

ABFR-KAN: Kolmogorov-Arnold Networks for Functional Brain Analysis

Tyler Ward¹, Abdullah Imran¹

Department of Computer Science, University of Kentucky, Lexington, KY 40506, USA

Abstract

Functional connectivity (FC) analysis is widely used for computer-aided brain disorder diagnosis, but conventional atlas-based approaches can introduce structural bias and overlook subject-specific brain function. Addressing this, we propose **Advanced Brain Function Representation with Kolmogorov-Arnold Networks (ABFR-KAN)**, a transformer-based framework that combines globally shared randomized anchor patches with subject-specific multi-scale patch sampling and Kolmogorov-Arnold Networks (KANs). The randomized anchors reduce structural bias and improve anatomical conformity while providing a consistent reference space across subjects. In contrast, per-subject patch sampling captures individual functional variability and improves the reliability of FC estimation. We evaluate the proposed method on resting-state functional magnetic resonance imaging (rs-fMRI) data from the Autism Brain Imaging Data Exchange (ABIDE) I and II datasets (281 subjects and 211 subjects, respectively). Extensive experiments, including 5-fold cross-validation, cross-site evaluation, and ablation studies across varying model backbones and KAN configurations, demonstrate that ABFR-KAN consistently outperforms state-of-the-art baselines for autism spectrum disorder (ASD) classification. These results highlight the potential of combining advanced brain function representation strategies with KAN-based architectures for robust neuroimaging analysis. Our code is available at <https://github.com/tbwa233/ABFR-KAN>.

Keywords

Brain MRI, Classification, Functional Connectivity, Kolmogorov-Arnold Network, Transformer, ASD

Article informations

<https://doi.org/10.59275/j.melba.2024-AAAA>

©2026 Ward and Imran. License: CC-BY 4.0

Volume 2026, Received: 2025-12-29, Published yyyy-m2-d2

Corresponding author: aimran@uky.edu

Special issue: Medical Imaging with Deep Learning (MIDL) 2025

Guest editors: Lisa Koch, Ronald M. Summers, Chen Chen, Yan Zhuang



1. Introduction

ACCURATE and early diagnosis of brain disorders is crucial for effective medical intervention and treatment planning [Kumar \(2025\)](#). A widely adopted method for aiding diagnosis of such conditions is the analysis of functional connectivity (FC), which measures relationships between brain regions using blood-oxygen-level-dependent (BOLD) signals captured during resting-state functional magnetic resonance imaging (rs-fMRI) [Du et al. \(2018\)](#). However, FC analysis (FCA) often relies on pre-defined atlases to define regions-of-interest (ROIs) within the brain, which carry several limitations. This approach can lead to subjective selection bias, disregard for individual specificity, and a lack of interaction between brain regions and FCA [Liu et al. \(2024a\)](#). Despite research into various methods of addressing these issues, such as data-driven

[Reeves et al. \(2025\)](#), individualized [Hakonen et al. \(2025\)](#), and multi-atlas [Han et al. \(2025\)](#) setups, a definitive resolution to all of the challenges associated with atlas-based parcellation techniques has not yet emerged.

Additional drawbacks of traditional FCA methods are the high dimensionality and complexity of the functional representations, and finding solutions addressing these issues is an active area of research [Wang et al. \(2018\)](#); [Vidaurre \(2021\)](#). In the age of artificial intelligence (AI), many researchers have explored FCA methods that incorporate deep learning (DL) [Qu et al. \(2021\)](#), such as multi-layer perceptrons (MLPs) that take a flattened FC matrix as an input [Hou et al. \(2025\)](#). MLP-based methods have demonstrated effectiveness in downstream tasks such as brain disorder classification [Sachdeva et al. \(2024\)](#), but they share similar drawbacks as non-AI-based methods for FCA.

For example, a connectome for N brain regions is of-

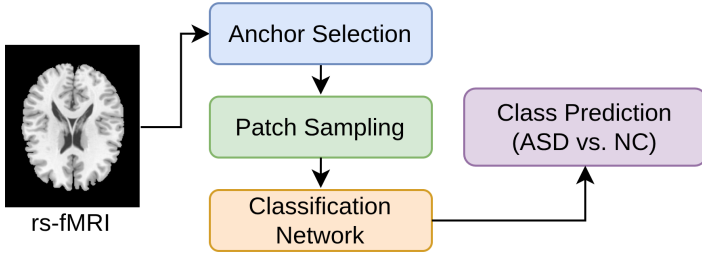


Figure 1: The general workflow of our proposed ABFR-KAN framework. In the anchor selection stage, a fixed set of reference patches is randomly sampled from grey matter (GM) regions to serve as anatomically unbiased functional anchors. In the patch sampling stage, multiple sets of brain patches are sampled at different spatial scales from each subject, and their FC to the anchors is computed to form robust, position-aware functional representations. These representations are then embedded with spatial information and processed by the classification network to produce a binary prediction of whether a subject belongs to the autism spectrum disorder (ASD) or normal control (NC) cohort.

ten represented by an $N \times N$ matrix or an $\frac{N(N-1)}{2}$ -length feature vector for unique pairs [Kaiser \(2011\)](#). This means that for a moderately sized parcellation (e.g., $N = 100$), there could be upwards of 10,000 features per subject. This makes MLPs trained on data such as this highly prone to overfitting due to the curse of dimensionality, as the extremely high-dimensional feature space requires estimating a vast number of parameters from relatively limited samples [Smucny et al. \(2022\)](#). In addition to this risk of overfitting, it is computationally burdensome to use MLPs for FCA, as MLPs that take a full FC matrix as input must process all pairwise connections simultaneously, which is inefficient.

Recently, Kolmogorov-Arnold Networks (KANs) [Liu et al. \(2024c\)](#) have emerged as an alternative to traditional MLPs, leveraging learnable activation functions on edges rather than fixed activation functions on nodes. Inspired by the Kolmogorov-Arnold representation theorem, KANs replace conventional weight matrices with univariate functions parameterized as splines, offering improved expressiveness and flexibility in function approximation [Liu et al. \(2024c\)](#). This design enables KANs to model complex transformations more efficiently while maintaining better interpretability and scaling properties compared to [Liu et al. \(2024c\)](#). Additionally, KANs have demonstrated potential in computer vision-related tasks [Cheon \(2024\)](#). Based on this, we hypothesize that replacing MLPs with KANs in brain disorder diagnosis models can better capture intricate relationships in FC patterns, leading to more robust and individualized diagnoses.

In this paper, we propose *ABFR-KAN* (Advanced Brain Function Representation with Kolmogorov-Arnold Networks), a novel workflow for brain disorder diagnosis (Fig. 1). Build-

ing upon state-of-the-art methods, we propose novel sampling and function representation strategies and investigate the impact of KANs under various configurations in transformer networks. Our specific contributions in the present paper are summarized as:

- Randomized anchor patch selection, which helps avoid structural bias, boosts individual-specific representations, and increases robustness and generalizability by reducing dependence on atlas-based parcellation.
- Multi-scale sampling of patches from a subject’s brain, aimed at creating multiple functional representations for the same subject, introducing variance while preserving meaningful FC information.
- Extensive experimentation demonstrating the effectiveness of replacing traditional MLP components in transformer networks with KANs.

This manuscript is an extension of our paper “Improving brain disorder diagnosis with advanced brain function representation and Kolmogorov-Arnold Networks” presented at the 2025 Medical Imaging with Deep Learning (MIDL) conference [Ward and Imran \(2025\)](#). This manuscript substantially extends the work by adding a thorough literature review, additional experiments across new model backbones and KAN variants, and a more detailed description of the methods and results discussion.

2. Related Works

2.1 Brain Function Representation

There generally exist three different setups for brain disorder analysis using an atlas: single-atlas, multi-atlas, and individual-specific atlas. In the single-atlas setup, a fixed, predefined brain atlas is used to partition the brain into ROIs, and FC is computed between these regions. This results in a consistent graph structure across subjects, enabling straightforward comparison but potentially limiting representational flexibility due to reliance on a single spatial scale and fixed anatomical definitions [Dickie et al. \(2017\)](#). An example of a single-atlas approach is BrainGNN [Li et al. \(2021\)](#), which constructs subject-specific brain graphs using the Desikan-Killiany [Desikan et al. \(2006\)](#) atlas, where nodes correspond to ROIs and edges represent pairwise functional connectivity derived from fMRI time series. The model then applies graph neural network (GNN) operations along with ROI-aware convolution and pooling mechanisms to learn discriminative biomarkers while maintaining interpretability.

In contrast, the multi-atlas setup leverages multiple atlases simultaneously to capture complementary views of brain organization at different spatial resolutions or anatomical perspectives. Since no single atlas fully characterizes

the complexity of brain structure and function, combining multiple atlases can provide richer and more robust representations [Dickie et al. \(2017\)](#). In this setting, FC networks are constructed separately for each atlas and then fused at either the feature level or the graph level to form a more comprehensive representation. An example of this can be seen in the work of [Lee et al. \(2024\)](#), who leveraged a spectral GNN-based framework that constructs a holistic functional connectivity network by integrating ROI signals across multiple atlases [Kennedy et al. \(1998\)](#); [Craddock et al. \(2012\)](#); [Rolls et al. \(2020\)](#), capturing both intra-atlas relationships and inter-atlas interactions. Their approach incorporates both early fusion (combining ROI signals before graph construction) and late fusion (aggregating model outputs), enabling the model to exploit multi-scale connectivity patterns for improved disease classification.

The individual-specific atlas setup aims to overcome limitations of group-level atlases by tailoring brain parcellations to each subject. Instead of relying on fixed ROI definitions, this approach constructs personalized functional regions based on subject-specific fMRI data, thereby capturing inter-subject variability in brain organization. This can lead to more precise and discriminative functional connectivity representations, particularly for neurological disorders where individual differences are significant [Taimouri and Ravindra \(2025\)](#). For example, PFC-DBGNN-STAA by [Cui et al. \(2023\)](#) first generates a personalized FC template for each subject by aligning ROIs based on their individual FC patterns. It then combines both individual-specific and group-level representations within a dual-branch GNN and further enhances them using spatio-temporal attention mechanisms to capture both spatial dependencies and temporal dynamics in fMRI signals.

As an alternative to using pre-defined atlases for ROI parcellation, several data-driven approaches have been proposed. For example, attention-guided hybrid deep learning networks have been used to localize discriminative brain regions automatically for Alzheimer's disease and MCI diagnosis [Lian et al. \(2020\)](#). RandomFR [Liu et al. \(2024a\)](#) is an innovative approach for brain function representation and operates via a randomized selection of brain patches as well as the use of novel function and position description methods. RandomFR serves as the main inspiration for the research presented in this paper.

2.2 Kolmogorov-Arnold Networks

For decades, there has been debate regarding the optimal number of hidden layers/nodes when designing neural networks (NNs) [Stathakis \(2009\)](#). Such open questions regarding hidden layers have contributed to the lack of interpretability within NNs, thus giving them their reputation of being “black boxes” [Tan et al. \(2015\)](#). The Kolmogorov-

Arnold representation theorem has long been thought to explain the use of more than one hidden layer in NNs [Hecht-Nielsen \(1987\)](#), although this is not a universally held belief [Schmidt-Hieber \(2021\)](#), and many investigations of building NNs based on this theorem occurred before the rise of DL [Liu et al. \(2024c\)](#). Addressing this, [Liu et al. \(2024c\)](#) introduced their aptly named Kolmogorov-Arnold Network, or KAN, which provided a generalized form of the original Kolmogorov-Arnold representation to arbitrary widths and depths, boosting its practicality in a world driven by DL.

The core idea of the KAN framework is fundamentally different from traditional MLPs. In an MLP, learnable linear weight matrices are applied along edges, followed by fixed nonlinear activation functions at nodes. In contrast, KANs eliminate linear weights entirely. Each connection is instead parameterized by a learnable univariate function, and nodes perform only summation of incoming signals. Thus, KANs are not simply MLPs with learnable activation functions, but rather replace all scalar weights with function-valued parameters defined on edges. In practice, each edge encodes a univariate function implemented as a B-spline with trainable coefficients, and a KAN layer can be viewed as a matrix of such functions. The network is then formed as a composition of these nonlinear function matrices.

Because KANs learn both the compositional structure of multivariate functions via network connectivity and flexible univariate transformations via spline-parameterized edge functions, they achieve high expressiveness. At the same time, each learned component corresponds to an explicit, low-dimensional function that can be directly visualized or approximated symbolically, improving interpretability compared to standard MLPs. However, this expressive design introduces computational inefficiencies. Unlike MLPs, where each connection involves a single scalar multiplication, KANs require evaluating a learnable function on every edge. As a result, intermediate representations must store and propagate the outputs of many per-edge function evaluations rather than simple weighted sums, increasing both memory usage and computational cost during both forward and backward passes.

Addressing these performance bottlenecks, many improvements of KANs have been proposed. Particularly, in terms of identifying more effective and efficient activation functions compared to the initial residual activation function. One such improvement lies in the Efficient-KAN [Blealtan and Dash \(2024\)](#) architecture, which leverages the linearity of B-spline basis functions instead of computing each spline activation independently for every edge via expensive tensor expansion. Specifically, Efficient-KAN restructures the computation by applying the fixed basis functions to the input once and then combining the results via a learnable weight matrix, akin to matrix multiplication. This avoids high-dimensional tensor expansions and makes the forward

and backward passes both memory- and compute-efficient.

An alternative implementation, FastKAN [Li \(2024\)](#), replaces the 3-order B-spline basis in the original KAN with radial basis functions (RBFs). [Li \(2024\)](#) demonstrates that this change has a positive impact, improving forward speed compared to Efficient-KAN by $33\times$. This indicates that Gaussian RBFs are capable of approximating the B-spline basis, which is the bottleneck of KAN and Efficient-KAN.

FasterKAN, by [Delis \(2024\)](#), further improves the speed of KAN by using approximations of the B-spline via the reflectional switch function, inspired by [Deng \(2009\)](#), modified to have reflectional symmetry. Designing the activation function in this manner allows FasterKAN to approximate the 3rd-order B-splines used in the original KAN implementation. Experimentation with this setup revealed that FasterKAN achieved speeds $1.5\times$ faster than FastKAN. At the time this research was performed, FasterKAN was the most efficient KAN framework.

Since then, other KAN frameworks such as RBF-KAN [Sidharth \(2024\)](#), which uses RBFs as the basis function similar to FastKAN, and ChebyKAN [Guo et al. \(2024\)](#), which employs Chebyshev polynomials in place of B-splines, have demonstrated competitive performance with FasterKAN. Additionally, hybrid structures such as those combining B-splines with Gaussian RBFs [Ta \(2024\)](#) or wavelet functions [Seydi et al. \(2024\)](#) are interesting and have demonstrated good performance across a variety of tasks. Finding the best basis function and architecture for KAN is still an active research topic, and best practices are expected to remain dynamic and ever-changing.

Up to this point, our literature survey of KANs has focused on the history and development process. To shift focus to the application side, the original use case of KANs reported in [Liu et al. \(2024c\)](#) lay in improving AI + Science tasks such as function fitting, partial differential equation (PDE) solving, symbolic regression, and scientific discovery. The original authors, as well as subsequent work [Liu et al. \(2024b\)](#), have demonstrated the effectiveness of KANs in this regard, but it is far from the only use case.

Early work in expanding the applicability of KANs came in the form of assessing their suitability for computer vision tasks [Azam and Akhtar \(2024\)](#). It was discovered that on simple classification/segmentation tasks using the MNIST, CIFAR-10, and CamVid datasets, KAN-based vision models outperformed MLP and convolutional neural network (CNN)-based approaches in both accuracy and scalability, demonstrating their suitability for computer vision tasks. For computer vision applications, KANs have also demonstrated strong performance in medical imaging tasks [Li et al. \(2025\)](#). Our work in this paper leverages VisionKAN [Chen et al. \(2024\)](#), a repository investigating whether KANs can replace MLPs in Vision Transformers. The next subsection will discuss the use of KANs specifically for imaging-based

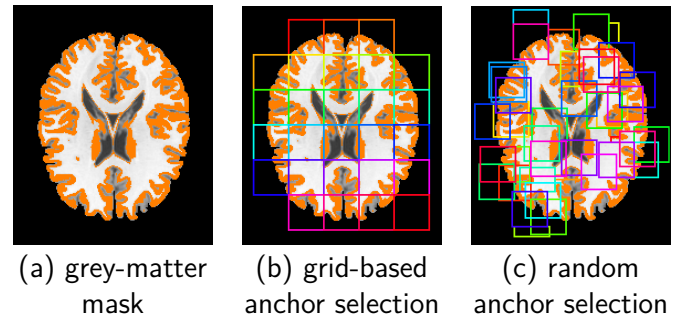


Figure 2: (a) The GM mask (highlighted in orange and overlaid over the *ch2bet* template produced by [Holmes et al. \(1998\)](#)) from which our anchor patches are selected. (b) The baseline grid-based anchor selection process. (c) Our randomized anchor selection process, which reduces structural bias and enhances individual specificity.

brain diagnostic tasks.

