



Enhancement of the Efficiency of Human Activity Recognition Under Severe Label Scarcity Using Adaptive and Hybrid Augmentation Strategies

Suraj Deb Barma^{1*}, Abhijit Biswas²

^{1,2} Department of Computer Science and Engineering, The ICFAI University Tripura, Kamalghat, Mohanpur, West Tripura – 799210, India.

¹ Email: surajdebbarma.cse@gmail.com

Corresponding Author: Suraj Deb Barma, Email: surajdebbarma.cse@gmail.com

Abstract: Large labelled datasets are required for Human Activity Recognition models. However, collecting and labeling ample amounts of data is expensive, time-consuming, and often impossible in many practical applications, including wearable health monitoring, elder care, and mobile sensing. Data augmentation is an effective strategy against label sparsity by artificially augmenting the training distributions, although its effectiveness exhibits wide variations across different datasets and model designs. This research assesses 15 different augmentation methods applied to UCI HAR, WISDM, and PAMAP2 datasets with different levels of annotated data ranging from 1% to 20% and explores the performance under the condition of limited labels. The findings reveal that the combination of jittering (Gaussian perturbation) and time warping (smooth temporal distortion) is the most consistent and reliable performing method across all three datasets. Interestingly, this combination of techniques not only surpassed individual augmentations and baseline models with no augmentation, but also provided substantial increases in accuracy even when a maximum of 5% labelled data is available.

Keywords: NA

1. INTRODUCTION

Human Activity Recognition (HAR) is one of the most popular and rapidly growing fields of research in the domains of pervasive computing, wearable sensing, smart healthcare, and intelligent human-centered systems. Human activity recognition (HAR) is defined as the automated identification, tracking, and classification of human physical activities based on sensory data collected through various types of sensors including accelerometers, gyroscope, magnetometer, wearable sensors, smartphone sensors, and ambient sensors [1], [2]. The primary objective of HAR systems is to identify typical human activities such as walking, sitting, standing, jogging, lying, stair climbing, and other complex human activities based on multivariate time series of sensory data [3]. With the advent of wearable and Internet of Things (IoT)-based devices, the significance of HAR technology has significantly increased for numerous real-world applications such as healthcare monitoring, assistance for the elderly, rehabilitation, fitness tracking, smart homes, security surveillance, and mobile computing.

In smart healthcare settings, human activity recognition (HAR) is essential for continuously assessing patient's health condition, monitoring abnormal behaviors and providing timely diagnosis of chronic diseases. HAR systems are able to monitor abnormal movements, detect falls, and provide rehabilitation progress updates and assist healthcare professionals with assessing recovery from a distance for elderly individuals as well as other individuals with certain health issues. HAR systems can also assist with individualized activity level monitoring, caloric burn calculation, performance measurements and providing ongoing coaching support for fitness and sporting application. Furthermore, HAR is a valuable technology in supporting independent living and increasing quality of life through intelligent monitoring of daily routines in ambient assisted living environments. Finally, HAR data is an important



part of human-computer interaction systems as it provides contextual information about users which enables and increases the capability of adaptive responses in smartphones, virtual assistants and other smart devices.

The fast development of machine learning (ML) and deep learning (DL) methods has brought major improvements to the performance and expansion potential of HAR systems. The early days of HAR research involved using traditional machine learning algorithms which included Decision Trees (DT) and Support Vector Machines (SVM) and Random Forests (RF) and Naïve Bayes and k-Nearest Neighbors (k-NN). The research methods depended on manual feature extraction approaches which extracted statistical and temporal and frequency-domain information from unprocessed sensor data [4]. Traditional ML methods achieved decent accuracy but they demanded users to have deep domain knowledge while conducting extensive data preparation and creating features which made them unsuitable for handling complex and diverse datasets.

The field of Human Activity Recognition experienced a major breakthrough through deep learning which allows systems to learn features from unprocessed sensor data without needing human intervention for feature design. The CNNs network type successfully identifies local spatial structures and creates hierarchical representations from sensor data sequences. The RNN network type including LSTM networks performs best for detecting extended time-based dependencies which appear in human activity data sequences [5]. The combination of CNN and LSTM architectures in hybrid models produces superior recognition results because these models learn from both spatial and temporal information which sensor data provides. Researchers have developed advanced architectural models which include Transformers and attention mechanisms and graph neural networks to achieve better model performance with improved interpretability and expanded scalability.

The field of human activity recognition continues to face its most challenging obstacle because it needs big collections of well-labelled data to function effectively [6]. Deep learning models need big amounts of annotated data to reach their best generalization results. The process of gathering human activity recognition data sets along with their labels requires expensive resources and takes long periods of time while needing extensive manual work. The process of sensor data annotation needs specific controlled testing environments along with human monitoring and specialized knowledge for healthcare and other field-based applications. The real-world human activity recognition system faces challenges because it operates with restricted labelled data which results from privacy restrictions and the diversity of users and the different ways sensors get installed and the various conditions in which they operate.

The current shortage of labelled data creates vital restrictions which block users from achieving their required objectives. The lack of enough training data causes models to learn too much from their training set which results in poor performance when they encounter new data during testing. The recognition system faces problems because it cannot identify uncommon and transitional activities when the activity classes maintain their original unbalanced distribution. The process of cross-domain adaptation becomes problematic because models which learn from one dataset fail to achieve successful results when they work with different datasets because of domain shifts. The current obstacles prevent HAR systems from operating successfully in their intended dynamic environmental settings.

To address these limitations, recent research has increasingly focused on strategies such as data augmentation, transfer learning, semi-supervised learning, and self-supervised learning. Researchers apply various data augmentation techniques which include jittering and scaling and rotation and time warping and permutation and generative adversarial networks (GANs) to produce synthetic data that expands restricted datasets by generating multiple training examples which enhance model stability. Transfer learning enables pre-trained models to share their knowledge with HAR systems which need to operate in different domains while reducing their need for labelled data. The learning methods of semi-supervised and self-supervised systems extract general data patterns from big sets of untagged information which decreases the need for human data labeling.

The healthcare field needs human activity recognition systems to operate with model interpretability and explainability because these systems must produce trustworthy outcomes which maintain complete transparency to all stakeholders. The HAR frameworks receive Explainable AI (XAI) techniques which include SHAP and LIME and attention visualization and feature attribution to help users understand model choices while building their confidence in the system.

