robert, Mitchell Joanna, Nestor Peter J., Hodges John R. (2006) **Mild Cognitive Impairment: Applicability of Research Criteria in a Memory Clinic and Characterization of Cognitive Profile** *Psychological Medicine* <https://doi.org/10.1017/S0033291705006744>
2. Awad Nesrine, Tsiakas Maria, Gagnon Michèle, Mertens Valérie B., Hill Eva, Messier Claude (2004) **Explicit and Objective Scoring Criteria for the Taylor Complex Figure Test** *Journal of Clinical and Experimental Neuropsychology* **26**:405–15

3. Berry David T. R., Allen Rebecca S., Schmitt Frederick A. (1991) **Rey-Osterrieth Complex Figure: Psychometric Characteristics in a Geriatric Sample** *Clinical Neuropsychologist* <https://doi.org/10.1080/13854049108403298>
4. Buolamwini Joy, Gebru Timnit (2018) **Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification** *Proceedings of Machine Learning Research*
5. Canham R. O., Smith S. L., Tyrrell A. M. (2000) **"Automated Scoring of a Neuropsychological Test: The Rey Osterrieth Complex Figure."** *Proceedings of the 26th Euromicro Conference. EUROMICRO 2000 Informatics: Inventing the Future* <https://doi.org/10.1109/eurmic.2000.874519>
6. Ensign Danielle, Friedler Sorelle A., Neville Scott, Scheidegger Carlos, Venkatasubramanian Suresh (2017) **Runaway Feedback Loops in Predictive Policing** *arXiv* <https://doi.org/10.48550/arXiv.1706.09847>
7. Fastenau P. S., Bennett J. M., Denburg N. L. (1996) **Application of Psychometric Standards to Scoring System Evaluation: Is 'New' Necessarily 'Improved'?** *Journal of Clinical and Experimental Neuropsychology* **18**:462–72
8. Franzen Michael D (2000) **Validity as Applied to Neuropsychological Assessment** *Reliability and Validity in Neuropsychological Assessment* https://doi.org/10.1007/978-1-4757-3224-5_5
9. Freedman Morris, Leach Larry, Kaplan Edith, Winocur Gordon, Shulman Kenneth, Delis Dean C. (1994) **Clock Drawing: A Neuropsychological Analysis**
10. Groth-Marnat Gary (2000) **Neuropsychological Assessment in Clinical Practice: A Guide to Test Interpretation and Integration** *Wiley*
11. Harbi Zainab, Hicks Yulia, Setchi Rossitza (2016) **Clock Drawing Test Digit Recognition Using Static and Dynamic Features** *Procedia Computer Science* <https://doi.org/10.1016/j.procs.2016.08.166>
12. Hsiao Chun-Ju, Hing Esther, Ashman Jill (2014) **"Trends in Electronic Health Record System Use among Office-Based Physicians: United States, 2007-2012."** *National Health Statistics Reports* **75**:1–18
13. Huygelier Hanne, Moore Margaret Jane, Demeyere Nele, Gillebert Céline R. (2020) **Non-Spatial Impairments Affect False-Positive Neglect Diagnosis Based on Cancellation Tasks** *Journal of the International Neuropsychological Society: JINS* **26**:668–78
14. Kim Hyungsin, Cho Young Suk, Do Ellen Yi-Luen (2010) **"Computational Clock Drawing Analysis for Cognitive Impairment Screening."** *Proceedings of the Fifth International Conference on Tangible Embedded, and Embodied Interaction* <https://doi.org/10.1145/1935701.1935768>
15. Kingma Diederik P., Ba Jimmy (2014) **Adam: A Method for Stochastic Optimization** *December* <https://doi.org/10.48550/arXiv.1412.6980>
16. Levine Andrew J., Miller Eric N., Becker James T., Selnes Ola A., Cohen Bruce A. (2004) **Normative Data for Determining Significance of Test-Retest Differences on Eight Common Neuropsychological Instruments** *The Clinical Neuropsychologist* <https://doi.org/10.1080/1385404049052420>