2.3 KANs for Brain Diagnostic Tasks

[Wang et al. \(2024a\)](#) explored the effectiveness of various KAN-based neural networks (KAN, FastKAN, and ChebyKAN) in detecting cognitive load from multi-channel EEG signals. They found that on a 19-channel EEG dataset constructed during a mental arithmetic task, the best KAN-based model outperformed the accuracy of existing state-of-the-art deep learning approaches like EEGNet and CNN-BiLSTM by 1.04% and 3.22%, respectively, with [Liu et al. \(2024c\)](#)'s original KAN implementation outperforming the FastKAN and ChebyKAN approaches. Additional KAN variants were explored on EEG data by [Thant and Panitanarak \(2025\)](#), and it was found that Wavelet KAN outperformed the other KAN variants. [Hasan et al. \(2024\)](#) also explored the use of KANs on EEG data, designing a KAN-based model for hand movement classification from motor imagery EEG. The authors found that their KAN-based approach outperformed the conventional MLP-based approach by a magnitude of 19.5%. Contrary to these studies, where the use of KAN-based architectures was found to be beneficial, [Hasan et al. \(2025\)](#) compared the performance of KAN to that of LSTMs for epileptic seizure prediction from EEG data and found the performance of KAN to be drastically below that of LSTM, by a margin of 15.22% accuracy.

[Ding et al. \(2025\)](#) explored the efficacy of KANs for Alzheimer's disease (AD) diagnosis by integrating KAN into graph convolutional networks (GCNs) to enhance diagnostic accuracy and interpretability. Evaluated on the ADNI dataset, the authors found that their GCN-KAN model outperformed traditional GCNs by 4-8% in classification accuracy while providing interpretable insights into key brain regions associated with AD. The applicability of hybrid KAN + GCN approaches for AD diagnosis was further validated

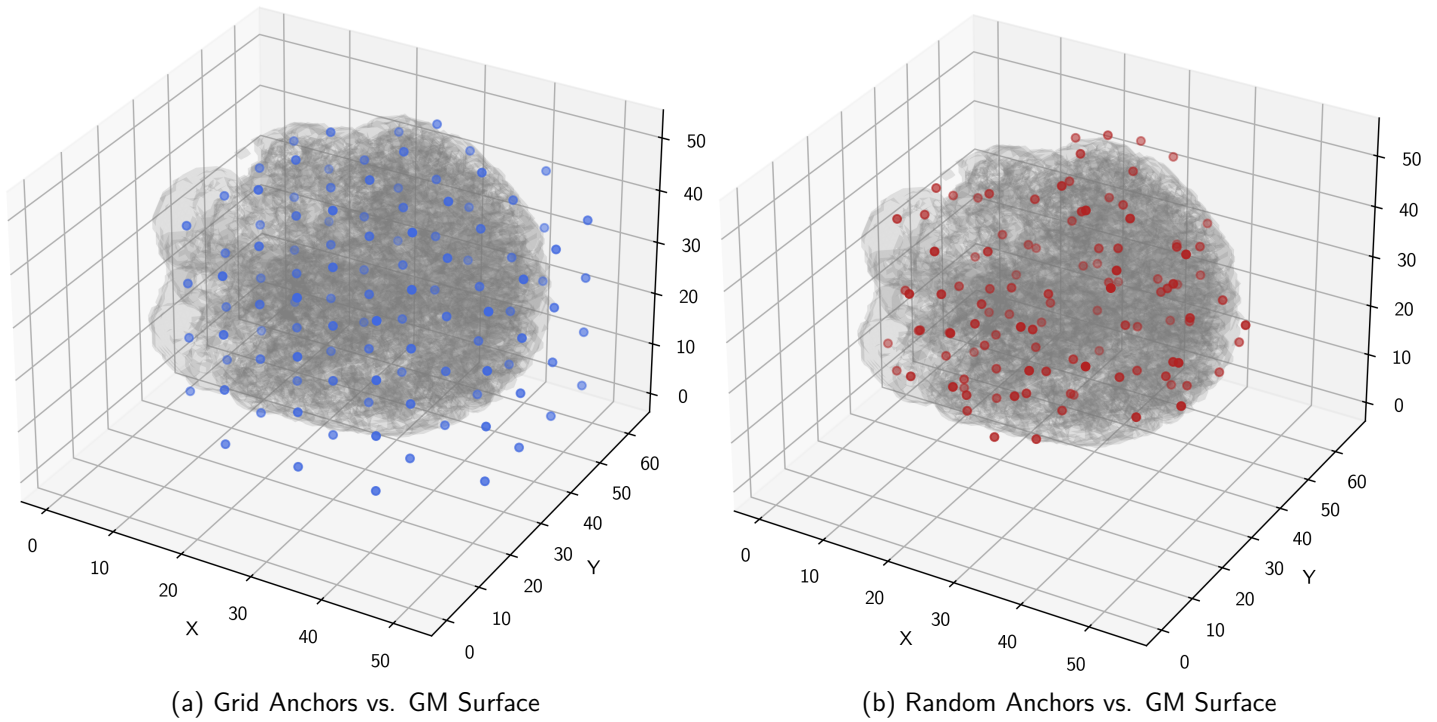


Figure 3: Visualization of anchor patch centers overlaid on the GM mask (represented as a 3D point cloud of voxels). Anchors selected with a grid-based strategy (left) exhibit a rigid lattice distribution that extends into non-GM regions and fails to conform to cortical geometry. In contrast, randomly selected anchors (right) are distributed more adaptively within the GM and better follow the brain's anatomical shape.

by Wang (2025), who found that a KAN + GCN model outperformed traditional random forest (RF), support vector machine (SVM), MLP, and CNN approaches on the ADNI dataset. Verma et al. (2025) also explored KANs for AD diagnosis, integrating KAN layers as densely connected layers into the VGG-19 architecture. Their hybrid architecture, dubbed VGG-KAN, achieved a classification accuracy of 95.67% on the OASIS dataset and an accuracy of 97.86% on the ADNI dataset.

Wang et al. (2025) demonstrated the feasibility of KAN and its variants for chemical exchange saturation transfer (CEST)-MRI data analysis. On a CEST-MRI dataset of 27 subjects, the authors explored the performance of MLP, KAN, and Lorentzian-KAN (LKAN) against a traditional multi-pool Lorentzian fitting (MPLF) method in generating various CEST parameters. They found that KAN and LKAN both outperformed MLP across the experiments. Sticking with MRI as the modality, Penkin and Krylov (2024) examined the capabilities of three different KAN bottlenecks (linear, B-spline, and Chebyshev polynomials) as deep feature extractors for MRI reconstruction, finding that Chebyshev polynomials performed the best. Giupponi et al. (2025) evaluated the performance of KANs against CNNs for brain age prediction from MRIs, and found that KAN-based models reduced estimation errors, highlighting their potential for improving brain age assessment. Overall, these works span a wide range of tasks and model architectures,

illustrating that KANs are not tied to a single network design but can be integrated into diverse neural architectures.

3. Methods

Following Liu et al. (2024a) for brain function representation, our proposed ABFR-KAN workflow is divided into three stages: **sampling, function representation, and classification network**. In the sampling stage, anchor patches are selected from the GM region of rs-fMRI scans. The GM mask is obtained by segmenting the structural brain image aligned with the functional brain image. In the function representation stage, each sampled brain patch is described using two types of information: where it is located in the brain and how it functionally relates to other regions. The spatial location is encoded using its coordinates in Montreal Neurological Institute (MNI) space. To capture functional relationships, we compare the BOLD signal of the sampled patch with the BOLD signals of a set of anchor patches. Specifically, we compute the Pearson correlation between the time series of the sampled patch and each anchor patch. These correlation values form a vector that summarizes how strongly the sampled patch is functionally connected to all anchors. Collecting these vectors for all sampled patches produces a functional representation matrix, where each row corresponds to one sampled patch and each column corresponds to one anchor patch. In the

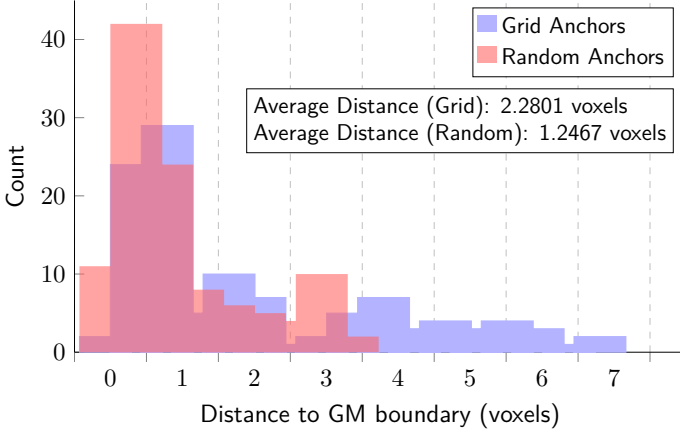


Figure 4: Histogram of distances from anchor patch centers to the nearest GM boundary voxel. Our proposed anchor selection method (red bars) produces anchor patches significantly closer to the GM surface than anchors selected on a set grid (blue bars, the baseline). This indicates that randomly selected anchor patches reduce spatial bias and achieve stronger conformity to anatomy compared to the baseline.

classification network, embeddings based on the fusion of the function and position descriptions are passed to the classifier for final predictions.

3.1 Random Anchor Selection

Preliminaries: $\Omega \subset \mathbb{Z}^3$ denotes a discrete 3D brain volume in MNI space with spatial dimensions:

$$\Omega = \{0, \dots, X-1\} \times \{0, \dots, Y-1\} \times \{0, \dots, Z-1\}. \quad (1)$$

$M : \Omega \rightarrow \{0, 1\}$ is a GM mask, where $M(\mathbf{v}) = 1$ indicates that voxel $\mathbf{v} \in \Omega$ belongs to GM. Cubic patches of side length $s \in \mathbb{N}$ are defined by their top-left coordinate:

$$\mathbf{t} = (t_x, t_y, t_z) \in \Omega, \quad (2)$$

with the associated voxel set:

$$\mathcal{P}(\mathbf{t}) = \left\{ (x, y, z) \in \Omega \mid \begin{array}{l} t_x \leq x < t_x + s, \\ t_y \leq y < t_y + s, \\ t_z \leq z < t_z + s \end{array} \right\}. \quad (3)$$

We define the GM support of a patch as the functionally relevant tissue a patch actually contains. We quantify this as:

$$G(\mathbf{t}) = \sum_{\mathbf{v} \in \mathcal{P}(\mathbf{t})} M(\mathbf{v}). \quad (4)$$

A patch is considered valid if $G(\mathbf{t}) \geq \tau$, where $\tau > 0$ is a predefined minimum GM overlap threshold. For this research, a τ of 100 was used, meaning that for a patch to

be considered valid, it must contain at least 100 GM voxels. The geometric center of a patch is given as:

$$\mathbf{c}(\mathbf{t}) = \mathbf{t} + \left(\frac{s}{2}, \frac{s}{2}, \frac{s}{2}\right). \quad (5)$$

$\mathcal{A} = \{\mathcal{P}(\mathbf{t}_j)\}_{j=1}^H$ denotes a set of anchor patches, which are used as fixed reference regions for defining functional relations like correlations with sampled brain patches. $\mathcal{A}$ is considered valid if each patch is spatially contained within Ω and satisfies the GM constraint $G(\mathbf{t}_j) \geq \tau$.

Proposed Method: In our proposed anchor selection mechanism, anchor patches are generated via random sampling. Specifically, anchor top-left coordinates are sampled from a discrete uniform distribution:

$$\mathbf{t} \sim \text{Unif}([0, X-s] \times [0, Y-s] \times [0, Z-s]), \quad (6)$$

subject to the GM constraint $G(\mathbf{t}) \geq \tau$. Sampling continues until a predefined number H of valid anchors ($\mathcal{A}_{\text{rand}}$) is obtained. Our proposed method differs from Liu et al. (2024a) in the anchor patch selection process (grid-based vs random). Fig. 2 highlights the differences between the two approaches.

$$\text{ROI} = [(x_{\min}, x_{\max}) \times (y_{\min}, y_{\max}) \times (z_{\min}, z_{\max})]. \quad (7)$$

Given a stride vector $\delta = (\delta_x, \delta_y, \delta_z)$, they construct candidate anchor locations as:

$$\mathbf{t}_{ijk} = (x_{\min} + i\delta_x + o_x, y_{\min} + j\delta_y + o_y, z_{\min} + k\delta_z + o_z), \quad (8)$$

where $i, j, k \in \mathbb{N}$ index the grid and (o_x, o_y, o_z) are deterministic offsets chosen to approximately center the grid within the ROI. The resulting anchor set, $\mathcal{A}_{\text{grid}}$, forms an equidistant tiling of the ROI up to boundary effects.

This approach, while reasonably effective, limits models in terms of flexibility and adaptability because the same grid is used for every subject in a dataset, imposing a structural bias and a disregard for individual specificity, as a subject may have functionally distinct regions that do not align well with predefined anchor patches. Additionally, selecting anchor patches in this ROI, grid-based manner poses the risk of many patch centers falling outside of the GM region of the brain (as demonstrated in Fig. 3(a)). For FCA, this poses a problem, as the goal should be to extract the most functionally relevant patches from the GM.

Our randomized approach to anchor selection fixes this (see Fig. 3)(b), offering anchor patches that are more diverse across patients, as well as more conform to brain anatomy. This ensures that more functionally relevant patches are extracted from the anchor regions. The better anatomical conformity is validated by Fig. 4. As seen, our proposed anchor selection approach produces anchor patches whose centers are 45.32% closer to the GM boundary compared to those produced by the baseline approach.

3.2 Multi-Scale Patch Sampling

Preliminaries: rs-fMRI data is represented as a 4D array $\mathbf{X} \in \mathbb{R}^{T \times X \times Y \times Z}$, where $\mathbf{X}_t(\mathbf{v})$ denotes the BOLD signal at time $t \in \{1, \dots, T\}$ and voxel $\mathbf{v} \in \Omega$. All volumes are assumed to be in the same MNI coordinate system as the previously selected anchor patches. $L: \Omega \rightarrow \{0, 1, \dots, H\}$ denotes the anchor label image, where $L(\mathbf{v}) = j$ indicates that voxel $\mathbf{v}$ belongs to anchor patch j , and $L(\mathbf{v}) = 0$ denotes background. For anchor j , the GM-restricted support is defined as:

$$\mathcal{A}_j = \mathbf{v} \in \Omega | L(\mathbf{v}) = j. \quad (9)$$

For a patch side length s , a sampled patch is indexed by its center coordinate:

$$\mathbf{p} = (p_x, p_y, p_z) \in \Omega, \quad (10)$$

which induces a cubic voxel set:

$$\mathcal{P}(\mathbf{p}; s) = \left\{ \mathbf{v} \in \Omega \left| \begin{array}{l} |v_x - p_x| < \frac{s}{2}, \\ |v_y - p_y| < \frac{s}{2}, \\ |v_z - p_z| < \frac{s}{2} \end{array} \right. \right\}. \quad (11)$$

with boundary clipping implicitly enforced by the condition $\mathbf{v} \in \Omega$. A sampled patch is considered valid if its GM support satisfies $G(\mathbf{p}; s) \geq \theta$. In this research, the threshold is set to $\theta = 1$.