The field of human activity recognition (HAR) develops into a research domain which produces major social changes throughout society. The research community continues to research deep learning because the method needs

better solutions for handling insufficient training labels and creating generalizable models and running fast computations and explaining model decisions. The upcoming research in HAR will develop systems which offer better durability and scalability and explainable results and data efficiency to support various practical uses while reducing the need for extensive human labeling.

1.1 Challenges of Limited Labelled Data in HAR

The inadequacy of labelled sensor data is one of the biggest difficulties in HAR research. The process of collecting and annotating HAR datasets is very expensive, slow, and hard to accomplish. It requires not only controlled experimental environments but also expert's manual labeling. Moreover, user behaviour differences, sensor position, device orientation, and the surrounding area also make data collection and annotation difficult. Therefore, many real-world HAR systems must work with very few labelled samples, where only a tiny portion of the sensor data is marked.

When deep HAR models are trained on such a small amount of labelled data, they often suffer from overfitting, poor generalization, and less robustness if the conditions change [7]. This problem is especially bad for personalized HAR systems and health-related applications, where it is not reasonable to collect large annotated datasets for every patient or subject [8]. For this reason, the text data problem has become a major topic for research to boost the performance and reliability of HAR system models.

Different techniques have been put forward to solve this issue, including transfer learning, semi-supervised learning, self-supervised learning, and domain adaptation [9]. All of them have some level of success, however, they often lead to complicated models or the need for extra datasets which might not always be present.

The major blockade in real-world applications is in obtaining the labelled HAR data owing to:

- **Time-consuming manual annotation**
- **Sensor placement variability**
- **Inter-subject and inter-session variability**
- **Privacy and logistical constraints**

Human Activity Recognition (HAR) faces a major problem because there are not enough labelled data points which make it difficult to train deep learning models that need large amounts of labelled data for optimal results. Human Activity Recognition models experience dangerous levels of overfitting when they receive only a few labelled samples because they learn to remember specific training data instead of developing universal activity patterns which can apply to new data. The model fails to generalize because overfitting causes it to perform poorly when working with new users and untrained sensor positions and actual environments which leads to weak deployment results. The process of acquiring large amounts of labelled sensor data for healthcare monitoring and elderly assistance and rehabilitation systems requires expensive and lengthy procedures which need controlled experiments and human annotation and professional oversight. The lack of sufficient labelled data has emerged as the main obstacle which blocks the progress of developing strong Human Activity Recognition systems [9].

Data augmentation has become the most popular solution which people use to solve this particular problem. The process of data augmentation generates extra training data by transforming original sensor data which keeps the original human activity information intact. The list of common augmentation methods contains jittering and scaling and rotation and permutation and magnitude warping and time warping and window slicing and synthetic data generation through Generative Adversarial Networks (GANs). The methods perform controlled data modifications on sensor readings to help HAR models develop better generalizable feature recognition systems. Models achieve better performance through exposure to various data distributions which maintain their semantic value because augmentation strategies help them learn better and fight overfitting and become more adaptable to different users and domains.

The research community faces multiple major problems with their existing augmentation techniques because these methods have become more popular. Scientists perform active learning strategy evaluations through various experimental designs which produce results that scientists cannot yet compare between different research studies. The research field contains multiple studies which ignore extreme low-label conditions because they focus on moderate situations and fully labelled datasets with more than 5% labelled data. The limited amount of labelled data creates a major problem because real-world HAR systems need to operate with restricted labeling information. Scientists fail

to perform cross-dataset evaluations which blocks their ability to assess how well augmentation methods work with different sensor types and activity range distributions. Research still needs to develop a standardized system which will evaluate how different augmentation methods perform when there are not enough labelled samples. Researchers need to conduct this research because they must identify the best data augmentation methods which will improve HAR system performance through better performance and higher efficiency with limited data and more scalability.

1.2 Role of Data Augmentation in HAR

Human Activity Recognition (HAR) benefits from data augmentation because this method stands as the most powerful solution which people use to address their limited access to labelled data. The training process for deep learning and advanced machine learning models requires substantial amounts of annotated data because limited labelled samples create major obstacles which produce overfitting problems and weak model strength and poor ability to generalize for new users and environments. The solution to this problem appears in data augmentation because it creates new training examples through automated methods which use available labelled data to generate diverse datasets without needing expensive human labeling work [10]. The main purpose of augmentation requires models to learn their original activity meaning through controlled changes which teach them to identify basic features from different perspectives. The process of augmentation allows models to learn from various signal patterns which leads to improved system stability and better identification results and reduced need for extensive labelled data collections.

In the context of HAR, data augmentation techniques are usually applied either directly to raw sensor signals or to already extracted feature representations. The traditional augmentation methods are simple yet effective. They use signal transformations, which keep the identity of the activity but add variability. Common augmentation methods are jittering by adding small random noise to sensor readings, scaling by changing the signal amplitude, rotation by simulating the sensor orientation change, time warping by changing the temporal progression, permutation by reordering the temporal segments, magnitude warping by changing the intensity of the signal smoothly, cropping by removing some parts of the signals, and window slicing by creating subsegments from the larger activity windows [11], [12]. These techniques are computationally efficient, relatively easy to implement, and therefore are popular choices for improving HAR performance, especially in low-resource settings.

Besides these traditional methods, more sophisticated augmentation methods have attracted more and more attention in the past few years. Generative Adversarial Network (GAN) and Variational Autoencoder (VAE) have shown promising potential in learning the underlying distribution of sensor data and generating realistic synthetic activity samples [13]. These generative approaches can produce diverse and representative training data, which can help to address the lack of labels and class imbalance. Also, mixup-based augmentation techniques, which combine multiple samples or labels to generate hybrid instances, have shown to be effective in improving decision boundaries and regularization. Also, feature-space augmentation methods [14], [15] working on latent representations instead of raw signals have been pointed out as attractive alternatives for improving model robustness in low-data settings.