17. Meyers J. E., Bayless J. D., Meyers K. R. (1996) **Rey Complex Figure: Memory Error Patterns and Functional Abilities** *Applied Neuropsychology* **3**:89–92
18. Meyers John E., Volbrecht Marie (1999) **Detection of Malingerers Using the Rey Complex Figure and Recognition Trial** *Applied Neuropsychology* https://doi.org/10.1207/s15324826an0604_2
19. Nazar Haris Bin, Moetesum Momina, Ehsan Shoaib, Siddiqi Imran, Khurshid Khurram, Vincent Nicole, McDonald-Maier Klaus D. (2017) **Classification of Graphomotor Impressions Using Convolutional Neural Networks: An Application to Automated Neuro-Psychological Screening Tests**. 2017 14th IAPR International Conference on Document Analysis and Recognition (ICDAR) <https://doi.org/10.1109/icdar.2017.78>
20. Osterrieth Paul Alexandre (1944) **Le test de copie d'une figure complexe: Contribution à l'étude de la perception et de la mémoire** *Delachaux et Niestlé*
21. Paszke Adam, Gross Sam, Massa Francisco, Lerer Adam, Bradbury James, Chanan Gregory, Killeen Trevor, et al. (2019) **PyTorch: An Imperative Style, High-Performance Deep Learning Library** *In Proceedings of the 33rd International Conference on Neural Information Processing Systems* <https://doi.org/10.48550/arXiv.1912.01703>
22. Petilli Marco A., Daini Roberta, Saibene Francesca Lea, Rabuffetti Marco (2021) **Automated Scoring for a Tablet-Based Rey Figure Copy Task Differentiates Constructional, Organisational, and Motor Abilities** *Scientific Reports* **11**
23. Rabin Laura A., Barr William B., Burton Leslie A. (2005) **Assessment Practices of Clinical Neuropsychologists in the United States and Canada: A Survey of INS, NAN, and APA Division 40 Members** *Archives of Clinical Neuropsychology: The Official Journal of the National Academy of Neuropsychologists* **20**:33–65
24. Rabin Laura A., Paolillo Emily, Barr William B. (2016) **Stability in Test-Usage Practices of Clinical Neuropsychologists in the United States and Canada Over a 10-Year Period: A Follow-Up Survey of INS and NAN Members** *Archives of Clinical Neuropsychology: The Official Journal of the National Academy of Neuropsychologists* **31**:206–30
25. Regard Marianne, Strauss Esther, Knapp Paul (1982) **Children's Production on Verbal and Non-Verbal Fluency Tasks** *Perceptual and Motor Skills* <https://doi.org/10.2466/pms.1982.55.3.839>
26. Shin Min-Sup, Park Sun-Young, Park Se-Ran, Seol Soon-Ho, Kwon Jun Soo (2006) **Clinical and Empirical Applications of the Rey-Osterrieth Complex Figure Test** *Nature Protocols* <https://doi.org/10.1038/nprot.2006.115>
27. Somerville J., Tremont G., Stern R. A. (2000) **The Boston Qualitative Scoring System as a Measure of Executive Functioning in Rey-Osterrieth Complex Figure Performance** *Journal of Clinical and Experimental Neuropsychology* **22**:613–21
28. Trautmann Sebastian, Rehm Jürgen, Hans-ulrich Wittchen (2016) **"The Economic Costs of Mental Disorders."** *EMBO Reports* <https://doi.org/10.15252/embr.201642951>
29. Trojano Luigi, Gainotti Guido (2016) **Drawing Disorders in Alzheimer's Disease and Other Forms of Dementia** *Journal of Alzheimer's Disease* <https://doi.org/10.3233/jad-160009>

30. Vogt J., Kloosterman H., Vermeent S., Van Elswijk G., Dotsch R., Schmand B. (2019) **Automated Scoring of the Rey-Osterrieth Complex Figure Test Using a Deep-Learning Algorithm** *Archives of Clinical Neuropsychology* <https://doi.org/10.1093/arclin/acz035.04>
31. Watkins Marley W (2017) **The Reliability of Multidimensional Neuropsychological Measures: From Alpha to Omega** *The Clinical Neuropsychologist* **31**:1113–26
32. Webb Sam S., Moore Margaret Jane, Yamshchikova Anna, Kozik Valeska, Duta Mihaela D., Voiculescu Irina, Demeyere Nele (2021) **Validation of an Automated Scoring Program for a Digital Complex Figure Copy Task within Healthy Aging and Stroke** *Neuropsychology* **35**:847–62

Editors

Reviewing Editor

Juan Zhou

National University of Singapore, Singapore, Singapore

Senior Editor

Yanchao Bi

Beijing Normal University, Beijing, China

Reviewer #1 (Public Review):

Summary:

The authors aimed to develop and validate an automated, deep learning-based system for scoring the Rey-Osterrieth Complex Figure Test (ROCF), a widely used tool in neuropsychology for assessing memory deficits. Their goal was to overcome the limitations of manual scoring, such as subjectivity and time consumption, by creating a model that provides automatic, accurate, objective, and efficient assessments of memory deterioration in individuals with various neurological and psychiatric conditions.