For any voxel set $\mathcal{S} \subseteq \Omega$ with nonzero cardinality, the mean BOLD time series is defined as:

$$\mu(\mathcal{S})_t = \frac{1}{|\mathcal{S}|} \sum_{\mathbf{v} \in \mathcal{S}} \mathbf{X}_t(\mathbf{v}). \quad (12)$$

The anchor time series for anchor j is given by $\mathbf{a}_j = \mu(\mathcal{A}_j)$. Similarly, the time series for sampled patch i , centered at $\mathbf{p}_i$, is computed over GM only as:

$$\mathbf{b}_i = \mu(\mathcal{P}(\mathbf{p}_i; s) \cap \{\mathbf{v} | M(\mathbf{v}) = 1\}). \quad (13)$$

To retain spatial information, the center coordinate of each sampled patch is normalized by the image dimensions, yielding:

$$\tilde{\mathbf{p}}_i = \left(\frac{p_{ix}}{X}, \frac{p_{iy}}{Y}, \frac{p_{iz}}{Z} \right) \in [0, 1]^3. \quad (14)$$

The position-aware patch representation, $\mathbf{u}$, is formed by concatenating the normalized position with the patch time series:

$$\mathbf{u}_i = ([\tilde{\mathbf{p}}_i \ \mathbf{v}_i]) \in \mathbb{R}^{3+T}. \quad (15)$$

Given N sampled patches, the patch-to-anchor FC matrix $\mathbf{F} \in \mathbb{R}^{N \times H}$ is defined by:

$$F_{ij} = \text{Corr}(\mathbf{v}_i, \mathbf{a}_j), \quad (16)$$

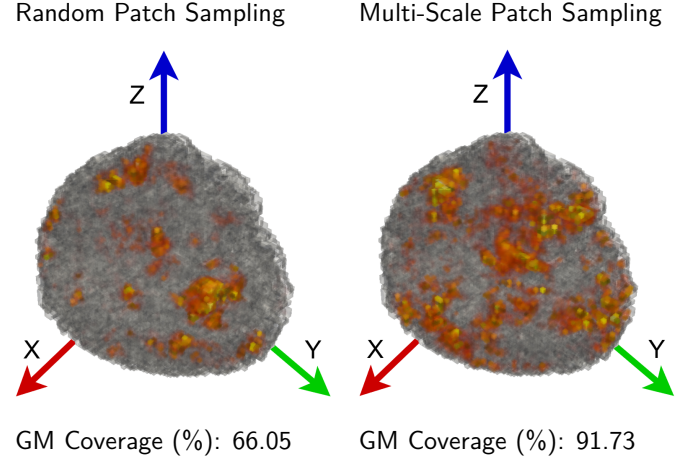


Figure 5: GM coverage of the baseline random patch sampling strategy and the proposed multi-scale patch sampling. Visualized are the centers of the extracted patches, with hotter (indicated by the color yellow) regions in the heatmap indicating areas in the GM captured in multiple patches. Multi-scale patch sampling yields a substantially denser and more uniform coverage ($\uparrow 25.68\%$) of cortical GM compared to single-pass random sampling.

where $\text{Corr}(\cdot, \cdot)$ denotes Pearson correlation. Each row of $\mathbf{F}$ represents the functional relationship between one sampled patch and all anchor patches.

Proposed Method: Our proposed patch sampling approach performs multiple independent sampling iterations with different patch sizes. Fig. 10 shows a comparison with the baseline method. Specifically, we use $R = 3$ iterations with patch sizes $s_1 = 8$, $s_2 = 12$, and $s_3 = 16$. An ablation study demonstrates the impact of the number of iterations (Fig. 8). In each iteration, 256 valid patches are sampled using the same acceptance rule, producing iteration-specific FC matrices, $\mathbf{F}^{(r)}$, and position-aware representations, $\{\mathbf{u}_i^{(r)}\}$. We obtain the final FC representation by averaging the FC matrices across iterations, while aggregating the position-aware representations via concatenation.

The baseline method Liu et al. (2024a) that our proposed approach aims to improve upon simply samples patches once per subject using a fixed patch size of $s = 8$. Candidate patch centers are drawn uniformly from the brain volume and accepted if they satisfy the GM support constraint. Sampling terminates when $N = 256$ valid patches have been obtained. Anchor time series are computed first, followed by the extraction of sampled patch time series and the construction of the FC matrix, $\mathbf{F}$.

Our proposed multi-scale sampling approach enhances the robustness of extracted rs-fMRI features by reducing the impact of single-scale patch selection biases and improving anatomical coverage (as demonstrated in Fig. 5). Additionally, our proposed approach effectively reduces the

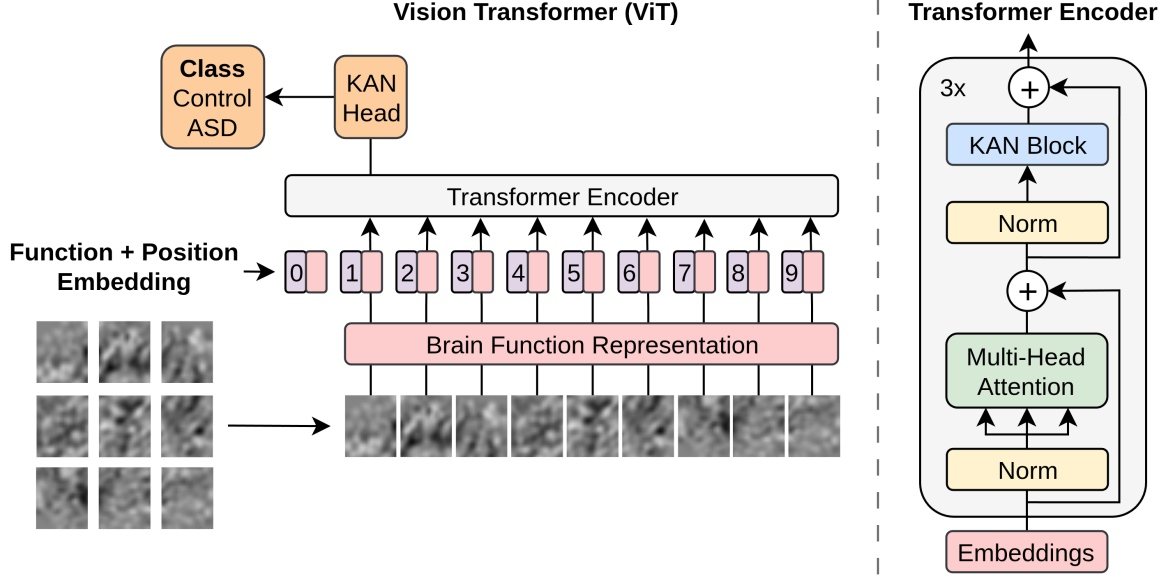


Figure 6: The classification network used in our proposed ABFR-KAN framework is a ViT-style encoder-only network. The transformer is fed patches derived from rs-fMRIs that are embedded with spatial position information and passed through the encoder. The binary classification prediction (control or ASD) is produced by the KAN head. In the encoder, we replace the MLP block with a KAN block.

sampling variance in FC estimation (Fig. 8). The reduction of this variance should improve feature stability, reduce the risk of overfitting to sampling artifacts, and improve reproducibility between sites.

3.3 Classification Network

For our proposed ABFR-KAN approach, we utilized a vision transformer (ViT)-style architecture similar to the one introduced by Dosovitskiy et al. (2020). The specific ViT architecture is visualized in Fig. 6 and is discussed in detail in this section.

Inputs: For each subject, the model receives a set of sampled brain patches represented by their functional relationships to anchor patches and their spatial locations. Specifically, after patch sampling and anchor-based FC computation, each subject is represented by a patch-anchor FC matrix, $\mathbf{F} \in \mathbb{R}^{N \times H}$, where N is the number of sampled patches and H is the number of anchor patches. Each row corresponds to the FC profile of one sampled patch with respect to all anchors. A corresponding set of normalized spatial coordinates, $\{\tilde{\mathbf{p}}_i\}_{i=1}^N \subset [0, 1]^3$ encodes the spatial location of each sampled patch center. These components jointly characterize both the functional and spatial properties of the sampled patches.

Patch Selection: Before entering the transformer encoder, a learnable Top-K pooling operation is applied to reduce the number of patch tokens and emphasize the most informative ones. Each patch feature vector, which is a row of $\mathbf{F}$, $\mathbf{f}_i \in \mathbb{R}^H$, is assigned an importance score through a learnable

projection followed by a nonlinearity. Patches are then ranked by these scores, and only the top 80% of patches are retained. The same selection is applied consistently to both the functional features and their corresponding spatial coordinates. This operation encourages the model to focus on functionally salient patches that are most discriminative for the classification task.

Positional Encoding: To retain spatial information after Top-K pooling, each retained patch’s normalized spatial coordinate is projected into the same embedding space as the functional features using a learnable linear transformation. The resulting positional embedding is added to the patch’s functional feature vector. This additive fusion allows the transformer to jointly reason about FC patterns and spatial context, without imposing any predefined spatial ordering or grid structure Dosovitskiy et al. (2020).

Transformer Encoder: The resulting set of patch embeddings is processed by a multi-layer transformer encoder composed of stacked self-attention and feedforward blocks. Within each self-attention layer, patches are treated as tokens and allowed to attend to one another. In contrast to pairwise FC, self-attention captures global dependencies across all sampled patches simultaneously, allowing the network to model distributed functional patterns associated with ASD. Each transformer layer follows a standard pre-normalization design in that layer normalization is applied before both the attention and feedforward blocks, and residual connections are used to ensure stable gradient flow and preserve information across layers. Through successive layers, patch representations are progressively refined,

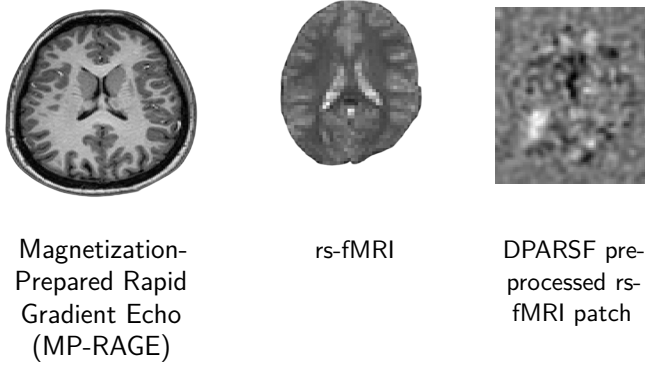


Figure 7: Raw and preprocessed data for a single patient in the ABIDE I dataset.

integrating information from across the brain.

Global Aggregation and Classification: After being passed through the encoder, the patch-level representations are aggregated into a single subject-level embedding by mean pooling across the patch dimension. This pooling operation produces a fixed-length representation regardless of the number of retained patches. The aggregated embedding is then passed through a lightweight classification head consisting of layer normalization followed by a linear classifier. The output is a two-dimensional logit vector corresponding to the ASD and control classes.

4. Experimental Evaluation

4.1 Data

We evaluated our proposed ABFR-KAN using pre-processed neuroimaging data from the Autism Brain Imaging Data Exchange (ABIDE) [Craddock et al. \(2013\)](#); [Di Martino et al. \(2014\)](#) I and II.

ABIDE I: The preprocessed ABIDE I repository contains data collected from a total of 1,112 patients at various sites. In this work, we trained and evaluated our ABFR-KAN approach on data from the New York University (NYU) Langone Medical Center and the University of Michigan (UM) Functional MRI Center. Both sites were processed using the Data Processing Assistant for Resting-State fMRI (DPARSF) [Yan and Zang \(2010\)](#) method. In total, 64 male (age range: 7-39) and 9 female (age range: 10-38) patients with ASD diagnoses were selected from the NYU site, along with 72 male (age range: 6-31) and 26 female (age range: 8-29) patients from the control group. From the UM site, we selected a balanced dataset consisting of information from 55 patients in the control group (46 male, aged 8 to 18; 9 female, aged 9 to 18) and 55 in the ASD group (38 male, aged 8 to 18; 17 female, aged 9 to 19).

ABIDE II: The ABIDE II repository contains data from a total of 1,114 patients, collected from 19 sites. In this

work, we select data from the largest ABIDE II site, the Kennedy Krieger Institute (KKI). Preprocessing on this site is consistent with that of the NYU and UM sites, which has already been discussed. In this data subset, there were a total of 211 patients, with an age range of 8-13 years old. This subset represents the most imbalanced of the data subsets we explore in this work, with 56 patients having an ASD diagnosis, and 155 belonging to the control group.

Within each data subset, ASD diagnoses were made using a combination of clinical judgment and standardized “gold standard” instruments, specifically the Autism Diagnostic Observation Schedule (ADOS) [Lord et al. \(2000\)](#) and the Autism Diagnostic Interview–Revised (ADI-R) [Rutter et al. \(2003\)](#), with the diagnostic categories based on the Diagnostic and Statistical Manual of Mental Disorders, Fourth Edition, Text Revision (DSM-IV-TR) [American Psychiatric Association \(2000\)](#) criteria. The bias towards the male sex that is shown in both data sites is consistent with ASD literature, where it has been shown that males are about four times as likely to be diagnosed with ASD compared to females [Leow et al. \(2024\)](#). A visualization of the data types extracted from the ABIDE I dataset is shown in Fig. 7.

4.2 Implementation Details

Machine Configuration: The ABFR-KAN model was implemented in PyTorch and trained on an Intel (R) Xeon (R) w7-2475X, 2600MHz machine with dual NVIDIA A4000X2 GPUs (32GB).

Baselines: We compare our proposed ABFR-KAN against state-of-the-art (SOTA) models for brain disorder diagnosis. These include a support vector machine (SVM) [Cortes and Vapnik \(1995\)](#); a CNN-based method in BrainNetCNN [Kawahara et al. \(2017\)](#); a graph attention network (GAT) [Veličković et al. \(2018\)](#); GCN-based methods in the original GCN [Qin et al. \(2022\)](#) and MVS-GCN [Wen et al. \(2022\)](#); a GNN-based approach in BrainGNN [Li et al. \(2021\)](#), and transformer-based approaches in KD-Transformer [Zhang et al. \(2022\)](#) and RandomFR [Liu et al. \(2024a\)](#).

Model Backbones: We investigate the performance of ABFR-KAN under multiple backbones for the classification network. Specifically, we experiment using a CNN-based backbone in EfficientNetV2 [Tan and Le \(2021\)](#), transformer-based backbones using a ViT [Dosovitskiy et al. \(2020\)](#) and Data-efficient image Transformer (DeiT) [Touvron et al. \(2021\)](#), and a Mamba-based backbone in Vision Mamba (Vim) [Zhu et al. \(2024\)](#).

KAN Variants: We explored five different KAN variants in our ABFR-KAN framework: Efficient-KAN [Blealtan and Dash \(2024\)](#), FastKAN [Li \(2024\)](#), FasterKAN [Delis \(2024\)](#), Wav-KAN [Bozorgasl and Chen \(2024\)](#), and ChebyKAN [Guo et al. \(2024\)](#). Efficient-KAN, FastKAN, and FasterKAN

Table 1: Performance comparison of ABFR-KAN against state-of-the-art methods on the NYU and UM sites of the ABIDE I dataset and the KKI site of the ABIDE II dataset. Data samples were obtained via our proposed random anchor selection and multi-scale patch sampling strategy. Mean $\pm$ stdev scores are reported by averaging the results of each fold. Best and second-best results are **bolded** and underlined, respectively. (†) indicates baselines that ABFR-KAN improves upon ACC with statistical significance.