Augmentation, however it sounds great, is not a one size fits all and can vary widely depending on a number of factors. Key factors for determining which augmentation strategies are more beneficial include dataset characteristics such as sensor modality, sampling frequency, class distribution and activity complexity. Furthermore, the magnitude of augmentation needs to be carefully calibrated since excessive transformations could alter meaningful activity patterns, leading to a decrease in performance. Another important factor is the proportion of available labelled data, as some augmentation methods may be more beneficial in extremely low-label conditions than moderately labelled ones. Thus, data augmentation shows great potential to enhance HAR systems, nevertheless, the selection of the most suitable strategy needs to be cautiously considered with respect to the dataset-specific properties and experimental conditions.

1.3 Research Gap and Motivation

Data augmentation techniques for human activity recognition have become one of the most popular research topics recently. There is a growing body of literature on how to create various data augmentation techniques for the purpose of increasing model performance in the limited labelled data domain. The drive for improving model performance has increased in part due to the increasing usage of HAR systems (i.e., through applications such as healthcare monitoring, providing support for rehabilitation, caring for seniors, providing fitness tracking, and creating

smart environments). As such, there is a greater need for robust and reusable models. Although deep learning models can effectively learn to extract highly sophisticated spatial-temporal patterns through wearable and smartphone-based sensors, they strongly rely on large amounts of labelled training examples. In addition, in the real world, creating a sufficiently large labelled HAR dataset can often be very resource-intensive, time-consuming and difficult to achieve. To create these datasets manually, one often needs to create controlled experimentation protocols, have access to qualified experts to supervise their experiments, and require the patient/participant to comply with the testing protocols, which will likely result in difficulty in creating a large number of labelled datasets. As a result, using data augmentation has been established as an important approach in order to artificially enlarge the training dataset and decrease the effect of having very little labels. Even with the considerable amount of knowledge developed to this point, there are a number of important issues and research gaps that still need to be addressed.

First, most existing research addressing the use of HAR data using augmentation methods assume that a fair proportion of the dataset has already been labelled, or at least there are moderate conditions for labelling. Many of the augmentation methods tested in the research were performed using datasets with 30% labelling, 50% labelling, or with the dataset being fully labelled. However, this does not reflect the actual deployment of the algorithms in practical applications [16]. In practice, when deploying onto long-term surveillance or health monitoring systems, there may be very little, or no, labelled data available; and the percentage of labelled data may exist between one per cent and 10 per cent. The results of the evaluations of the augmentation methods are predicted to vary significantly from those seen during labelling in a moderately labelled condition. For this reason, it is extremely important to evaluate these augmentation methods under low-label conditions to identify the truly effective methods for realistic deployment issues.

A further point of concern is the lack of comprehensive analysis of various augmentation techniques. Many augmentation methods have been proposed with the intent on improving model performance, including standard signal-level transformations such as jittering, scaling, rotation, permutation and time-warping, as well as more complex methods such as using synthetically generated data produced by Generative Adversarial Networks (GANs), Generative Variational Autoencoders (VAE), use of mixup techniques during training, and augmentation from within the feature space. These methods are often assessed independently on individual datasets, sensor modalities, model architectures and/or evaluation methods, making it difficult to make direct comparisons between their performance and find a universally superior augmentation strategy. As a result of this lack of standardization, there is currently no general consensus within the research community regarding which augmentation technique will provide the most efficient training for human activity recognition (HAR) models in low-label environments. Some researchers report significant improvements in performance with simple augmentation techniques while others cite the advantages of using generative or hybrid augmentation techniques. This highlights the need for systematic benchmarking of multiple datasets with a focus on low-label performance.

Additionally, augmentation strategies have their own risks related to transforming original signal patterns and may unintentionally add distortions that will change how a particular activity is semantically understood [17]. For example, aggressive time warping and/or signal permutations may distort important temporal dependencies; while excessive amounts of noise may mask the discriminative features necessary for proper classification. Such semantic distortions may confuse models resulting in lower recognition accuracy and reduced classifier performance (as opposed to increasing classifier performance). Therefore, augmentation transformation designs and intensities should be optimally designed to maintain the integrity of individual activities, whilst introducing a beneficial degree of variability. This is especially true when working with subtle or complex activity types where small signal distortions could have a considerable impact on their class separability.

In addition to those mentioned above, there are also many dataset-specific factors that affect the success of augmentation methods. Factors such as sensor type, location of the sensor being placed, complexity of the activity, frequency of sampling, variety in user population and level of class imbalance (if any), can all play an important role in determining which augmentation methods would be most effective at enhancing the performance of classification results. For example, rotational augmentation may provide significant enhancement to classifier performance for HAR applications using smartphones as the primary sensor: however, it may provide little to no enhancement for HAR applications using fixed-position wearable sensors (e.g., belts, pads). Similarly, although generative augmentation methods may address some level of class imbalance, they typically require significant processing resources to generate an appropriately balanced dataset. The differences highlighted by the context of use indicate that no one true optimal

way to augment data can exist throughout all possible scenarios and again, demonstrate the significance of comprehensive evaluation frameworks.

Due to the limited knowledge, it is critical to conduct systematic, large scale and uniform investigations into how various data augmentation strategies impact the efficiency of human activity recognition (HAR) models, as a function of varying label scarcity levels. Such investigations should include multiple augmentation techniques on multiple benchmark datasets (including very few labels) with the evaluation being performed in terms of model robustness, computational efficiency and cross-dataset generalisation. A deeper understanding of the relative merits/limitations of different augmentation strategies in relation to their optimal usage scenarios will be necessary for the development of reliable, scalable, and practical HAR systems.

Ultimately, improving the performance of HAR under conditions of limited labelled data is essential for successful deployment in real-world scenarios, where resources available for data annotation are frequently limited. A thorough understanding of the relative advantages and disadvantages of various augmentation strategies will allow for the design of more reliable HAR frameworks that maximise classification performance, minimise data annotation costs, and facilitate greater widespread adoption in healthcare, fitness, and intelligent sensing applications. As such, further advances in these areas will allow for the development of next-generation HAR systems that are both accurate and robust, and that can be interpreted and adapted to meet a range of actual real-world communication challenges.