Strengths:

Comprehensive Data Collection:

The authors collected over 20,000 hand-drawn ROCF images from a wide demographic and geographic range, ensuring a robust and diverse dataset. This extensive data collection is critical for training a generalizable and effective deep learning model.

Advanced Deep Learning Approach:

Utilizing a multi-head convolutional neural network to automate ROCF scoring represents a sophisticated application of current AI technologies. This approach allows for detailed analysis of individual figure elements, potentially increasing the accuracy and reliability of assessments.

Validation and Performance Assessment:

The model's performance was rigorously evaluated against crowdsourced human intelligence and professional clinician scores, demonstrating its ability to outperform both groups. The inclusion of an independent prospective validation study further strengthens the credibility of the results.

Robustness Analysis Efficacy:

The model underwent a thorough robustness analysis, testing its adaptability to variations in rotation, perspective, brightness, and contrast. Such meticulous examination ensures the

model's consistent performance across different clinical imaging scenarios, significantly bolstering its utility for real-world applications.

Weaknesses:

Insufficient Network Analysis for Explainability:

The paper does not sufficiently delve into network analysis to determine whether the model's predictions are based on accurately identifying and matching the 18 items of the ROCF or if they rely on global, item-irrelevant features. This gap in analysis limits our understanding of the model's decision-making process and its clinical relevance.

Generative Model Consideration:

The critique suggests exploring generative models to model the joint distribution of images and scores, which could offer deeper insights into the relationship between scores and specific visual-spatial disabilities. The absence of this consideration in the study is seen as a missed opportunity to enhance the model's explainability and clinical utility.

Appraisal and discussion:

By leveraging a comprehensive dataset and employing advanced deep learning techniques, they demonstrated the model's ability to outperform both crowdsourced raters and professional clinicians in scoring the ROCF. This achievement represents a significant step forward in automating neuropsychological assessments, potentially revolutionizing how memory deficits are evaluated in clinical settings. Furthermore, the application of deep learning to clinical neuropsychology opens avenues for future research, including the potential automation of other neuropsychological tests and the integration of AI tools into clinical practice. The success of this project may encourage further exploration into how AI can be leveraged to improve diagnostic accuracy and efficiency in healthcare.

However, the critique regarding the lack of detailed analysis across different patient demographics, the inadequacy of network explainability, and concerns about the selection of median crowdsourced scores as ground truth raises questions about the completeness of their objectives. These aspects suggest that while the aims were achieved to a considerable extent, there are areas of improvement that could make the results more robust and the conclusions stronger.

<https://doi.org/10.7554/eLife.96017.1.sa2>

Reviewer #2 (Public Review):

Summary:

The authors aimed to develop and validate a machine-learning-driven neural network capable of automatic scoring of the Rey-Osterrieth Complex Figure. They aimed to further assess the robustness of the model to various parameters such as tilt and perspective shift in real drawings. The authors leveraged the use of a huge sample of lay workers in scoring figures and also a large sample of trained clinicians to score a subsample of figures. Overall, the authors found their model to have exceptional accuracy and perform similarly to crowdsourced workers and clinicians with, in some cases, less degree of error/score dispersion than clinicians.

Strengths:

The authors used very large data; including a large number of Rey-Osterrieth Complex Figures, a huge crowdsourced human worker sample, and a large clinician sample.

The authors deeply describe their model in relatively accessible terms.

The writing style of the paper is accessible, scientific, and thorough.

Pre-registration of the prospectively collected new data was acceptable.

Weaknesses:

There is no detail on how the final scoring app can be accessed and whether it is medical device-regulated.

No discussion on the difference in sample sizes between the pre-registration of the prospective study and the results (e.g., aimed for 500 neurological patients but reported data from 288).

Details in pre-registration and paper regarding samples obtained in the prospective study were lacking.

Demographics for the assessment of the representation of healthy and non-healthy participants were not present.

The authors achieved their aims and their results and conclusions are supported by strong methods and analyses. The resulting app produced in this work, if suitable for clinical practice, will have impact in automated scoring, which many clinicians will be exceptionally happy with.

<https://doi.org/10.7554/eLife.96017.1.sa1>

Reviewer #3 (Public Review):

Summary:

This study presented a valuable inventory of scoring a neuropsychological test, ROCFT, with constructing an artificial intelligence model.

Strengths:

They constructed huge samples collected among multi-center international researchers and tested the model precisely using internet data.

The model scored the test with excellent ability, surpassing even experts. The product can run an application on a tablet, which helps clinicians and patients.

Their method of building the model of deep learning and testing will apply to tests in all fields, not just the psychological field.

Weaknesses:

The considerable effort and cost to make the model only for an existing neuropsychological test.

<https://doi.org/10.7554/eLife.96017.1.sa0>