NYU Site Classification Performance						
Method	ACC	AUC	F1	PRE	SEN	SPE
SVM (†)	0.5615 $\pm$ 0.0301	0.4740 $\pm$ 0.1130	0.6750 $\pm$ 0.0237	0.5888 $\pm$ 0.0317	0.7974 $\pm$ 0.0681	0.2495 $\pm$ 0.1005
BrainNetCNN (†)	0.5438 $\pm$ 0.0390	0.5151 $\pm$ 0.0464	0.6527 $\pm$ 0.0305	0.5790 $\pm$ 0.0211	0.7495 $\pm$ 0.0550	0.2689 $\pm$ 0.0508
GAT (†)	0.5545 $\pm$ 0.0107	0.5500 $\pm$ 0.0115	0.6556 $\pm$ 0.0118	0.5894 $\pm$ 0.0181	0.7414 $\pm$ 0.0426	0.3061 $\pm$ 0.0490
GCN (†)	0.5758 $\pm$ 0.0336	0.5731 $\pm$ 0.0364	0.6806 $\pm$ 0.0290	0.5984 $\pm$ 0.0301	0.7896 $\pm$ 0.0300	0.2888 $\pm$ 0.0405
BrainGNN (†)	0.5731 $\pm$ 0.0132	0.5571 $\pm$ 0.0238	0.7285 $\pm$ 0.0106	0.5731 $\pm$ 0.0132	1.0000 $\pm$ 0.0000	0.0000 $\pm$ 0.0000
MVS-GCN (†)	0.6376 $\pm$ 0.0811	0.6163 $\pm$ 0.0746	0.7563 $\pm$ 0.0403	0.6269 $\pm$ 0.0739	<u>0.9700 $\pm$ 0.0600</u>	0.1990 $\pm$ 0.2282
KD-Transformer (†)	0.6610 $\pm$ 0.0272	0.6137 $\pm$ 0.0692	0.7225 $\pm$ 0.0216	0.6942 $\pm$ 0.0830	0.7768 $\pm$ 0.1014	0.6437 $\pm$ 0.0438
RandomFR	<u>0.7079 $\pm$ 0.0388</u>	<u>0.6544 $\pm$ 0.0583</u>	<u>0.7734 $\pm$ 0.0379</u>	<u>0.6986 $\pm$ 0.0435</u>	0.8784 $\pm$ 0.1024	<u>0.6783 $\pm$ 0.0486</u>
ABFR-KAN (†)	0.7427 $\pm$ 0.0342	0.7041 $\pm$ 0.0708	0.8003 $\pm$ 0.0200	0.7267 $\pm$ 0.0469	0.8984 $\pm$ 0.0646	0.7159 $\pm$ 0.0434
UM Site Classification Performance						
Method	ACC	AUC	F1	PRE	SEN	SPE
SVM (†)	0.6000 $\pm$ 0.0603	0.4866 $\pm$ 0.1278	0.6978 $\pm$ 0.0490	0.5655 $\pm$ 0.0834	0.9385 $\pm$ 0.0754	0.2802 $\pm$ 0.0971
BrainNetCNN (†)	0.5362 $\pm$ 0.0379	0.5563 $\pm$ 0.0526	0.5378 $\pm$ 0.0549	0.5460 $\pm$ 0.1043	0.5552 $\pm$ 0.0996	0.5404 $\pm$ 0.0998
GAT (†)	0.5435 $\pm$ 0.0422	0.5741 $\pm$ 0.0556	0.6042 $\pm$ 0.0481	0.5405 $\pm$ 0.0927	0.7108 $\pm$ 0.0846	0.3981 $\pm$ 0.0758
GCN (†)	0.5601 $\pm$ 0.0462	0.5898 $\pm$ 0.0579	0.5714 $\pm$ 0.0657	0.5609 $\pm$ 0.0792	0.6085 $\pm$ 0.1320	0.5337 $\pm$ 0.0557
BrainGNN (†)	0.4392 $\pm$ 0.0563	0.5521 $\pm$ 0.0365	0.3502 $\pm$ 0.2860	0.2722 $\pm$ 0.2275	0.5246 $\pm$ 0.4440	0.4964 $\pm$ 0.4393
MVS-GCN (†)	0.6727 $\pm$ 0.0445	0.6357 $\pm$ 0.0635	0.6953 $\pm$ 0.0571	0.6498 $\pm$ 0.0667	0.7744 $\pm$ 0.1488	0.5509 $\pm$ 0.2138
KD-Transformer (†)	0.6727 $\pm$ 0.0182	<u>0.6486 $\pm$ 0.0464</u>	<u>0.7233 $\pm$ 0.0356</u>	0.6194 $\pm$ 0.0487	0.8730 $\pm$ 0.0367	0.6674 $\pm$ 0.0381
RandomFR (†)	<u>0.7182 $\pm$ 0.0530</u>	0.6426 $\pm$ 0.1395	0.6999 $\pm$ 0.1311	<u>0.7169 $\pm$ 0.1104</u>	0.7397 $\pm$ 0.2197	<u>0.6984 $\pm$ 0.0780</u>
ABFR-KAN	0.7727 $\pm$ 0.0287	0.7184 $\pm$ 0.0620	0.7682 $\pm$ 0.0482	0.7840 $\pm$ 0.1086	<u>0.7837 $\pm$ 0.1314</u>	0.7768 $\pm$ 0.0270
KKI Site Classification Performance						
Method	ACC	AUC	F1	PRE	SEN	SPE
SVM (†)	0.6182 $\pm$ 0.0680	0.4538 $\pm$ 0.1260	0.0800 $\pm$ 0.1600	0.2000 $\pm$ 0.4000	0.0500 $\pm$ 0.1000	1.0000 $\pm$ 0.0000
BrainNetCNN (†)	0.4973 $\pm$ 0.0884	0.4358 $\pm$ 0.1379	0.2801 $\pm$ 0.1082	0.3455 $\pm$ 0.1553	0.2429 $\pm$ 0.0943	0.6714 $\pm$ 0.1057
GAT (†)	0.5756 $\pm$ 0.0746	0.5593 $\pm$ 0.0906	0.3496 $\pm$ 0.0883	0.4795 $\pm$ 0.1163	0.2972 $\pm$ 0.1098	0.7673 $\pm$ 0.1339
GCN (†)	0.5818 $\pm$ 0.0533	0.4717 $\pm$ 0.1318	0.1344 $\pm$ 0.1231	0.3609 $\pm$ 0.2960	0.0835 $\pm$ 0.0789	<u>0.9140 $\pm$ 0.0381</u>
BrainGNN (†)	0.6000 $\pm$ 0.0445	0.4974 $\pm$ 0.0357	0.0000 $\pm$ 0.0000	0.0000 $\pm$ 0.0000	0.0000 $\pm$ 0.0000	1.0000 $\pm$ 0.0000
MVS-GCN (†)	0.6000 $\pm$ 0.0445	0.4914 $\pm$ 0.1078	0.0000 $\pm$ 0.0000	0.0000 $\pm$ 0.0000	0.0000 $\pm$ 0.0000	1.0000 $\pm$ 0.0000
KD-Transformer (†)	0.6364 $\pm$ 0.0575	0.3757 $\pm$ 0.0642	0.2133 $\pm$ 0.1759	0.5000 $\pm$ 0.4472	0.1400 $\pm$ 0.1158	0.5557 $\pm$ 0.0509
RandomFR	<u>0.7455 $\pm$ 0.0600</u>	<u>0.6248 $\pm$ 0.1552</u>	<u>0.6457 $\pm$ 0.1855</u>	<u>0.7076 $\pm$ 0.1774</u>	<u>0.6200 $\pm$ 0.2159</u>	0.7148 $\pm$ 0.1162
ABFR-KAN	0.7636 $\pm$ 0.0445	0.6505 $\pm$ 0.0869	0.6588 $\pm$ 0.1329	0.8033 $\pm$ 0.1675	0.6300 $\pm$ 0.2015	0.7388 $\pm$ 0.0648

all utilize spline-based basis functions, Wav-KAN utilizes fixed wavelet transforms, and ChebyKAN uses truncated Chebyshev polynomial expansions on normalized inputs.

Training: A 5-fold cross-validation strategy was used to assess the model’s performance. On the NYU site, the four folds used for training contained 34 samples each, while the fold kept out of the training process contained 35 samples. For the UM site, the split was 22 samples in each fold, and on KKI, there were 42 samples in the folds used for training and 43 samples in the one kept out of training. For classification, we minimize the cross-entropy loss. The model was trained for 100 epochs, using the AdamW optimizer with a learning rate of 0.001. We used cosine annealing with warm restarts as the learning rate scheduler, along with a weight decay of 0.0001.

Evaluation: Model performance is gauged using traditional metrics for classification tasks, namely accuracy (ACC), area under the curve (AUC), F1 score (F1), precision (PRE), sensitivity (SEN), and specificity (SPE). For the non-AUC metrics, class labels are obtained by applying an argmax over the softmax output, corresponding to a decision threshold of 0.5 in the binary setting. For AUC computation, we use the continuous softmax probability of the positive class rather than binarized predictions, allowing evaluation across all possible decision thresholds.

4.3 Results

Table 2: Cross-site classification performance of ABFR-KAN compared to SOTA baselines. The data were extracted via our proposed random anchor selection and multi-scale patch sampling strategy. Mean $\pm$ stdev scores are reported by averaging the results of five runs. Best and second-best results are **bolded** and underlined, respectively. (†) indicates baselines that ABFR-KAN improves upon the ACC in a statistically significant manner.

Classification Performance (Train: NYU, Test: UM)						
Method	ACC	AUC	F1	PRE	SEN	SPE
SVM	<u>0.5527 $\pm$ 0.0260</u>	0.5170 $\pm$ 0.0706	0.5623 $\pm$ 0.0582	0.5479 $\pm$ 0.0204	0.5855 $\pm$ 0.1050	0.5200 $\pm$ 0.0676
BrainNetCNN (†)	0.5055 $\pm$ 0.0473	0.5067 $\pm$ 0.0452	0.5413 $\pm$ 0.1714	0.4850 $\pm$ 0.0588	0.6764 $\pm$ 0.2857	0.3345 $\pm$ 0.2336
GAT (†)	0.4909 $\pm$ 0.0359	0.5070 $\pm$ 0.0468	0.3322 $\pm$ 0.2777	0.3817 $\pm$ 0.1562	0.4145 $\pm$ 0.4223	0.5673 $\pm$ 0.3718
GCN (†)	0.5073 $\pm$ 0.0253	0.4976 $\pm$ 0.0397	0.2652 $\pm$ 0.2633	0.5101 $\pm$ 0.3165	0.2836 $\pm$ 0.3200	0.7309 $\pm$ 0.2801
BrainGNN (†)	0.5000 $\pm$ 0.0000	0.5869 $\pm$ 0.0093	0.6667 $\pm$ 0.0000	0.5000 $\pm$ 0.0000	1.0000 $\pm$ 0.0000	0.0000 $\pm$ 0.0000
MVS-GCN (†)	0.4945 $\pm$ 0.0159	0.5498 $\pm$ 0.0668	0.5682 $\pm$ 0.1883	0.4788 $\pm$ 0.0450	0.8073 $\pm$ 0.3418	0.1818 $\pm$ 0.3120
KD-Transformer (†)	0.5164 $\pm$ 0.0278	0.5035 $\pm$ 0.0619	0.3717 $\pm$ 0.2104	0.5029 $\pm$ 0.1373	0.3673 $\pm$ 0.2705	<u>0.6655 $\pm$ 0.2463</u>
RandomFR	0.6109 $\pm$ 0.0668	<u>0.6364 $\pm$ 0.0673</u>	<u>0.7224 $\pm$ 0.0379</u>	<u>0.6840 $\pm$ 0.0521</u>	0.7689 $\pm$ 0.0647	0.5662 $\pm$ 0.0714
ABFR-KAN	0.5517 $\pm$ 0.0462	0.7181 $\pm$ 0.0619	0.7473 $\pm$ 0.0320	0.7036 $\pm$ 0.0489	0.8021 $\pm$ 0.0598	0.6147 $\pm$ 0.0632
Classification Performance (Train: UM, Test: NYU)						
Method	ACC	AUC	F1	PRE	SEN	SPE
SVM (†)	0.5427 $\pm$ 0.0563	0.4909 $\pm$ 0.0618	0.5827 $\pm$ 0.1626	0.5846 $\pm$ 0.0365	0.6286 $\pm$ 0.2331	0.4274 $\pm$ 0.1875
BrainNetCNN (†)	0.4585 $\pm$ 0.0413	0.4805 $\pm$ 0.0402	0.2753 $\pm$ 0.1844	0.5343 $\pm$ 0.1367	0.2204 $\pm$ 0.1977	0.7781 $\pm$ 0.1785
GAT (†)	0.5392 $\pm$ 0.0183	0.4872 $\pm$ 0.0210	0.6758 $\pm$ 0.0298	0.5654 $\pm$ 0.0061	0.8449 $\pm$ 0.0883	0.1288 $\pm$ 0.0834
GCN (†)	<u>0.5556 $\pm$ 0.0414</u>	0.4979 $\pm$ 0.0243	0.6931 $\pm$ 0.0574	0.5696 $\pm$ 0.0176	0.8939 $\pm$ 0.1327	0.1014 $\pm$ 0.0816
BrainGNN (†)	0.4643 $\pm$ 0.0530	<u>0.5703 $\pm$ 0.0128</u>	0.2212 $\pm$ 0.2616	0.3561 $\pm$ 0.2925	0.2347 $\pm$ 0.3561	<u>0.7726 $\pm$ 0.3551</u>
MVS-GCN (†)	0.4795 $\pm$ 0.0409	0.5420 $\pm$ 0.0252	0.3473 $\pm$ 0.1927	<u>0.5887 $\pm$ 0.0392</u>	0.3122 $\pm$ 0.2825	0.7041 $\pm$ 0.2939
KD-Transformer (†)	0.5181 $\pm$ 0.0530	0.5001 $\pm$ 0.0531	0.5617 $\pm$ 0.1255	0.5747 $\pm$ 0.0464	0.5939 $\pm$ 0.2410	0.4164 $\pm$ 0.2306
RandomFR (†)	0.5252 $\pm$ 0.0330	0.4866 $\pm$ 0.1155	0.6064 $\pm$ 0.0159	0.5873 $\pm$ 0.0445	0.6279 $\pm$ 0.0512	0.4628 $\pm$ 0.0603
ABFR-KAN	0.5707 $\pm$ 0.0147	0.5714 $\pm$ 0.0500	0.6752 $\pm$ 0.0322	0.6128 $\pm$ 0.0397	0.7536 $\pm$ 0.0551	0.5019 $\pm$ 0.0584

4.3.1 Single-Site Performance

Table 1 shows the performance of our proposed ABFR-KAN compared to SOTA baselines on the NYU and UM sites of the ABIDE I dataset. In the table, ABFR-KAN is represented by the best-performing KAN variation and configuration as determined by ablation experiments we performed (See Section 4.3.4). For the NYU site, this means that the FastKAN variant was used in a KAN-KAN configuration, where FastKAN replaced the MLP components in both the ViT encoder and the classification head. For the UM site, ABFR-KAN is customized as the WavKAN variant, using an MLP-KAN configuration, where WavKAN replaces only the MLP classification head. Each of the compared-against baseline methods was trained and evaluated on the same data (that is, data that had been processed using our proposed random anchor selection and multi-scale patch sampling approach) as ABFR-KAN, to ensure fair comparison.

From the results, it is clear that ABFR-KAN consistently outperforms both the direct strong baseline in RandomFR, as well as all other compared baseline methods. When trained on the NYU site, ABFR-KAN outperforms RandomFR in terms of ACC, AUC, F1, PRE, SEN, and SPE by 3.48%, 4.97%, 2.69%, 2.81%, 2.00%, and 3.76%, respectively. When trained on the UM site, these improvements become 5.45%, 7.58%, 6.83%, 6.71%, 4.40%, and 7.84%,

respectively. On the KKI site, the improvements are by 1.81%, 2.57%, 1.31%, 9.57%, 1.00%, and 2.40%, respectively. On the third and fourth best-performing baselines (KD-Transformer and MVS-GCN, respectively), ABFR-KAN outperforms KD-Transformer on average by 7.94% on the NYU site, 6.657% on the UM site, and 30.4% on the KKI site, while outperforming MVS-GCN on average by 13.03% on the NYU site, 11.81% on the UM site, and 35.9% on the KKI site. Across the metrics for NYU and UM, the only metric where ABFR-KAN is deemed not superior is SEN; however, it should be noted that on the NYU site, the baseline models with the highest SEN values, BrainGNN and MVS-GCN, show strong evidence of false positive predictions, as indicated by their comparatively low SPE values. Compared to these, ABFR-KAN, with an average SEN of 0.8984 and SPE of 0.7159, demonstrates much better balance in meaningfully distinguishing patients with ASD from the control. The same relationship is demonstrated between SEN and SPE for the best-performing baselines (SVM, BrainGNN, MVS-GCN) on KKI in this regard, just in reverse, where the high SPE values for these three baselines indicate universal false negative predictions for these models, which is well addressed by our ABFR-KAN approach.

Statistical Analysis: To validate the statistical significance of our reported improvements over baselines, we utilized one-tailed paired t-tests over the reported perfor-

Table 3: Impact of various classification backbones within the ABFR-KAN framework on data from the UM site. Mean $\pm$ stdev scores are reported by averaging the results of 5-fold cross-validation. Best and second-best results are **bolded** and underlined, respectively.

Backbone	ACC	AUC	F1	PRE	SEN	SPE
EfficientNetV2	0.6818 $\pm$ 0.0643	0.6240 $\pm$ 0.0933	<u>0.7142 $\pm$ 0.0686</u>	0.6422 $\pm$ 0.0771	0.8172 $\pm$ 0.1130	<u>0.6772 $\pm$ 0.0645</u>
ViT	0.7182 $\pm$ 0.0530	0.6426 $\pm$ 0.1395	0.6999 $\pm$ 0.1311	0.7169 $\pm$ 0.1104	0.7397 $\pm$ 0.2197	0.6984 $\pm$ 0.0780
DeiT	0.6909 $\pm$ 0.0782	<u>0.6242 $\pm$ 0.0934</u>	0.7126 $\pm$ 0.1011	<u>0.6595 $\pm$ 0.0877</u>	<u>0.8205 $\pm$ 0.1947</u>	0.6766 $\pm$ 0.0779
Vim	<u>0.6909 $\pm$ 0.0340</u>	0.5856 $\pm$ 0.0764	0.7227 $\pm$ 0.0882	0.6309 $\pm$ 0.0505	0.8594 $\pm$ 0.1569	0.6740 $\pm$ 0.0560

Table 4: Impact of various anchor selection and patch sampling methods on the UM site using the baseline MLP-only method. Mean $\pm$ stdev scores are reported by averaging the results of 5-fold cross-validation. Best and second-best results are **bolded** and underlined, respectively.