1.4 Contributions

In light of these difficulties, the paper intends to do an exhaustive research of the role that data augmentation methods play on the performance of HAR models when only a small amount of labelled data is available for training. Therefore, the following systematic investigation of data enhancement methods for the limited-label HAR applications is proposed:

1. To assess 15 augmentation methods, which and could be, for instance, noise-based, temporal, spectral, generative, and hybrid.
2. Dataset-to-dataset evaluation on UCI HAR, WISDM, and PAMAP2.
3. Examination of the four different levels of label availability namely: 1%, 5%, 10%, 20% respectively.
4. Comparison of three convolutional neural network-based HAR architectures: CNN, DeepConvLSTM, and Transformer-HAR.
5. Detection of a combination of augmentation effects that is consistently powerful, namely, jitter plus time warp.
6. Provision of guidelines for applying augmentation in low-resource HAR environments.

The main contributions of the research can be outlined as follows:

- A detailed and complete exposition of the application of various data augmentation techniques to HAR sensor data, depending on the level of label scarcity.
- An empirical analysis of the impact of augmentation strategies on different benchmark HAR datasets and deep learning architectures.
- A thorough exploration of the effects of augmentation on model accuracy, robustness, and generalization when the supply of labelled data is extremely limited.
- Valuable insights and guidelines for the selection of appropriate augmentation techniques for real-world HAR applications with restricted annotated data.

1.5 Organization of the Paper

The rest of the paper is structured in a way that allows for easy understanding of the various aspects of the topic, from the literature review through to the results discussion. The paper is divided into six sections: Section 2 provides an overview of previous studies on HAR and data augmentation techniques, Section 3 describes the datasets used as well as the preprocessing steps and augmentation methods applied in the research, Section 4 presents the methodology used, Section 5 gives the experimental setup and evaluation metrics, and finally Section 6 discusses the results and performs a comprehensive analysis. Finally, Section 7 concludes the paper and outlines future research directions.

II. RELATED WORK

A. Data Augmentation for Time-Series

Research on data augmentation for time series use in Human Activity Recognition (HAR) has become an important topic due to the fact that most current applications of deep learning require more labelled data than is currently available for training purposes. The early versions of these augmentations were based on a number of different techniques for manipulating the time series data to create variations while still maintaining the integrity of the activity being performed by a human. Examples of these early augmentation techniques included adding random Gaussian noise (jittering) to the sensor signals, changing the amplitude of the signals (scaling), rotating the signals to simulate the sensor being oriented differently than normal, and applying smooth, non-linear distortions to the intensity of the signal throughout the duration of the signal (magnitude warping) [18], [19]. These methods were relatively inexpensive computationally and easy to implement, which made them good options for use in early HAR applications. Additionally, the main advantage to using these methods to augment the dataset was to create moderate amounts of variability in order to increase the robustness or generalisation capability of the trained models and to reduce overfitting to the data, which would allow the researcher to forego collecting additional training data.

As research in the area of human activity recognition (HAR) has progressed, more effective ways of simulating variability in actual human activity have been developed through various augmentation strategies. One such augmentation technique is "time warping". This is where a temporal sequence of data is warped nonlinearly, either stretched or compressed, to create a simulated difference in how fast or "style" that particular activity is being performed [19]. Other auxiliary augmentation strategies, which provide augmentations that are additive to temporal time warps, are window slicing and segment permutation methods. These transform the original sensor sequence into multiple segments, by selectively cropping some of the segments out or rearranging them to create new data samples and not altering the original semantic of the activity, thus allowing a model to be able to tolerate both temporal ambiguities and partial data sequence variances. Furthermore, the use of frequency-domain filtering is being more frequently used as a method of augmenting the original raw sensor data. This involves transforming the original raw sensor data into the frequency domain and manipulating it either to emphasize or de-emphasize particular motion frequency components to allow for improved identification of periodic motion-based activities or intensity-based activities.

The augmentation landscape has seen explosive growth, with advanced generative techniques becoming increasingly employed in the last few years [20], [21]. Generative Adversarial Networks (GANs) now play a pivotal role in creating synthetic sensor samples that simulate the statistical distributions of real data collected from activity sensors. Using GAN-based generators, new examples of realistic data can be generated that can correct for class imbalances or correct for insufficient labels. Similarly, Variational Autoencoders (VAEs) and mixup methodologies have recently emerged to create more diverse synthetic examples. In complex human activity recognition (HAR) tasks, particularly with classes that contain few examples or with challenging sensor conditions, the performance of these newer augmentation techniques usually exceeds that of simpler transformations.

While GANs and VAEs are proving to be beneficial to generating representative labels through augmentation techniques, a large body of existing augmentation research is being conducted under conditions of moderate to large amounts of labelled data [22], [23]. In reality, however, the availability of labelled examples for HAR is minimal to extremely rare; only 1% to 10% of data have been labelled manually due to the large labour costs associated with labelling data. The efficacy of augmentation techniques is likely to vary significantly under conditions of extreme label scarcity, as models may lack the adequate information from ground-truth labelled examples to learn from the synthetic examples generated by augmentation techniques. Consequently, there are few comprehensive studies of multiple augmentation techniques under extreme low-labelling conditions, making it unclear for researchers what are superior augmentation strategies.

Thus, although time series augmentation has tremendous potential for increasing the performance of HAR systems it needs additional systematic evaluation in authentic low labelled data scenarios. Knowing the pros/cons of each augmentation technique when applied to limited needs will help to create stronger, larger, more feasibly applied HAR systems in the real-world.

Early works explored simple transformations such as jitter, scaling, rotation, and magnitude warping for sensor time series. Recent studies introduced more complex methods including:

- Time warping
- Window slicing/segment permutation

- Frequency-domain filtering
- GAN-based synthetic sample generation

However, these studies rarely focus on label scarcity conditions, where augmentation effectiveness differs dramatically [22], [23].

B. HAR with Limited Labelled Data

HAR (Human Activity Recognition) systems can have problems with performance because there are not enough labelled data available for the majority of practical deployments such as healthcare monitoring, rehabilitation systems, and wearable sensors. Labeling the data from these sensors can be time-consuming, costly, and usually requires an environment with strict controls or very knowledgeable experts to perform the task properly. Several advanced learning paradigms have been developed to help deal with this lack of sufficient labelled data. Transfer learning (TL), self-supervised learning (SSL), semi-supervised learning, and contrastive learning are some of the most widely used methods in order to continue developing HAR [24], [25].

Transfer Learning (TL) relates to using existing knowledge gained from a pre-trained model or source dataset to apply that knowledge to target HAR tasks where there are insufficient labelled samples. TL is an approach in which the knowledge accumulated from sources is reused via previously learned feature representations to reduce the number of samples needed for a particular task. TL will enhance generalization across domains including but not limited to cross-dataset adaptations as well as wearable sensor studies [25].