Anchor Selection	Patch Sampling	ACC	AUC	F1	PRE	SEN	SPE
Grid-based	Random	<u>0.7000 $\pm$ 0.0545</u>	0.6347 $\pm$ 0.0840	0.7348 $\pm$ 0.0290	0.6609 $\pm$ 0.0353	0.8395 $\pm$ 0.1040	0.6899 $\pm$ 0.0871
Random	Random	0.6636 $\pm$ 0.0364	0.5639 $\pm$ 0.0962	<u>0.7030 $\pm$ 0.0752</u>	0.6124 $\pm$ 0.0545	<u>0.8359 $\pm$ 0.1410</u>	0.6427 $\pm$ 0.0400
Grid-based	Multi-scale	0.7000 $\pm$ 0.0617	0.6342 $\pm$ 0.1051	0.6226 $\pm$ 0.3131	0.5349 $\pm$ 0.2689	0.7470 $\pm$ 0.3784	0.6674 $\pm$ 0.0967
Random	Multi-scale	0.7182 $\pm$ 0.0530	0.6426 $\pm$ 0.1395	0.6999 $\pm$ 0.1311	0.7169 $\pm$ 0.1104	0.7397 $\pm$ 0.2197	0.6984 $\pm$ 0.0780

mance metrics. On the NYU site, ABFR-KAN demonstrates statistically significant improvement over KD-Transformer in terms of ACC, F1 (both $p < 0.01$), and SPE ($p < 0.05$) and over MVS-GCN on ACC, AUC, PRE (all $p < 0.05$), and SPE ($p < 0.01$). On the UM site, ABFR-KAN demonstrates statistically significant improvement over RandomFR in terms of ACC ($p < 0.01$) and SPE ($p < 0.05$), over KD-Transformer on ACC, F1, PRE, and SPE (all $p < 0.01$), and over MVS-GCN in terms of ACC, AUC (both $p < 0.01$), PRE, and SPE (both $p < 0.05$). This trend of ABFR-KAN offering statistically significant improvements continues for the poorly performing baselines across both ABIDE I sites. On the KKI site from ABIDE II, ABFR-KAN demonstrated statistically significant improvements on all metrics but SPE for the non-RandomFR baselines.

4.3.2 Cross-Site Generalizability

Table 2 shows the cross-site generalizability of ABFR-KAN compared to baseline methods. We explore cross-site generalizability in two directions: first, we examine how well models trained on the NYU site generalize when tested on data from the UM site; then we explore the opposite, where models trained on UM are tested on NYU. From the results, it is evident that despite an expected drop in performance compared to the in-domain cross-validation experiments, ABFR-KAN is the superior model in terms of generalizability across both settings compared to the other baselines, albeit they are more competitive in this setting than in the single-site experiments.

In both settings, ABFR-KAN achieves the strongest performance of the compared models on three of the six metrics. In the NYU $\rightarrow$ UM setup, ABFR-KAN achieves

the highest values across AUC, F1, and PRE. On UM $\rightarrow$ NYU, ABFR-KAN achieves the best performance on ACC, AUC, and PRE. As with the discussion of the single-site performance of the models, a caveat needs to be added when viewing the results of both GCN and BrainGNN on the NYU $\rightarrow$ UM experiment, as well as BrainNetCNN and GCN on UM $\rightarrow$ NYU. In both instances, the relationship between the reported SEN and SPE values for these models indicates issues in terms of clinical viability, as there is a high indication of both false positives (BrainGNN on NYU $\rightarrow$ UM, GCN on UM $\rightarrow$ NYU) and missed positives (GCN on NYU $\rightarrow$ UM, BrainNetCNN on UM $\rightarrow$ NYU). ABFR-KAN, in comparison, achieves a much better balance between R and SPE in both settings, indicating better clinical viability.

However, it should be noted that while ABFR-KAN achieved the strongest AUC in the NYU $\rightarrow$ UM cross-site setting (0.7181), the corresponding classification accuracy remained relatively low at 0.5517, which was below the RandomFR baseline (0.6109). This discrepancy between AUC and ACC indicates that the model retains meaningful discriminative capability in terms of ranking ASD versus control subjects, but that the learned decision boundary does not transfer robustly across acquisition sites. Since AUC evaluates separability across all possible thresholds whereas ACC depends on a fixed threshold of 0.5, these results suggest that ABFR-KAN produces informative probability estimates under domain shift, but that calibration deteriorates when evaluated on unseen sites.

This behavior, along with the general asymmetry of the cross-site results, with models trained on NYU generalizing differently to UM compared to the reverse direction, is to be expected due to the disparate data distributions between the two sites. As described in Section 4.1, the NYU and

Table 5: Impact of various KAN variants within the ABFR-KAN framework. Mean $\pm$ stdev scores are reported by averaging the scores of 5-fold cross-validation. Best and second-best results are **bolded** and underlined, respectively. ($\dagger$) indicates KAN variants that improved upon the ACC of the baseline in a statistically significant manner, and ($\ddagger$) denotes the KAN variant that most improved upon the ACC of the baseline in a statistically significant manner.

NYU Site Classification Performance						
KAN Variant	ACC	AUC	F1	PRE	SEN	SPE
Baseline (MLP)	0.7079 $\pm$ 0.0388	0.6544 $\pm$ 0.0583	0.7734 $\pm$ 0.0379	0.6986 $\pm$ 0.0435	0.8784 $\pm$ 0.1024	0.6783 $\pm$ 0.0486
Efficient-KAN	0.7079 $\pm$ 0.0341	0.6286 $\pm$ 0.0826	0.7701 $\pm$ 0.0525	0.7022 $\pm$ 0.0490	0.8795 $\pm$ 0.1501	0.6793 $\pm$ 0.0352
FastKAN ($\ddagger$)	0.7427 $\pm$ 0.0342	0.7041 $\pm$ 0.0708	0.8003 $\pm$ 0.0200	<u>0.7267 $\pm$ 0.0469</u>	<u>0.8984 $\pm$ 0.0646</u>	0.7159 $\pm$ 0.0434
FasterKAN ($\dagger$)	<u>0.7259 $\pm$ 0.0716</u>	0.6602 $\pm$ 0.0359	<u>0.7954 $\pm$ 0.0429</u>	0.7039 $\pm$ 0.0629	0.9184 $\pm$ 0.0248	0.6925 $\pm$ 0.0838
Wav-KAN	0.6842 $\pm$ 0.0218	<u>0.6672 $\pm$ 0.0331</u>	0.7078 $\pm$ 0.0290	0.7585 $\pm$ 0.0524	0.6732 $\pm$ 0.0759	0.6842 $\pm$ 0.0311
ChebyKAN	0.6904 $\pm$ 0.0447	0.6611 $\pm$ 0.0493	0.7315 $\pm$ 0.0615	0.7204 $\pm$ 0.0350	0.7558 $\pm$ 0.1306	0.6793 $\pm$ 0.0363
UM Site Classification Performance						
KAN Variant	ACC	AUC	F1	PRE	SEN	SPE
Baseline (MLP)	0.7182 $\pm$ 0.0530	0.6426 $\pm$ 0.1395	0.6999 $\pm$ 0.1311	<u>0.7169 $\pm$ 0.1104</u>	0.7397 $\pm$ 0.2197	0.6984 $\pm$ 0.0780
Efficient-KAN	0.7091 $\pm$ 0.0464	<u>0.6769 $\pm$ 0.1095</u>	0.7108 $\pm$ 0.1075	0.6715 $\pm$ 0.0684	0.7688 $\pm$ 0.1637	0.6956 $\pm$ 0.0626
FastKAN	0.7091 $\pm$ 0.0464	0.6471 $\pm$ 0.0287	0.7123 $\pm$ 0.0926	0.6683 $\pm$ 0.0482	0.7745 $\pm$ 0.1641	0.6845 $\pm$ 0.0381
FasterKAN	0.7091 $\pm$ 0.0464	0.6203 $\pm$ 0.0490	0.7093 $\pm$ 0.1027	0.6813 $\pm$ 0.0761	<u>0.7859 $\pm$ 0.2114</u>	0.6777 $\pm$ 0.0278
Wav-KAN ($\ddagger$)	0.7727 $\pm$ 0.0287	0.7184 $\pm$ 0.0620	0.7682 $\pm$ 0.0482	0.7840 $\pm$ 0.1086	0.7837 $\pm$ 0.1314	0.7768 $\pm$ 0.0270
ChebyKAN ($\dagger$)	<u>0.7273 $\pm$ 0.0498</u>	0.6345 $\pm$ 0.1052	<u>0.7443 $\pm$ 0.0838</u>	0.6801 $\pm$ 0.0374	0.8412 $\pm$ 0.1742	<u>0.7061 $\pm$ 0.0794</u>
KKI Site Classification Performance						
KAN Variant	ACC	AUC	F1	PRE	SEN	SPE
Baseline (MLP)	0.7455 $\pm$ 0.0600	0.6248 $\pm$ 0.1552	0.6457 $\pm$ 0.1855	0.7076 $\pm$ 0.1774	0.6200 $\pm$ 0.2159	0.7148 $\pm$ 0.1162
Efficient-KAN	0.6727 $\pm$ 0.0445	0.5110 $\pm$ 0.1895	0.3455 $\pm$ 0.3007	0.4533 $\pm$ 0.3942	0.3300 $\pm$ 0.3219	0.5983 $\pm$ 0.0889
FastKAN	0.6909 $\pm$ 0.0445	0.5143 $\pm$ 0.2367	0.3931 $\pm$ 0.3248	0.4167 $\pm$ 0.3416	0.3800 $\pm$ 0.3250	0.6257 $\pm$ 0.1042
FasterKAN	0.7091 $\pm$ 0.0680	0.5910 $\pm$ 0.1871	0.5400 $\pm$ 0.1718	<u>0.7833 $\pm$ 0.1944</u>	0.5100 $\pm$ 0.2871	0.6812 $\pm$ 0.0974
Wav-KAN $\ddagger$	0.7636 $\pm$ 0.0445	0.6505 $\pm$ 0.0869	0.6588 $\pm$ 0.1329	0.8033 $\pm$ 0.1675	0.6300 $\pm$ 0.2015	0.7388 $\pm$ 0.0648
ChebyKAN	0.6000 $\pm$ 0.0445	0.4744 $\pm$ 0.0887	0.0000 $\pm$ 0.0000	0.0000 $\pm$ 0.0000	0.0000 $\pm$ 0.0000	0.5000 $\pm$ 0.0000

UM cohorts differ in multiple respects, including age ranges, sex distributions, cohort balance, and scanner acquisition characteristics. In particular, the NYU cohort contains a broader age range and a more imbalanced class distribution, whereas the UM cohort is more tightly age-constrained and balanced. The domain shift caused by these differences means that the properties learned from one site may not be equally transferable to the other.

Statistical Analysis: Using the same one-tailed paired t-test as before, we investigate the statistical significance of our proposed ABFR-KAN framework over existing methods. On the NYU $\rightarrow$ UM setup, ABFR-KAN achieves a statistically significant improvement over RandomFR in AUC, F1, and PRE (all $p < 0.01$). On UM $\rightarrow$ NYU, ABFR-KAN improves in a statistically significant manner over GCN in terms of AUC ($p < 0.05$) and SPE ($p < 0.01$), over GAT on ACC, PRE (both $p < 0.05$), AUC, and SPE (both $p < 0.01$), and over BrainGNN in ACC, AUC, PRE (all $p < 0.05$), and F1 ($p < 0.01$). Like before, the trend of ABFR-KAN offering statistically significant improvements continues across lesser-performing baselines in both setups.

Table 6: Performance of ABFR-KAN against several atlas-based techniques for ASD diagnosis on the NYU site. Best and second-best results are **bolded** and underlined, respectively.

Model	ACC	AUC	SEN	SPE
Spera et al. (2019)	0.6300	0.7500	0.4800	0.8300
ElNakieb et al. (2021)	0.6978	—	<u>0.7200</u>	0.6700
Yang et al. (2024)	<u>0.7294</u>	—	0.7083	<u>0.7470</u>
Ours	0.7427	<u>0.7041</u>	0.8984	0.7159

4.3.3 Comparison to Atlas-Based Techniques

In order to assess the performance of ABFR-KAN against “gold-standard” atlas-based techniques, Table 6 shows a direct comparison against several such techniques. The compared-against methods are: Spera et al. (2019), whose method utilizes the Harvard-Oxford anatomical atlas; ElNakieb et al. (2021), which uses the Johns Hopkins University white matter atlas; and Yang et al. (2024), which uses the Brodmann atlas. The results demonstrate that ABFR-KAN is capable of performing comparably to, or even outperforming the atlas-based techniques on certain metrics.

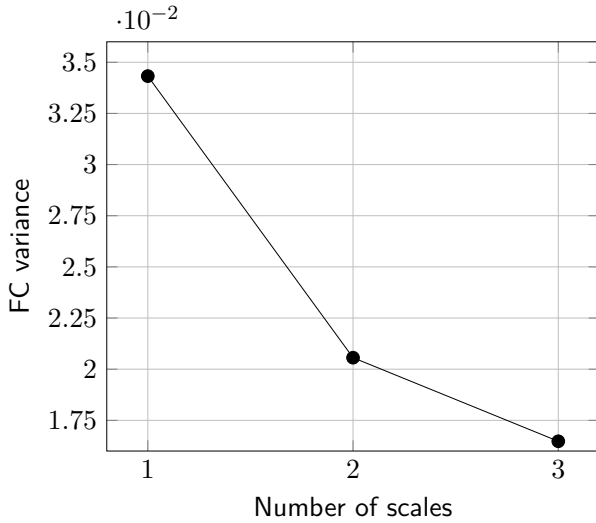


Figure 8: Impact of differing scale values in our proposed multi-scale patch sampling method on reducing sampling variance in FC estimation. Increasing the number of iterations beyond single-pass sampling (the baseline) demonstrably reduces sample variance, which offers benefits that include improved feature stability, a reduction in the risk of overfitting to sampling artifacts, and improved reproducibility between sites.

4.3.4 Ablation Experiments

In order to justify the existence of each of our proposed components within our ABFR-KAN framework, as well as to determine the best-performing KAN variants and configurations, we conduct a series of ablation experiments. Tables 3 and 4 show the impact of various classification backbones and anchor selection/patch sampling configurations within the ABFR-KAN framework, respectively. These experiments were conducted on the UM site of ABIDE I, as the same experiments on the NYU site were previously conducted and thoroughly analyzed in our previous work [Ward and Imran \(2025\)](#).

From the results shown in Table 3, it is clear that the use of ViT as the classification backbone offers the best performance, with it achieving top performance across four of the six methods. This validates our previous findings [Ward and Imran \(2025\)](#), where we determined that ViT was a better choice than DeiT. The Mamba-based method Vim achieves the best performance on the F1 and R metrics, but overall, it performs worse than ViT. The CNN-based method, EfficientNetV2, performs the worst overall, which aligns with the established knowledge that transformer-based methods often outperform CNN-based ones [Dosovitskiy et al. \(2020\)](#).

Table 4 demonstrates that our proposed architectural configuration involving the use of random anchor selection followed by multi-scale patch sampling is superior compared to the baseline grid-based anchor selection approach followed by random patch sampling. Again, this validates

our findings in our previous work [Ward and Imran \(2025\)](#). Interestingly, the results also show that simply applying randomized anchor selection while keeping the random patch sampling can perform worse than the original baseline configuration that uses neither proposed component. We hypothesize that this occurs because randomized anchor selection increases inter-subject variability and representational diversity, which, without accompanying multi-scale sampling, may lead to less stable FC estimation and reduced robustness due to the additional stochasticity introduced by randomized anchors. As for why the use of grid-based anchors and multi-scale patch sampling also exhibits worse performance compared to the baseline, we posit that grid-based anchors, while structurally biased, may provide a spatially consistent reference structure that can partially compensate for the limitations of single-pass patch sampling.

We briefly mentioned this in Section 3.2, but would now like to draw attention to Fig. 8, which both demonstrates that increasing the number of iterations reduces the FC sampling variance, which has numerous benefits, and serves as an ablation experiment on the number of iterations themselves. Table 5 similarly investigates the relative behavior of various KAN variants when incorporated into the ABFR-KAN framework, with the primary goal being comparative analysis rather than formal hyperparameter or architecture optimization. A detailed analysis of the computational complexity and efficiency of each variation is provided in Section 4.4, based on Figure 9. For brevity, we report only the results of the best-performing configuration involving these KAN variants in this table. Additional experiments that showcase the full range of possible configurations for each of these variants on both ABIDE I sites are reported in Appendix B.

From our experiments, we found that the KAN variants perform differently across the three sites. This is, perhaps, to be expected given the diversity in patient populations between the two sites (See Section 4.1) and varying acquisition parameters between them. On the NYU site, the FastKAN variant in a KAN-KAN configuration (where FastKAN blocks replaced the MLP components in both the ViT encoder and the classification head) came out on top, outperforming the second-best method (FasterKAN, also in a KAN-KAN configuration) by 1.41%, on average. On the UM site, Wav-KAN (in an MLP-KAN configuration where only the MLP classification head was replaced by Wav-KAN) was clearly superior to the other KAN variants, achieving a 4.51% improvement on average over the second-best variant, ChebyKAN. On KKI, Wav-KAN again demonstrated supremacy, outperforming the second-best FasterKAN by 7.18% on average.

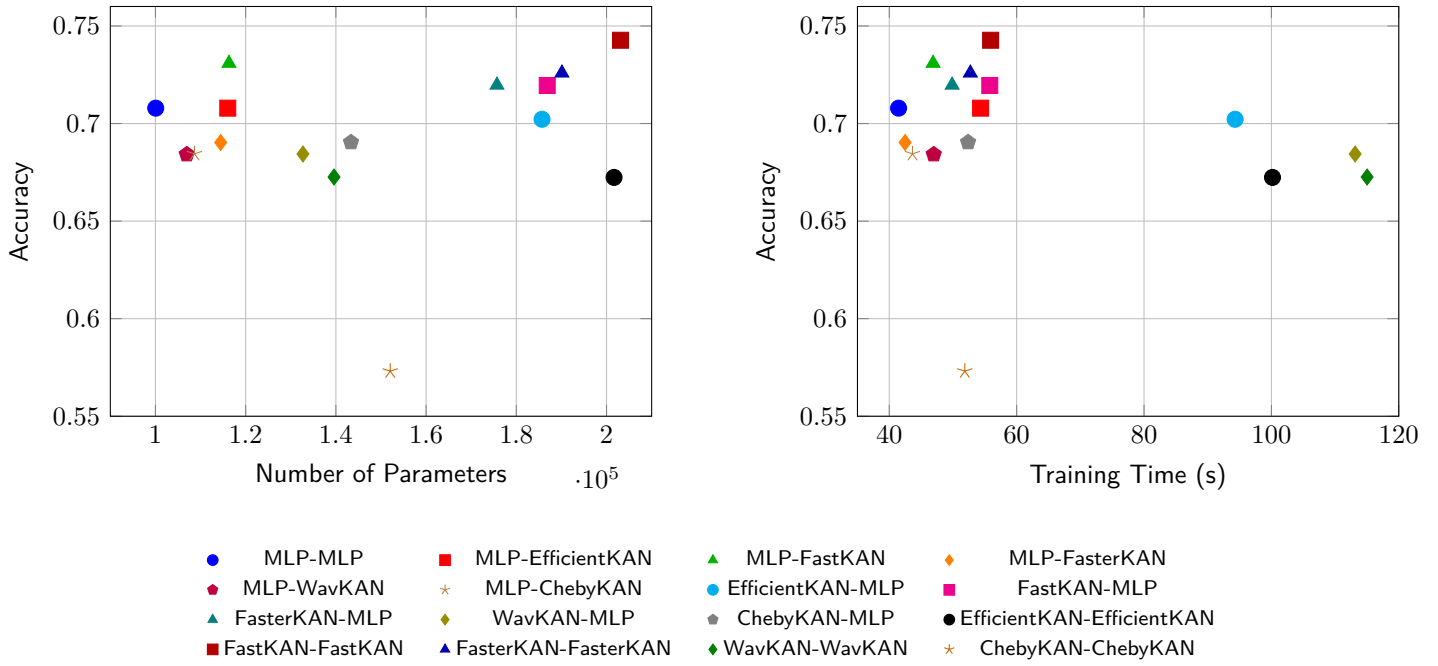


Figure 9: Trade-off analysis between classification performance on the NYU site, model complexity, and computational cost across different KAN architectural configurations. Left: Accuracy versus number of trainable parameters. Right: Accuracy versus total training time across all five folds of 5-fold cross-validation.

4.4 Discussion

Reducing the Reliance on Atlases: A core contribution of our work is the removal of reliance on predefined atlases through randomized anchor selection. Unlike the baseline grid-based approach, which imposes a rigid spatial arrangement, our method randomly selects anchor patches from valid GM regions, resulting in improved anatomical conformity and reduced structural bias, as demonstrated in Figs. 4 and 3. Importantly, these anchors are sampled once at the dataset level and are subsequently fixed and shared across all subjects, providing a common reference space for functional connectivity computation. Thus, the individual-specificity of the resulting representations does not arise from anchor selection itself, but from the independently sampled local brain patches used to construct each subject’s functional representation.

The use of randomized rather than grid-based anchors may nevertheless influence downstream performance by altering the spatial reference against which subject-specific patches are characterized. Our ablation experiments suggest that randomized anchor selection is most effective when paired with multi-scale patch sampling. In isolation, randomized anchors can introduce additional variability into the functional representation and may reduce stability, whereas multi-scale sampling helps mitigate this effect by aggregating multiple stochastic estimates of functional connectivity. This interaction may explain why the full combination of randomized anchor selection and multi-scale patch sampling performs better than configurations using

either component alone. Overall, randomized anchor selection should therefore be viewed primarily as a mechanism for reducing atlas-related structural bias and improving anatomical conformity, while subject-level patch sampling provides the source of individual-specific functional variability.

Complexity and Efficiency of the Explored KAN Variants: Fig. 9 shows the trade-off in training time and number of parameters for the five different KAN variants we tested in this work (Efficient-KAN, FastKAN, FasterKAN, Wav-KAN, ChebyKAN). We show this trade-off across each of the three architectural configurations we investigated (MLP-KAN, KAN-MLP, KAN-KAN), and the baseline (MLP-MLP). As expected, the baseline method using only MLPs is both the fastest to train (41.46 seconds in total), and uses the smallest number of parameters (100,066). This is consistent with known characteristics of KANs, with them often increasing parameter counts with the trade-off of improved performance Wang et al. (2024b).

Despite the increase in model complexity and decrease in efficiency, we argue that the performance trade-off is well justified by the results reported in Tables 1 and 2, especially. Take, for instance, Wav-KAN in an MLP-KAN configuration (denoted by the red pentagon in Fig. 9). This configuration offered the best performance on the single-site UM experiment, improving on average over the baseline method by 6.56%, while adding only 6,882 parameters and 5.52 seconds to the total training time. However, if training complexity and efficiency are top priorities, practitioners must carefully weigh whether KAN integration is worth the

extra overhead.

The Role of KANs: Replacing traditional MLP components with KAN variants plays a key role in the performance gains offered by ABFR-KAN. KANs offer increased expressiveness through learnable univariate functions while reducing over-parameterization relative to dense MLPs operating on high-dimensional FC features. The ablation studies demonstrate that KAN-based networks consistently outperform MLP-only baselines, with performance boosts depending on both the KAN variant and its placement within the architecture. Interestingly, no single KAN variant dominates across all experimental settings. FastKAN performs best on the NYU site, while Wav-KAN achieves superior results on the UM site, suggesting that the optimal basis function may be data- and domain-dependent. This observation aligns with recent findings that KAN design choices remain task-specific [Becerra-Suarez et al. \(2025\)](#), and it motivates future work on adaptive or hybrid KAN configurations for neuroimaging applications.

Clinical Implications: Cross-site experiments reveal that ABFR-KAN generalizes better than existing methods when trained and tested across different imaging sites. From a clinical perspective, the improved balance between sensitivity and specificity is particularly important. Several baseline methods achieve high SEN at the cost of excessive false positives, limiting their practical utility. ABFR-KAN avoids this issue, indicating a stronger potential for real-world diagnostic support.

Limitations and Future Directions: Despite the promise of our ABFR-KAN method, there are several limitations to our work. *First*, we focus exclusively on rs-fMRI data from the ABIDE repository. While cross-site evaluation partially demonstrates the generalizability of our approach, in the future, this should be extended to validate additional datasets, disorders, and imaging modalities. *Second*, compared to the prevalence of ASD in the real-world population, the datasets used in this research are fairly balanced, leaving questions about how well the model would perform on a real population dataset. A potential solution to dealing with real, imbalanced datasets lies in the use of loss functions designed to handle class imbalance, such as Focal loss [Lin et al. \(2017\)](#). *Third*, although KANs offer improved interpretability in principle, we don't explicitly analyze the learned activation functions. Investigating whether these functions correspond to neurobiologically meaningful patterns would be an interesting avenue for future research. *Finally*, the KAN variant and configuration selection were performed empirically on a per-site basis. Although this enabled exploration of the relative effectiveness of different KAN formulations under varying data characteristics, the selection process was not fully isolated from the evaluation protocol. As a result, the reported performance estimates

may partially reflect site-specific optimization and could therefore overestimate generalization performance. Future work should employ a strictly nested cross-validation framework, where architecture and hyperparameter selection are performed exclusively within inner validation folds while keeping the outer test folds completely unseen, in order to eliminate potential information leakage and provide a more rigorous estimate of model generalizability.

5. Conclusions

In this work, we introduced ABFR-KAN, a novel framework for functional brain analysis that addresses the limitations of atlas-based parcellations in defining ROIs. By combining randomized, anatomy-aware anchor selection with multi-scale patch sampling, our approach reduces structural bias, improves GM coverage, and yields more robust and individualized functional representations. To make this approach useful for downstream diagnosis, we follow this data pipeline with a ViT-style classification network. In this network, we demonstrate that replacing traditional MLP components in the encoder and classification head with KANs leads to consistent and statistically significant performance gains for ASD diagnosis.

Extensive experiments on the ABIDE I and II datasets show that ABFR-KAN outperforms strong SOTA baselines in both single-site and cross-site evaluations. Ablation studies further confirm the necessity of combining randomized anchor selection with multi-scale patch sampling and KAN-based modeling, while also highlighting the importance of selecting the optimal classification backbone and KAN variant/configuration for the task at hand. Future work will explore extending ABFR-KAN to additional neurological and psychiatric disorders, as well as expanding experimentation to different datasets.

Ethical Standards

The work follows appropriate ethical standards in conducting research and writing the manuscript, following all applicable laws and regulations regarding the treatment of animals or human subjects.

Conflicts of Interest

We declare that we have no conflicts of interest.

Data availability

For this work, we trained/evaluated our models on publicly available rs-fMRI data from the ABIDE Preprocessed

repository [Craddock et al. \(2013\)](#). These data can be downloaded following the instructions at <http://preprocessed-connectomes-project.org/abide/download.html>. The code, along with pre-trained checkpoints and data samples, is available in our public GitHub repository, linked at the end of our abstract.

References

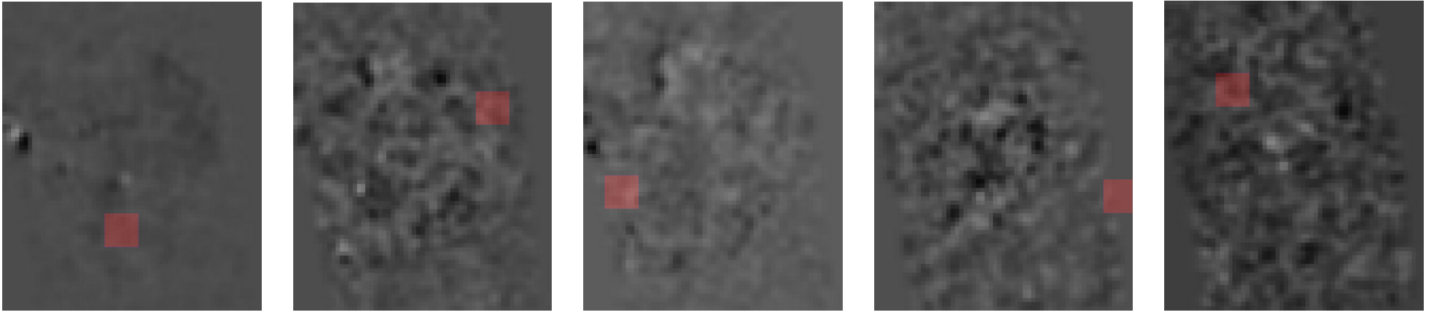
- American Psychiatric Association. *Diagnostic and Statistical Manual of Mental Disorders: DSM-IV-TR*. American Psychiatric Association, Washington, DC, 2000.
- Basim Azam and Naveed Akhtar. Suitability of KANs for computer vision: A preliminary investigation. *arXiv preprint arXiv:2406.09087*, 2024.
- Fray L Becerra-Suarez, Ana G Borrero-Ramírez, Edwin Valencia-Castillo, and Manuel G Forero. Mathematical generalization of Kolmogorov-Arnold networks (KAN) and their variants. *Mathematics*, 13(19):3128, 2025.
- Blealtan and Akaash Dash. Efficient-KAN: An efficient implementation of Kolmogorov-Arnold network. <https://github.com/Blealtan/efficient-kan>, 2024.
- Zavareh Bozorgasl and Hao Chen. Wav-KAN: Wavelet Kolmogorov-Arnold networks. *arXiv preprint arXiv:2405.12832*, 2024.
- Ziwen Chen, Saaketh Koundinya, and Wu Di. Vision-KAN: Exploring the possibility of KAN replacing MLP in vision transformer. <https://github.com/chenziwenhaoshuai/Vision-KAN.git>, 2024.
- Minjong Cheon. Demonstrating the efficacy of Kolmogorov-Arnold networks in vision tasks. *arXiv preprint arXiv:2406.14916*, 2024.
- Corinna Cortes and Vladimir Vapnik. Support-vector networks. *Machine Learning*, 20(3):273–297, 1995.
- Cameron Craddock, Yassine Benhajali, Carlton Chu, Francois Chouinard, Alan Evans, András Jakab, Budhachandra Singh Khundrakpam, John David Lewis, Qingyang Li, Michael Milham, et al. The Neuro Bureau Preprocessing Initiative: Open sharing of preprocessed neuroimaging data and derivatives. *Front. Neurosci.*, 7(27):5, 2013.
- R Cameron Craddock, G Andrew James, Paul E Holtzheimer III, Xiaoping P Hu, and Helen S Mayberg. A whole brain fMRI atlas generated via spatially constrained spectral clustering. *Hum. Brain Mapp.*, 33(8):1914–1928, 2012.
- Weigang Cui, Yulan Ma, Jianxun Ren, Jingyu Liu, Guolin Ma, Hesheng Liu, and Yang Li. Personalized functional connectivity based spatio-temporal aggregated attention network for MCI identification. *IEEE Trans. Neural Syst. Rehabil. Eng.*, 31:2257–2267, 2023.
- Athanasios Delis. FasterKAN. <https://github.com/AthanasiosDelis/faster-kan/>, 2024.
- Bo Deng. Conceptual circuit models of neurons. *J. Integr. Neurosci.*, 8(03):255–297, 2009.
- Rahul S Desikan, Florent Ségonne, Bruce Fischl, Brian T Quinn, Bradford C Dickerson, Deborah Blacker, Randy L Buckner, Anders M Dale, R Paul Maguire, Bradley T Hyman, et al. An automated labeling system for subdividing the human cerebral cortex on MRI scans into gyral based regions of interest. *NeuroImage*, 31(3):968–980, 2006.
- Adriana Di Martino, Chao-Gan Yan, Qingyang Li, Erin Denio, Francisco X Castellanos, Kaat Alaerts, Jeffrey S Anderson, Michal Assaf, Susan Y Bookheimer, Mirella Dapretto, et al. The Autism Brain Imaging Data Exchange: Towards a large-scale evaluation of the intrinsic brain architecture in autism. *Mol. Psychiatry*, 19(6):659–667, 2014.
- David Alexander Dickie, Susan D Shenkin, Devasuda Anblagan, Juyoung Lee, Manuel Blesa Cabez, David Rodriguez, James P Boardman, Adam Waldman, Dominic E Job, and Joanna M Wardlaw. Whole brain magnetic resonance image atlases: A systematic review of existing atlases and caveats for use in population imaging. *Frontiers in Neuroinformatics*, 11:1, 2017.
- Tianqi Ding, Dawei Xiang, Keith E Schubert, and Liang Dong. Explainable diagnosis of Alzheimer’s disease using graph Kolmogorov-Arnold networks. In *Southwest Data Science Conference*, pages 3–15. Springer, 2025.
- Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *International Conference on Learning Representations*, 2020.
- Yuhui Du, Zening Fu, and Vince D Calhoun. Classification and prediction of brain disorders using functional connectivity: Promising but challenging. *Front. Neurosci.*, 12:525, 2018.
- Yaser ElNakieb, Mohamed T Ali, Ahmed Soliman, Ali Mahmoud, Ahmed Shalaby, Andrew Switala, Mohammed Ghazal, Ashraf Khalil, Luay Fraiwan, Gregory Barnes, et al. Identifying brain pathological abnormalities of