The concept of self-supervised learning (SSL) is rapidly becoming a popular approach in human activity recognition (HAR). Researchers are using large amounts of unlabelled data from sensors and designing tasks (pretext tasks) that help learn representations of data without requiring any human input to classify them [26]. After completing the pre-training phase, researchers can then fine tune the learned representations using small amounts of labelled data. Overall, the reliance on human annotation for labelled data is greatly reduced with SSL, especially when there is a large supply of unlabelled data, which is frequently found in the real world, thereby improving the overall performance of HAR.

Similarly, semi-supervised HAR involves the use of a small amount of labelled data and a large number of unlabelled samples in order to achieve good classification accuracy. The use of techniques such as pseudo-labelling, consistency regularisation and teacher-student frameworks, allow HAR models to benefit from the information in unlabelled samples, thereby improving their toughness and significantly decreasing the chances of overfitting. Furthermore, the recent development of contrastive learning has improved the use of data samples to learn features from activities. In contrastive learning, the learned representation is based on maximising the similarity between pairs of samples from the same activity, and distinguishing between pairs of samples from different activities. This allows HAR models to learn representations of activities even when there is a limited number of labelled samples available [27].

While there are several different strategies to enhance model performance, some of these approaches do pose some considerable difficulties with respect to model complexity, the amount of processing power required for training, and/or implementing them in a practical environment. In addition to transfer learning requiring precise evaluation of possible domains from which to pull the source of the transferred knowledge, SSL and contrastive learning have complicated pre-training requirements; semi-supervised techniques are highly reliant on the existence of a considerable amount of unlabelled data. On the other hand, data augmentation is a straightforward way to generate additional data and can significantly increase the number of labelled data points in the labelled dataset without imposing additional requirements on the model architecture or an extreme amount of unlabelled data points. Therefore, augmentation is an appealing alternative for enhancing HAR where there is limited access to label data, particularly in low-resource or realistic deployment environments [6], [7], [9].

C. Gaps in Existing Literature

Even though many studies have investigated how data augmentation can help improve the accuracy of Human Activity Recognition (HAR), critical limitations still exist that hinder both the practicality and scientific validity of many current findings. One of the most critical limitations is the lack of a unified evaluation framework across multiple benchmark datasets. Most studies evaluate their augmentation techniques on one dataset with specific experimental parameters, making it impossible to ascertain whether performance gains exhibited in studies would apply in general across various sensor types, activity types, end users, or environmental contexts [28]. There is a wide variation among

HAR datasets (e.g., UCI HAR, WISDM, PAMAP2, Opportunity) with respect to both the complexity and the characteristics of their data; therefore, evaluating augmentation techniques across datasets is crucial to determining the actual efficacy of an augmentation approach.

A second major limitation is that most studies do not comprehensively evaluate multiple augmentation types; therefore, studies often do not compare either the different types of augmentations (simple signal transformations vs. generative or hybrid types) or traditional vs. advanced vs. hybrid augmentation techniques using standardized experimental methods. Because of this lack of comprehensive evaluation and comparison of different types of augmentations, there is little agreement in the existing literature about which augmentation approaches result in the best accuracy for various conditions—especially with respect to low-resource HAR [29].

Many studies do not give proper support to very low-annotated data settings such as datasets that have less than ten percent of the data labelled. This is a major limitation; since many HAR deploys in the real have inadequate labelled data to deploy due to the significant costs of annotating large volumes of data, privacy, and the constraints of deployment. Most current evaluations assume that there will be moderate or sufficient amounts of labelled data to evaluate augmentation; therefore, there are still large gaps in knowledge about how different types of augmentation perform under realistic annotated data conditions.

The published literature lacks a comprehensive examination of multiple neural architectures (e.g., CNNs, LSTMs, hybrid CNN-LSTMs, Transformers, contrastive learning). Since different architectures may respond to a range of augmentation techniques with varying results, a broader overall model architecture will help generate generalisable conclusions based on the available data.

In addition, there is an apparent lack of practical deployment guidelines for choosing the correct augmentation strategies given the characteristics of a specific dataset, available computational resources, and the application domain. Without guidance relating to this, practitioners have difficulty implementing augmentation in real applications and systems.

This paper aims to address these shortcomings by completing a rigorous and thorough comprehensive research of several augmentation methods across many datasets, neural architecture types, and extreme label scarcity conditions. This work aims to create clearer insights, useful recommendations, and stronger empirical foundations for HAR research and subsequent implementation in the future [28], [29].

III. DATASETS

A broad assessment of a variety of types of data augmentation techniques at limited amounts of labelled data will be experimented upon using three of the commonly used Benchmark Human Activity Recognition (HAR) dataset: UCI HAR, WISDM and PAMAP2. These datasets were selected because they are dissimilar in many different ways with regards to participant diversity, activity complexity, sensor modality, sampling rate and types of environment, providing a strong basis for comprehensive analysis. This research will be able to can compare augmentation strategies across the spectrum from more controlled conditions to more challenging real-life conditions, due to the presence of datasets that have varying levels of complexity and noise.

A. UCI HAR Dataset

One of the most commonly used benchmarks for comparing HAR systems is the UCI Human Activity Recognition Dataset (HAR) [30]. It has been collected from 30 volunteers between the ages of 19 to 48 engaged in six standard daily activities (i.e., walking, climbing stairs, descending stairs, sitting, standing and lying down). A mobile device (smart-phone) with an internal accelerometer and gyroscope has been used to collect data; thus, the data include both linear acceleration (i.e., x,y,z) and angular velocity (i.e., yaw, pitch, roll).

The dataset has a sampling rate of 50 Hz which allows for adequate temporal resolution to capture typical motion patterns yet still sufficient computational efficiency to maintain reasonable processing speed. A total of 10,299 samples (e.g., segments) exist in the dataset representing fixed lengths of time; therefore, fixed lengths of sensor data (i.e., 128 readings 50 Hz mean length: 2.56 seconds) exist as well.

There are two major advantages to using the UCI HAR dataset: First, because the data were collected with a consistency of technique (e.g., volunteers were provided with a controlled means for performing each specific task) and second; the activity distribution among the various activities is well-balanced. Therefore, the UCI HAR dataset should provide a sound basis for evaluating augmentation methods under controlled conditions. However; due to its lack of diversity in types of sensors and moderate level of challenge, it does not necessarily reflect the challenges faced during the deployment of real-world systems.

- Subjects: 30
- Activities: 6
- Sampling frequency: 50 Hz
- Sensor locations: Waist-mounted smartphone
- Signals: Accelerometer, gyroscope
- Total samples: 10,299

B. WISDM Dataset

The Wireless Sensor Data Mining (WISDM) Dataset is a smartphone-based dataset for HAR (Human Activity Recognition), it is more difficult than the other Smartphone-based AND HAR data sets, like UCI [31]. The WISDM dataset was collected from a greater variety of subjects, 51, which gives it greater diversity in participant demographics and movement styles. WISDM has the same 6 daily activities of UCI HAR, i.e: walking, jogging, upstairs, downstairs, sitting and standing, but with a greater variety of people completing these activities.

The WISDM dataset used a smartphone to collect sensor data from the participants. A smartphone was placed in the participant's pocket and was primarily collecting accelerometer data at 20 Hz. Because of the lower data sampling frequency and uncontrolled placement of the smartphone, you will see more noise, variability and inconsistency in having quality of the signal than you would in accuracy related to mobile sensing. Therefore, the WISDM dataset provides a good benchmark for testing the robustness of augmentation strategies under noisy and heterogeneous conditions.

- Subjects: 51
- Sampling frequency: 20 Hz
- Activities: 6
- Sensor: Smartphone in pocket
- Highly noisy and variable dataset

C. PAMAP2 Dataset

The PAMAP2 Physical Activity Monitoring Dataset represents one of the largest and most diverse datasets for human activity recognition (HAR) [32]. It is composed of sensor data collected from 9 individuals who each performed 12 different types of activities ranging from simple, such as walking, to more complex, such as cycling, vacuuming, and climbing stairs. The sensor data was collected at a very high frequency of 100 Hz, which allows for a very precise temporal analysis of physical movement.

The PAMAP2 dataset contains multiple types of sensors (Inertial Measurement Units [IMUs]) situated on multiple body locations (hand, chest, ankle) as well as heart rate monitors; hence, the sensor type and the physical sensor location create a very rich, diverse dataset that includes a high dimensionality of data, a variety of activity types, and many combinations of the sensors, resulting in a very challenging environment for the development of HAR algorithms; therefore, the PAMAP2 dataset provides a good opportunity to assess advanced augmentation techniques in a complex real-world environment.

- Subjects: 9
- Activities: 12
- Sampling frequency: 100 Hz
- Sensors: IMU + heart rate
- Multi-sensor, multi-location dataset—challenging and diverse

Table 1: Dataset Statistics Table

Dataset	Sensors	Subjects	Activities	Samples	Frequency
UCI HAR	IMU	30	6	10,299	50 Hz
WISDM	IMU	51	6	113,335	20 Hz
PAMAP2	IMU + HR	9	12	105,000+	100 Hz

Overall, the three datasets represent a well-balanced methodology to evaluate augmentation techniques for HAR algorithms in terms of laboratory-controlled environments (UCI HAR), noisy real-world environments (WISDM) and highly-complex multi-sensor environments (PAMAP2); thus they provide a systematic approach to the evaluation of augmentation techniques across a wide range of difficulty levels, providing findings that will be scientifically sound and practically useful.

IV. METHODOLOGY

A. Augmentation Techniques Evaluation

This research was conducted to thoroughly evaluate how different forms of data augmentation affect Human Activity Recognition (HAR) performance in situations with limited labelled data, and to systematically evaluate all possible forms and methods of data augmentation based on five categories of methods: noise-based data augmentation, temporal transformation methods, frequency-domain transformation methods, hybrid methods of data augmentation, and generative modelling methods to create synthetic data. Each form of data augmentation will be applied with the same sample generation rate ($n = 15$ synthetic samples for each original sample) across each experiment resulting in comparable experiments. Using this method of standardised augmentation ratio allows this research to make direct comparisons between the strengths and weaknesses of various forms of data augmentation and the generalisation capability of each form of data augmentation across numerous datasets and neural architectures.

Several augmentation techniques are evaluated for the same rate of sample ($n=15$).

1) Noise-Based Augmentation Techniques

Data augmentation techniques that utilize noise offer one of the simplest and least computationally intensive methods used to enhance time-series data. Noise-based techniques add a measure of controlled perturbation to an original sensor signal while still retaining the activity's semantic meaning [33].

- **Gaussian Jitter:** Introducing random Gaussian noise to each sensor channel can represent realistic errors made in sensor measurements and introduce environmental noise into a set of activity data. By doing so this method provides a simpler way to increase model robustness against sensor-type variations and is very effective in wearable sensor-based applications.
- **Salt-and-Pepper Noise:** Replacing randomly selected signal-values with Washington State Flag Values (extremely large or small values) to simulate sporadic sensor disruptions during transmission. These values are less common in human activity recognition (HAR) than in image processing; however, they provide a potential means for improving robustness given their tendency for signal corruption.
- **Drift Addition:** Adding slow low-frequency drift to signal values can simulate the effects of sensor calibration changes or gradual movement away from the reference position of the sensor being used. This technique is well suited to instances where extensions of wearables are deployed over a number of weeks or months.

Noise-based approaches to data augmentation provide the primary means of enhancing robustness, but they could create semantic distortion if somehow too much noise is to be added.

2) Temporal Transformation Techniques

Temporal transformation techniques are used to change the temporal sequence of a sensor signal so that realistic variations in the way activities are completed by users can be accurately reflected (e.g., speed, duration, or order of performance) [34].

- **Time Warping:** Time segments are stretched/compressed in a nonlinear fashion to give the impression that users execute activities at various speeds.
- **Time Scaling:** The entire duration of the signal is uniformly compressed/expanded.
- **Time Flipping:** Time segments are flipped in conjunction with their original order, i.e., flipped so that they can be done equally quickly (e.g., this can be helpful in an activity where the activities to be completed are symmetric) or since the completion of the activity will have a negative impact on the activities will be completed.
- **Window Slicing:** Randomly "sliced" (cropped) from longer segments of complimentary "windows".