- autism for classification using diffusion tensor imaging. In *Neural Engineering Techniques for Autism Spectrum Disorder*, pages 361–376. Elsevier, 2021.
- Alessandro Giupponi, Davide De Crescenzo, Marco Pinamonti, Manuela Moretto, Alessandra Bertoldo, Mattia Veronese, and Marco Castellaro. Accurate brain age prediction from MRI: Evaluating Kolmogorov-Arnold and convolutional networks. In *Medical Imaging with Deep Learning-Short Papers*, 2025.
- David Guo, Ismail Khalfaoui-Hassani, and iisak. ChebyKAN. <https://github.com/SynodicMonth/ChebyKAN/>, 2024.
- Maria Hakonen, Louisa Dahmani, Kaisu Lankinen, Jianxun Ren, Julianna Barbaro, Anna Blazejewska, Weigang Cui, Parker Kotlarz, Meiling Li, Jonathan R Polimeni, et al. Individual connectivity-based parcellations reflect functional properties of human auditory cortex. *Imaging Neurosci.*, 3:imag.a_00486, 2025.
- Kangfu Han, Dan Hu, Jiale Cheng, Tianming Liu, Andrea Bozoki, Dajiang Zhu, and Gang Li. Dual multi-atlas representation alignment for brain disorder diagnosis using morphological connectome. In *2025 IEEE 22nd International Symposium on Biomedical Imaging (ISBI)*, pages 1–4. IEEE, 2025.
- Mohsin Hasan, Xufeng Zhao, Wenjuan Wu, Jiafei Dai, Xudong Gu, and Asia Noreen. Long short-term memory and Kolmogorov Arnold Network theorem for epileptic seizure prediction. *Engineering Applications of Artificial Intelligence*, 154:110757, 2025.
- Tasnimul Hasan, ABM Aowlad Hossain, Shuvagata Biswas, and Farhana Tazrin. HandEEGnet: A Kolmogorov-Arnold network for hand movements classification from motor imagery EEG. In *2024 IEEE 3rd International Conference on Robotics, Automation, Artificial-Intelligence and Internet-of-Things (RAAICON)*, pages 234–237. IEEE, 2024.
- Robert Hecht-Nielsen. Kolmogorov's mapping neural network existence theorem. In *Proceedings of the International Conference on Neural Networks*, volume 3, pages 11–14. IEEE Press New York, NY, USA, 1987.
- Colin J Holmes, Rick Hoge, Louis Collins, Roger Woods, Arthur W Toga, and Alan C Evans. Enhancement of MR images using registration for signal averaging. *Journal of Computer Assisted Tomography*, 22(2):324–333, 1998.
- Jiacheng Hou, Zhenjie Song, and Ercan Engin Kuruoglu. BrainNetMLP: An efficient and effective baseline for functional brain network classification. In *1st MICCAI Workshop on Efficient Medical AI*, 2025.
- Marcus Kaiser. A tutorial in connectome analysis: Topological and spatial features of brain networks. *NeuroImage*, 57(3):892–907, 2011.
- Jeremy Kawahara, Colin J Brown, Steven P Miller, Brian G Booth, Vann Chau, Ruth E Grunau, Jill G Zwicker, and Ghassan Hamarneh. BrainNetCNN: Convolutional neural networks for brain networks; towards predicting neurodevelopment. *NeuroImage*, 146:1038–1049, 2017.
- David N Kennedy, Nicholas Lange, Nikos Makris, Julianna Bates, James Meyer, and Verne S Caviness Jr. Gyri of the human neocortex: An MRI-based analysis of volume and variance. *Cereb. Cortex*, 8(4):372–384, 1998.
- Sahil Kumar. Early disease detection using AI: A deep learning approach to predicting cancer and neurological disorders. *Int. J. Sci. Res. Manag.*, 13(04):2136–2155, 2025.
- Deok-Joong Lee, Dong-Hee Shin, Young-Han Son, Ji-Wung Han, Ji-Hye Oh, Da-Hyun Kim, Ji-Hoon Jeong, and Tae-Eui Kam. Spectral graph neural network-based multi-atlas brain network fusion for major depressive disorder diagnosis. *IEEE J. Biomed. Health Informat.*, 2024.
- Karen Q Leow, Mary A Tonta, Jing Lu, Harold A Coleman, and Helena C Parkington. Towards understanding sex differences in autism spectrum disorders. *Brain Research*, 1833:148877, 2024.
- Chenxin Li, Xinyu Liu, Wuyang Li, Cheng Wang, Hengyu Liu, Yifan Liu, Zhen Chen, and Yixuan Yuan. U-KAN makes strong backbone for medical image segmentation and generation. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 4652–4660, 2025.
- Xiaoxiao Li, Yuan Zhou, Nicha Dvornek, Muhan Zhang, Siyuan Gao, Juntang Zhuang, Dustin Scheinost, Lawrence H Staib, Pamela Ventola, and James S Duncan. BrainGNN: Interpretable brain graph neural network for fMRI analysis. *Med. Image Anal.*, 74:102233, 2021.
- Ziyao Li. Kolmogorov-Arnold networks are radial basis function networks. *arXiv preprint arXiv:2405.06721*, 2024.
- Chunfeng Lian, Mingxia Liu, Yongsheng Pan, and Dinggang Shen. Attention-guided hybrid network for dementia diagnosis with structural MR images. *IEEE Trans. Cybern.*, 52(4):1992–2003, 2020.
- Tsung-Yi Lin, Priya Goyal, Ross Girshick, Kaiming He, and Piotr Dollár. Focal loss for dense object detection. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 2980–2988, 2017.

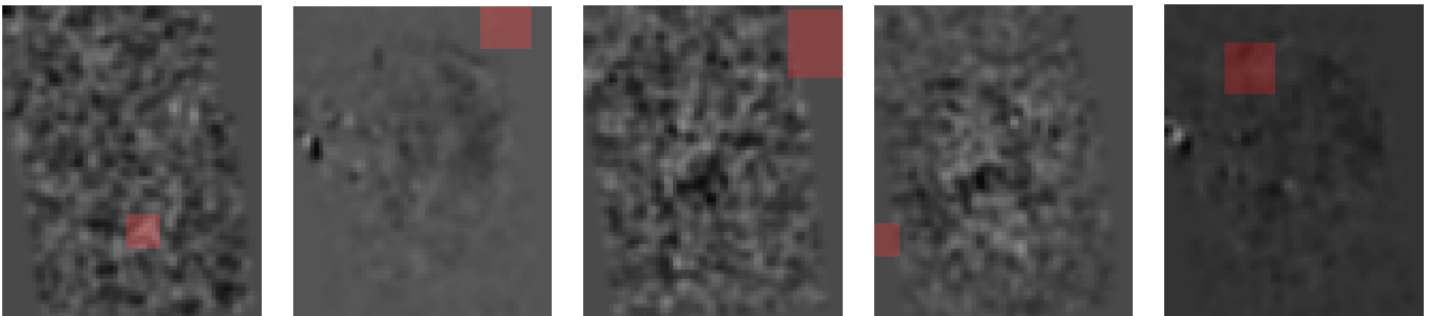
- Mengjun Liu, Huifeng Zhang, Mianxin Liu, Dongdong Chen, Zixu Zhuang, Xin Wang, Lichi Zhang, Daihui Peng, and Qian Wang. Randomizing human brain function representation for brain disease diagnosis. *IEEE Trans. Med. Imaging*, 2024a.
- Ziming Liu, Pingchuan Ma, Yixuan Wang, Wojciech Matusik, and Max Tegmark. KAN 2.0: Kolmogorov-Arnold networks meet science. *arXiv preprint arXiv:2408.10205*, 2024b.
- Ziming Liu, Yixuan Wang, Sachin Vaidya, Fabian Ruehle, James Halverson, Marin Soljačić, Thomas Y Hou, and Max Tegmark. KAN: Kolmogorov-Arnold networks. In *The Thirteenth International Conference on Learning Representations*, 2024c.
- Catherine Lord, Susan Risi, Linda Lambrecht, Edwin H Cook Jr, Bennett L Leventhal, Pamela C DiLavore, Andrew Pickles, and Michael Rutter. The Autism Diagnostic Observation Schedule—Generic: A standard measure of social and communication deficits associated with the spectrum of autism. *Journal of Autism and Developmental Disorders*, 30(3):205–223, 2000.
- Maksim Penkin and Andrey Krylov. Kolmogorov-Arnold networks as deep feature extractors for MRI reconstruction. In *Proceedings of the 2024 9th International Conference on Biomedical Imaging, Signal Processing*, pages 92–97, 2024.
- Kun Qin, Du Lei, Walter HL Pinaya, Nanfang Pan, Wenbin Li, Ziyu Zhu, John A Sweeney, Andrea Mechelli, and Qiyong Gong. Using graph convolutional network to characterize individuals with major depressive disorder across multiple imaging sites. *EBioMedicine*, 78, 2022.
- Gang Qu, Wenxing Hu, Li Xiao, Junqi Wang, Yuntong Bai, Beenish Patel, Kun Zhang, and Yu-Ping Wang. Brain functional connectivity analysis via graphical deep learning. *IEEE Trans. Biomed. Eng.*, 69(5):1696–1706, 2021.
- William D Reeves, Ishfaq Ahmed, Brooke S Jackson, Wenwu Sun, Celestine F Williams, Catherine L Davis, Jennifer E McDowell, Nathan E Yanasak, Shaoyong Su, and Qun Zhao. fMRI-based data-driven brain parcellation using independent component analysis. *J. Neurosci. Methods*, 417:110403, 2025.
- Edmund T Rolls, Chu-Chung Huang, Ching-Po Lin, Jianfeng Feng, and Marc Joliot. Automated anatomical labelling atlas 3. *NeuroImage*, 206:116189, 2020.
- Michael Rutter, Ann LeCouteur, and Catherine Lord. Autism Diagnostic Interview-Revised. *American Journal on Mental Retardation*, 2003.
- Jainy Sachdeva, Riyaansh Mittal, Jiya Mehta, Riya Jain, and Anmol Ranjan. Resolving autism spectrum disorder (ASD) through brain topologies using fMRI dataset with multi-layer perceptron (MLP). *Psychiatry Res. Neuroimaging*, 343:111858, 2024.
- Johannes Schmidt-Hieber. The Kolmogorov–Arnold representation theorem revisited. *Neural Netw.*, 137:119–126, 2021.
- Seyd Teymoor Seydi, Zavareh Bozorgasl, and Hao Chen. Unveiling the power of wavelets: A wavelet-based Kolmogorov-Arnold network for hyperspectral image classification. *arXiv preprint arXiv:2406.07869*, 2024.
- S. S. Sidharth. RBF-KAN. <https://github.com/sidhu2690/RBF-KAN/>, 2024.
- Jason Smucny, Ge Shi, and Ian Davidson. Deep learning in neuroimaging: Overcoming challenges with emerging approaches. *Front. Psychiatry*, 13:912600, 2022.
- Giovanna Spera, Alessandra Retico, Paolo Bosco, Elisa Ferrari, Letizia Palumbo, Piernicola Oliva, Filippo Muratori, and Sara Calderoni. Evaluation of altered functional connections in male children with autism spectrum disorders on multiple-site data optimized with machine learning. *Frontiers in Psychiatry*, 10:620, 2019.
- Dimitris Stathakis. How many hidden layers and nodes? *Int. J. Remote Sens.*, 30(8):2133–2147, 2009.
- Hoang-Thang Ta. BSRBF-KAN: A combination of B-splines and radial basis functions in Kolmogorov-Arnold networks. In *International Symposium on Information and Communication Technology*, pages 3–15. Springer, 2024.
- Monireh Taimouri and Vikram Ravindra. Characterizing changes to individual-specific brain signature with age. *Frontiers in Aging Neuroscience*, 17:1493855, 2025.
- Mingxing Tan and Quoc Le. EfficientNetV2: Smaller models and faster training. In *International Conference on Machine Learning*, pages 10096–10106. PMLR, 2021.
- Shawn Tan, Khe Chai Sim, and Mark Gales. Improving the interpretability of deep neural networks with stimulated learning. In *2015 IEEE Workshop on Automatic Speech Recognition and Understanding (ASRU)*, pages 617–623. IEEE, 2015.
- Aung Myo Thant and Thap Panitanarak. Emotion recognition through advanced signal fusion and Kolmogorov-Arnold networks. *IEEE Access*, 2025.
- Hugo Touvron, Matthieu Cord, Matthijs Douze, Francisco Massa, Alexandre Sablayrolles, and Hervé Jégou. Training

- data-efficient image transformers & distillation through attention. In *International Conference on Machine Learning*, pages 10347–10357. PMLR, 2021.
- Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Lio, and Yoshua Bengio. Graph attention networks. In *International Conference on Learning Representations*, 2018.
- Akshant Verma, Anshita Verma, and Ram Prakash Sharma. VGG-KAN: A hybrid approach for Alzheimer’s disease diagnosis using Kolmogorov Arnold network. *Authorea Preprints*, 2025.
- Diego Vidaurre. A new model for simultaneous dimensionality reduction and time-varying functional connectivity estimation. *PLoS Comput. Biol.*, 17(4):e1008580, 2021.
- Danny JJ Wang, Kay Jann, Chang Fan, Yang Qiao, Yu-Feng Zang, Hanbing Lu, and Yihong Yang. Neurophysiological basis of multi-scale entropy of brain complexity and its relationship with functional connectivity. *Front. Neurosci.*, 12:352, 2018.
- Guangshuai Wang, Tengfei Gao, Zhiyi Yang, Kai Zhang, Jingying Chen, and Dan Chen. A pilot study on detecting cognitive load using Kolmogorov-Arnold networks with EEG signals. In *2024 International Conference on Intelligent Education and Intelligent Research (IEIR)*, pages 1–8. IEEE, 2024a.
- Jiawen Wang, Pei Cai, Ziyang Wang, Huabin Zhang, and Jianpan Huang. CEST MRI data analysis using Kolmogorov-Arnold network (KAN) and Lorentzian-KAN (LKAN) models. *Magnetic Resonance in Medicine*, 2025.
- Yixuan Wang, Jonathan W Siegel, Ziming Liu, and Thomas Y Hou. On the expressiveness and spectral bias of KANs. *arXiv preprint arXiv:2410.01803*, 2024b.
- Yuheng Wang. Application of Kolmogorov-Arnold networks combined with graph convolutional networks in Alzheimer’s disease diagnosis. In *2025 5th International Conference on Consumer Electronics and Computer Engineering (ICCECE)*, pages 204–207. IEEE, 2025.
- Tyler Ward and Abdullah-Al-Zubaer Imran. Improving brain disorder diagnosis with advanced brain function representation and Kolmogorov-Arnold networks. In *Medical Imaging with Deep Learning*, 2025.
- Guangqi Wen, Peng Cao, Huiwen Bao, Wenju Yang, Tong Zheng, and Osmar Zaiane. MVS-GCN: A prior brain structure learning-guided multi-view graph convolution network for autism spectrum disorder diagnosis. *Computers in Biology and Medicine*, 142:105239, 2022.
- Chaogan Yan and Yufeng Zang. DPARSF: A MATLAB toolbox for "pipeline" data analysis of resting-state fMRI. *Front. Syst. Neurosci.*, 4:1377, 2010.
- Yang Yang, Detao Tang, Zhiwei Wang, Yifei Liu, Fulong Chen, Biao Jie, Tianjiao Ni, Chenglong Xu, Jintao Li, and Chao Wang. Identification of high-functioning autism spectrum disorders based on gray-white matter functional network connectivity. *Journal of Psychiatric Research*, 178:107–113, 2024.
- Jianjia Zhang, Luping Zhou, Lei Wang, Mengting Liu, and Dinggang Shen. Diffusion kernel attention network for brain disorder classification. *IEEE Transactions on Medical Imaging*, 41(10):2814–2827, 2022.
- Lianghui Zhu, Bencheng Liao, Qian Zhang, Xinlong Wang, Wenyu Liu, and Xinggang Wang. Vision Mamba: Efficient visual representation learning with bidirectional state space model. In *Proceedings of the 41st International Conference on Machine Learning*, 2024.