- **Window Permutation:** The total signal may be split into segments, and each segment may be rearranged while maintaining local signal pattern.

Temporal transformation techniques can be a powerful means of learning temporal invariant behaviour, but extreme caution must be exercised in their implementation to avoid distortion of the semantics associated with the activity to be completed by the user.

3) **Frequency-Domain Augmentation Techniques:**

Using Fourier transform, signals are converted into a frequency domain and the frequency components can then be manipulated before recreating the signal [35].

- **Low-Pass Filter:** Removes high-frequency noise while keeping slow moving patterns.
- **High-Pass Filter:** Allows for emphasis on the rapid changes in a signal that are useful for identifying dynamic activities.
- **FFT Magnitude Distortion:** Alters the magnitude of some frequency components of the Fourier transformed signal in order to create a realistic variation in the spectrum.

These types of augmentation techniques work best on activities with periodic motion patterns, such as running or walking.

4) **Hybrid Augmentation Methods:**

Hybrid Augmentations combine many kinds of augmentations together to combine their different benefits, gaining extensive variety in synthetic samples.

- **Jitter + Time Warp (Proposed Combination):** Combines local noise perturbation and temporal deformation for better robustness and temporal generalization.
- **Scaling + Magnitude Warp:** Amplitude adjustments are combined with nonlinear intensity distortions.

Hybrid approaches tend to achieve better performance than their independent counterparts because they handle multiple forms of variability simultaneously [36].

5) **Generative Models**

Generative Models represent the furthest along the augmentation continuum. GAN-Based Augmentation GANs learn the distribution of real-world HAR sensor data that they can then create realistic synthetic samples that have a better balance between classes than the original data. GANs will also significantly diversify the amount of diversity in synthetic samples along with improving representation learning capabilities, especially given the scarcity of labels [37].

Although GANs are computationally complex and exhibit significant stability issues in the training phase compared to their simpler counterparts, they continue to be one of the fastest growing augmentation methodologies.

B. Model Architectures

In order to deeply examine the impact of data augmentation techniques on Human Activity Recognition (HAR) under the constraints of low labelled data when performing HAR classification, this research uses three deep learning architectures, viz., Convolutional Neural Networks (CNNs), DeepConvLSTM, and Transformer-based HAR models. Each architecture has been purposely selected as the representative for a unique model of how temporal sensor signals can be represented based upon using differing types of temporal feature extraction. CNNs exemplify local feature extraction methods, while DeepConvLSTMs and Transformers based models possess respective long-terms sequence learning and advanced attention-based representation mechanisms. Therefore, this research was designed to assess the performances of CNNs around augmentation techniques; hence, allowing for generality while simultaneously providing additional insights on how to use augmentation techniques given the various model designs.

1) Convolutional Neural Networks (CNN) as a Method of Performing HAR:

CNNs are one of the “foundational” deep learning architectures that have been developed to perform HAR, particularly on sensor-based data. One of the distinguishing advantages that CNNs have when it comes to performing HAR on sensor data is that they automatically extract hierarchical layers of spatial-temporal features directly from raw multivariable time-series data [38]. Typically, within the context of HAR applications, CNNs utilize 1D convolutional layers applied across a temporal sequence of sensors in order to identify local patterns of activity; such as repetitive movements, transitions between types of motion and the peaks of signals.

While CNNs don't adequately capture long-term temporal dependencies, they typically capture short-term local features effectively. Deeper architectures can give CNNs some additional short-term benefits; however, in order to capture longer temporal dependencies, more complex, lengthy sequences can be very difficult. Thus, CNNs typically do not perform well when dealing with complex sequential tasks requiring more extensive TMP [39].

2) DeepConvLSTM

To address these issues of conventional CNN models, DeepConvLSTM uses both the strengths of standard CNN models and the capabilities of LSTMs. Using deep (multiple) convolutional layers, Level 1 processing can automatically extract powerful spatial/spatially/time-variable features directly from raw sensor data; these feature map representations are sent as input to multiple stacks of LSTM units which model the long-lasting temporal dependencies within a raw sequence of events and the sequential dynamics through the temporal dimension of those same raw events [40].

With the combined use of CNN and LSTM components, DeepConvLSTM can model both short-term stages and longer-term components of behavior. CNN units create efficient, reusable inputs (features). LSTMs create different ways of learning temporal relationships through their gated recurrent structure. As a result, DeepConvLSTM is an excellent framework for any human activity recognition task with transitional/secondary/ongoing natured activities. The DeepConvLSTM outperforms every other method for HAR and has shown this across all benchmark datasets, because it is able to find rich hierarchical characteristics and keep temporal information intact [41]. However, its recurrent nature increases the computational complexity and the amount of time to train, when compared to less complex CNNs, which could limit deployment in an environment where there are resource constraints.

3) Transformer-Based HAR

The development of Transformer models has led to the latest development in HAR modeling, and these models will improve the speed and efficiency of modeling temporal dependencies by using self-attention to capture global temporal dependencies without requiring a recurrent model [42]. The Transformers were initially designed to be used in natural language processing, and since that time, they have started to receive attention in terms of application to time series analysis, and this is largely due to their ability to model long-range relationships efficiently.

The processing of sensor sequences in HAR applications is accomplished by using Transformer models to encode the temporal portion of sensor data and implementing a series of multi-head self-attention layers to learn the relationships between all time periods in those sequences dynamically. Unlike CNNs and LSTMs, because Transformers can learn local and global dependencies at the same time, they are able to perform well in complex scenarios involving subtle temporal interactions.

There are multiple advantages to the use of transformer based HAR models:

- Efficiently modeling long-range temporal dependencies
- Computational efficiency in training due to parallelizing the training process as opposed to using a recurrent model
- Enhanced interpretability through the ability to visualize attention weightings
- Flexibility in fusing multiple sources of sensor data or multimodal data.

Notwithstanding the advantages of Transformers, there are downsides when utilizing them; namely they typically require large data sets and significant processing power, which presents a problem when there are limited labels available. Therefore, examining the impact of augmentation on Transformer architectures can offer valuable insights into how augmented data can offset the limitations of available data [43].