Appendix A. Multi-Scale Patch Sampling Visualization



(a) Random patch sampling



(b) Multi-scale patch sampling

Figure 10: (a) The random patch sampling process. Please note how the size of the patches is consistent. (b) The multi-scale patch sampling process, where each subject is processed three times as a form of data augmentation, with patch sizes varying from 8×8 , 12×12 , and 16×16 .

Appendix B. Additional Experiments on KAN Variants and Configurations

Table 7: Impact of various KAN variants and configurations within the ABFR-KAN framework on the NYU site. Mean $\pm$ stdev scores are reported by averaging the results of 5-fold cross-validation. Best and second-best results are **bolded** and underlined, respectively.

Baseline						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-MLP	0.7079 $\pm$ 0.0388	0.6544 $\pm$ 0.0583	0.7734 $\pm$ 0.0379	0.6986 $\pm$ 0.0435	0.8784 $\pm$ 0.1024	0.6783 $\pm$ 0.0486
Efficient-KAN						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-KAN	0.7079 $\pm$ 0.0341	0.6286 $\pm$ 0.0826	0.7701 $\pm$ 0.0525	0.7022 $\pm$ 0.0490	0.8795 $\pm$ 0.1501	0.6793 $\pm$ 0.0352
KAN-MLP	0.7022 $\pm$ 0.0433	0.6491 $\pm$ 0.0510	0.7665 $\pm$ 0.0516	0.6911 $\pm$ 0.0172	0.8689 $\pm$ 0.1115	0.6735 $\pm$ 0.0383
KAN-KAN	0.6724 $\pm$ 0.0573	<u>0.6340 $\pm$ 0.0417</u>	0.7363 $\pm$ 0.0223	0.7051 $\pm$ 0.0824	0.7958 $\pm$ 0.1146	0.6488 $\pm$ 0.0882
FastKAN						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-KAN	<u>0.7309 $\pm$ 0.0394</u>	0.6816 $\pm$ 0.0623	<u>0.7807 $\pm$ 0.0303</u>	0.7367 $\pm$ 0.0483	<u>0.8363 $\pm$ 0.0614</u>	<u>0.7120 $\pm$ 0.0499</u>
KAN-MLP	0.7195 $\pm$ 0.0424	<u>0.6837 $\pm$ 0.0481</u>	0.7679 $\pm$ 0.0423	<u>0.7328 $\pm$ 0.0504</u>	0.8179 $\pm$ 0.0968	0.7037 $\pm$ 0.0456
KAN-KAN	0.7427 $\pm$ 0.0342	0.7041 $\pm$ 0.0708	0.8003 $\pm$ 0.0200	0.7267 $\pm$ 0.0469	0.8984 $\pm$ 0.0646	0.7159 $\pm$ 0.0434
FasterKAN						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-KAN	0.6903 $\pm$ 0.0380	0.6355 $\pm$ 0.0689	0.7541 $\pm$ 0.0499	0.6929 $\pm$ 0.0472	0.8479 $\pm$ 0.1300	0.6630 $\pm$ 0.0398
KAN-MLP	<u>0.7197 $\pm$ 0.0364</u>	<u>0.6380 $\pm$ 0.0562</u>	<u>0.7735 $\pm$ 0.0546</u>	0.7154 $\pm$ 0.0283	<u>0.8589 $\pm$ 0.1312</u>	0.6961 $\pm$ 0.0272
KAN-KAN	0.7259 $\pm$ 0.0716	0.6602 $\pm$ 0.0359	0.7954 $\pm$ 0.0429	<u>0.7039 $\pm$ 0.0629</u>	0.9184 $\pm$ 0.0248	<u>0.6925 $\pm$ 0.0838</u>
Wav-KAN						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-KAN	<u>0.6842 $\pm$ 0.0218</u>	0.6672 $\pm$ 0.0331	<u>0.7078 $\pm$ 0.0290</u>	0.7585 $\pm$ 0.0524	<u>0.6732 $\pm$ 0.0759</u>	0.6842 $\pm$ 0.0311
KAN-MLP	0.6844 $\pm$ 0.0265	0.6191 $\pm$ 0.0999	0.7349 $\pm$ 0.0352	0.7177 $\pm$ 0.0612	0.7774 $\pm$ 0.1270	0.6720 $\pm$ 0.0364
KAN-KAN	0.6726 $\pm$ 0.0201	<u>0.6466 $\pm$ 0.0185</u>	0.6933 $\pm$ 0.0542	<u>0.7540 $\pm$ 0.0750</u>	0.6637 $\pm$ 0.1157	<u>0.6733 $\pm$ 0.0181</u>
ChebyKAN						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-KAN	<u>0.6845 $\pm$ 0.0442</u>	0.6526 $\pm$ 0.0459	0.7545 $\pm$ 0.0441	0.6803 $\pm$ 0.0407	0.8568 $\pm$ 0.1045	0.6546 $\pm$ 0.0408
KAN-MLP	0.6904 $\pm$ 0.0447	0.6611 $\pm$ 0.0493	<u>0.7315 $\pm$ 0.0615</u>	0.7204 $\pm$ 0.0350	0.7558 $\pm$ 0.1306	0.6793 $\pm$ 0.0363
KAN-KAN	0.5731 $\pm$ 0.0132	0.5054 $\pm$ 0.1340	0.7285 $\pm$ 0.0106	0.5731 $\pm$ 0.0132	1.0000 $\pm$ 0.0000	0.5000 $\pm$ 0.0000

Building on the discussion of our reported results in Section 4.3.1 and Table 1, in this appendix, we report the full results of all of our experiments on both the NYU and UM sites to investigate the efficacy of different KAN variants and architectural configurations within our ViT-style classification network. Table 7 reports the results of our experimentation across five different KAN variants (Efficient-KAN, FastKAN, FasterKAN, Wav-KAN, ChebyKAN) and three different architectural configurations (MLP-KAN, KAN-MLP, KAN-KAN) on the NYU site. Table 8 reports the results of experiments with these same variants and configurations on the UM site. For the configurations, MLP-KAN refers to setups where KAN blocks replace the MLP classification head, KAN-MLP refers to setups where the MLP blocks in the transformer encoder were replaced with KAN blocks, and in the KAN-KAN setup, KANs fully replaced

MLPs in the architecture.

On the NYU site, the ordering of the KAN variants, by best average performance, is: FastKAN $\rightarrow$ FasterKAN $\rightarrow$ Efficient-KAN $\rightarrow$ ChebyKAN $\rightarrow$ Wav-KAN. There is no clear winner in terms of the best architectural configuration, with two of the KAN variants (Efficient-KAN, Wav-KAN) achieving the best performance with the MLP-KAN configuration, another two (FastKAN, FasterKAN) achieving the best performance with the KAN-KAN configuration, and ChebyKAN achieving the best performance in the KAN-MLP configuration. That being said, the top two KAN variants, FastKAN and FasterKAN, use the KAN-KAN configuration, and we have previously found this configuration to perform well across thorough experimentation on the NYU site [Ward and Imran \(2025\)](#). Overall, FastKAN offers a 0.63% improvement over FasterKAN.

Table 8: Impact of various KAN variants and configurations within the ABFR-KAN framework on the UM site. Mean $\pm$ stdev scores are reported by averaging the results of 5-fold cross-validation. Best and second-best results are **bolded** and underlined, respectively.

Baseline						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-MLP	0.7182 $\pm$ 0.0530	0.6426 $\pm$ 0.1395	0.6999 $\pm$ 0.1311	0.7169 $\pm$ 0.1104	0.7397 $\pm$ 0.2197	0.6984 $\pm$ 0.0780
Efficient-KAN						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-KAN	<u>0.6818 $\pm$ 0.0407</u>	0.6250 $\pm$ 0.0754	0.6512 $\pm$ 0.1252	0.7325 $\pm$ 0.1444	0.6751 $\pm$ 0.2356	0.6615 $\pm$ 0.0507
KAN-MLP	0.7091 $\pm$ 0.0464	0.6377 $\pm$ 0.0709	0.6946 $\pm$ 0.1054	0.6918 $\pm$ 0.0756	0.7061 $\pm$ 0.1481	0.6887 $\pm$ 0.0577
KAN-KAN	0.7091 $\pm$ 0.0464	0.6769 $\pm$ 0.1095	0.7108 $\pm$ 0.1075	0.6715 $\pm$ 0.0684	0.7688 $\pm$ 0.1637	0.6956 $\pm$ 0.0626
FastKAN						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-KAN	0.6909 $\pm$ 0.0182	0.6274 $\pm$ 0.0966	0.7016 $\pm$ 0.0970	0.6497 $\pm$ 0.0397	<u>0.7927 $\pm$ 0.1909</u>	0.6637 $\pm$ 0.0314
KAN-MLP	<u>0.7000 $\pm$ 0.0464</u>	0.6563 $\pm$ 0.0927	0.7168 $\pm$ 0.0757	<u>0.6621 $\pm$ 0.0590</u>	0.7938 $\pm$ 0.1284	0.6913 $\pm$ 0.0551
KAN-KAN	0.7091 $\pm$ 0.0464	<u>0.6471 $\pm$ 0.0287</u>	<u>0.7123 $\pm$ 0.0926</u>	0.6683 $\pm$ 0.0482	0.7745 $\pm$ 0.1641	<u>0.6845 $\pm$ 0.0381</u>
FasterKAN						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-KAN	0.6818 $\pm$ 0.0575	0.5815 $\pm$ 0.0838	0.6955 $\pm$ 0.0824	0.6523 $\pm$ 0.0809	<u>0.7619 $\pm$ 0.1326</u>	0.6635 $\pm$ 0.0424
KAN-MLP	<u>0.6909 $\pm$ 0.0340</u>	0.6352 $\pm$ 0.0490	0.6716 $\pm$ 0.1045	0.6776 $\pm$ 0.0630	0.6852 $\pm$ 0.1791	<u>0.6671 $\pm$ 0.0426</u>
KAN-KAN	0.7091 $\pm$ 0.0464	<u>0.6203 $\pm$ 0.0490</u>	0.7093 $\pm$ 0.1027	0.6813 $\pm$ 0.0761	0.7859 $\pm$ 0.2114	0.6777 $\pm$ 0.0278
Wav-KAN						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-KAN	0.7727 $\pm$ 0.0287	0.7184 $\pm$ 0.0620	0.7682 $\pm$ 0.0482	<u>0.7840 $\pm$ 0.1086</u>	0.7837 $\pm$ 0.1314	0.7768 $\pm$ 0.0270
KAN-MLP	<u>0.7545 $\pm$ 0.0464</u>	0.6876 $\pm$ 0.0803	0.7204 $\pm$ 0.1090	0.8133 $\pm$ 0.1012	<u>0.7086 $\pm$ 0.2241</u>	<u>0.7312 $\pm$ 0.0448</u>
KAN-KAN	0.7364 $\pm$ 0.0445	<u>0.6921 $\pm$ 0.0547</u>	<u>0.7218 $\pm$ 0.0457</u>	0.7683 $\pm$ 0.0779	0.7032 $\pm$ 0.1246	0.7295 $\pm$ 0.0350
ChebyKAN						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
MLP-KAN	0.7091 $\pm$ 0.0464	0.6429 $\pm$ 0.0722	0.7007 $\pm$ 0.1118	0.6766 $\pm$ 0.0527	0.7409 $\pm$ 0.1825	0.6859 $\pm$ 0.0528
KAN-MLP	0.7273 $\pm$ 0.0498	<u>0.6345 $\pm$ 0.1052</u>	0.7443 $\pm$ 0.0838	<u>0.6801 $\pm$ 0.0374</u>	0.8412 $\pm$ 0.1742	0.7061 $\pm$ 0.0794
KAN-KAN	0.6182 $\pm$ 0.0464	0.5741 $\pm$ 0.0487	0.4880 $\pm$ 0.2272	0.7278 $\pm$ 0.1458	0.5099 $\pm$ 0.3765	0.5604 $\pm$ 0.0416

On the UM site, the ordering of the KAN variants by best average performance changes to: Wav-KAN $\rightarrow$ ChebyKAN $\rightarrow$ Efficient-KAN $\rightarrow$ FasterKAN $\rightarrow$ FastKAN. Wav-KAN in an MLP-KAN configuration achieves the best overall performance across all KAN variants and configurations. However, its comparatively poor performance under the same configuration on the NYU site indicates that the efficacy of KAN in replacing MLPs is heavily dependent on data characteristics. Also on the UM, there continues to be no consistent winner among the configurations, with Efficient-KAN and FasterKAN performing the best with a KAN-KAN configuration, FastKAN and ChebyKAN performing best in the KAN-MLP setting, and Wav-KAN performing the best under the MLP-KAN setting.

As a final experiment, we performed two experiments to examine how a hybrid configuration of FastKAN and Wav-KAN would perform across both the NYU and UM sites. The results of these experiments are shown in Table 9.

In this table, we are only showing hybrid variations of the KAN-KAN configuration; we are not attempting to combine the functionality of both KAN methods into one.

On both ABIDE I sites, the hybrid configurations do not manage to outperform the respective best-performing single KAN variants. On the NYU site, FastKAN-WavKAN and WavKAN-FastKAN both fall short of outperforming FastKAN-FastKAN, although both methods end up outperforming the baseline WavKAN-WavKAN setting, with WavKAN-FastKAN (so, WavKAN in the ViT encoder, FastKAN as the classification head) performing better than its inverse version. On the UM site, WavKAN-WavKAN offers the best performance, with the hybrid versions failing to surpass it. But again, the hybrid configurations do show some promise in outperforming the other baseline configuration, FastKAN-FastKAN, with FastKAN-WavKAN performing slightly better than WavKAN-FastKAN. From these results, two things are clear. *First*, the hybrid configura-

Table 9: Exploration of hybrid KAN configurations within the ABFR-KAN framework. Specifically, FastKAN and Wav-KAN hybrids. Mean $\pm$ stdev scores are reported by averaging the results of 5-fold cross-validation. Best and second-best results are **bolded** and underlined, respectively.

NYU Site Classification Performance						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
FastKAN-FastKAN	0.7427 $\pm$ 0.0342	0.7041 $\pm$ 0.0708	0.8003 $\pm$ 0.0200	0.7267 $\pm$ 0.0469	0.8984 $\pm$ 0.0646	0.7159 $\pm$ 0.0434
FastKAN-WavKAN	0.6897 $\pm$ 0.0324	0.6327 $\pm$ 0.1084	0.7472 $\pm$ 0.0379	0.7079 $\pm$ 0.0555	0.8147 $\pm$ 0.1335	0.6664 $\pm$ 0.0522
WavKAN-FastKAN	<u>0.7017 $\pm$ 0.0394</u>	<u>0.6797 $\pm$ 0.0606</u>	0.7444 $\pm$ 0.0408	<u>0.7327 $\pm$ 0.0445</u>	0.7658 $\pm$ 0.0921	<u>0.6924 $\pm$ 0.0388</u>
WavKAN-WavKAN	0.6726 $\pm$ 0.0201	0.6466 $\pm$ 0.0185	0.6933 $\pm$ 0.0542	0.7540 $\pm$ 0.0750	0.6637 $\pm$ 0.1157	0.6733 $\pm$ 0.0181
UM Site Classification Performance						
Configuration	ACC	AUC	F1	PRE	SEN	SPE
FastKAN-FastKAN	<u>0.7091 $\pm$ 0.0464</u>	0.6471 $\pm$ 0.0287	0.7123 $\pm$ 0.0926	<u>0.6683 $\pm$ 0.0482</u>	0.7745 $\pm$ 0.1641	0.6845 $\pm$ 0.0381
FastKAN-WavKAN	0.7000 $\pm$ 0.0617	0.6158 $\pm$ 0.1216	0.7220 $\pm$ 0.0934	0.6569 $\pm$ 0.0821	0.8291 $\pm$ 0.1836	0.6734 $\pm$ 0.0615
WavKAN-FastKAN	0.6909 $\pm$ 0.0340	0.7059 $\pm$ 0.0835	0.7177 $\pm$ 0.0439	0.6613 $\pm$ 0.0731	<u>0.8103 $\pm$ 0.1288</u>	<u>0.6979 $\pm$ 0.0370</u>
WavKAN-WavKAN	0.7364 $\pm$ 0.0445	<u>0.6921 $\pm$ 0.0547</u>	<u>0.7218 $\pm$ 0.0457</u>	0.7683 $\pm$ 0.0779	0.7032 $\pm$ 0.1246	0.7295 $\pm$ 0.0350

tions are not suitable for outperforming the best-performing single KAN-KAN configurations; and *second*, they do show some promise in boosting the performance of less optimally performing methods. This indicates that even though we deem hybrid KAN configurations not suitable for further pursuit for our specific task of ASD vs. NC from rs-fMRI data, there may still be use cases for them in other tasks, and this warrants further exploration.