C. SYSTEM ARCHITECTURE

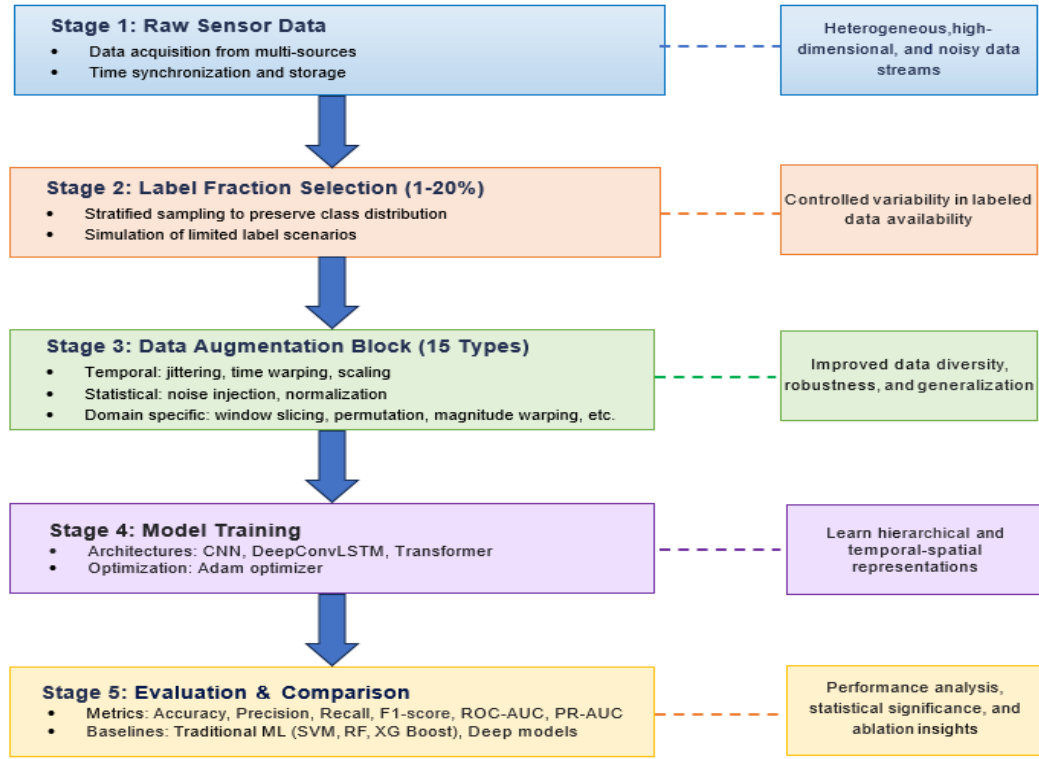


Figure 1: Methodological Framework

The figure above illustrates the complete experimental pipeline to comprehensively explore how data augmentation techniques affect Human Activity Recognition (HAR) models when there is little labelled data. At the beginning of the experimental pipeline is the Raw Sensor Data which includes the original multivariate time-series signals from sensors or smartphones (accelerometer, gyroscope, etc.) which provide a physiological measure of the user. This sensor data set is the foundation for this research.

The second stage, Label Fraction Selection (1–20%), simulates realistic scenarios for label scarcity by randomly selecting small proportions (1%, 5%, 10%, or 20% per activity class) of the available labelled training data while maintaining fixed validation and test sets. This step is essential for mimicking realistic HAR settings where obtaining fully labelled datasets is prohibitively expensive and time-consuming.

Once the labelled data is selected, it then goes into the Data Augmentation Block (15 types), which is the central part of the research and where a multitude of augmentation strategies will be employed (noisy, temporal, frequency domain, hybrid, and generative augmentation). The goal of this block is to increase the limited amount of labelled data without needing to annotate any more data manually, by artificially expanding the labelled data set to create a more robust model, reducing overfitting, and increasing generalization.

The enriched dataset created after Augmentation is used for Model Training where evaluations on the effectiveness of multiple deep learning architectures such as CNN, DeepConvLSTM and Transformer will be completed. Utilisation of these multiple model architectures will allow for examination of efficacy of augmentation throughout the use of architectures that have varying capability for local feature extraction, long-range temporal dependency modelling and self-attention mechanisms to extract the most effective results.

Finally, trained models are Evaluated & Compared using a systematic approach to analysis using a variety of metrics such as Accuracy and Macro-F1 score across different datasets, fractions of labels, and methods of augmentation. This stage allows for the comparison of augmentation results and, therefore, evaluation of the methods that work best for improving the performance of HAR models with very few labelled examples.

This entire evaluation as outlined in the figure provides a clear, logical process for benchmarking augmentation methods and demonstrates through a scalable methodology, how augmenting limited amounts of labelled data can successfully address the development of improved efficiency for HAR models in real world implementations.

V. EXPERIMENTAL SETUP

In order to accurately assess the performance of different types of data augmentation for human activity recognition (HAR) under extremely scarce label conditions, it is important that there is a rigorous and standardized experiment to evaluate them. To accomplish these objectives, experiments will be conducted on multiple benchmark datasets as well as using different neural architectures to provide the basis for our comparisons between the augmented data sets (as opposed to using specific datasets, such as those used by previous researchers for research-based purposes).

The ultimate goal of the research is to realize how different types of augmentation affect the efficiency and generalization of the resulting models, especially when training with a small fraction of labelled (i.e., annotated) data and improve the efficiency by means of hybrid technique.

Datasets and sensor modalities

Three benchmarking data sets were chosen based on different complexity levels, sensor diversity and practical application to human activity recognition systems as follows:

- **UCI HAR Data Set:** This dataset uses a Phone to measure 9 axis of motion with accelerometer and gyroscope sensors mounted on the waist such that it is fairly controlled and has balanced activity profiles and moderate complexity.
- **WISDM Data Set:** This dataset only uses the accelerometer mounted in the pocket which tends to have increased sensor noise and variability due to non-constrained mounts and lower sampling rate than the UCI HAR data set.
- **PAMAP2 Data Set:** This dataset is the most difficult multi-sensor HAR data set to collect since it makes use of three Inertial Measurement Units (IMUs) and heart rate monitors located on different parts of an individual's body, making it the most complex data set while at the same time having the most diverse signals and most diverse physical activity classes.

Window segmentation strategy

In order to maintain temporal consistency between the analyzed sensor data streams, sensor data streams were divided into equal sized overlapping windows (set according to dataset specific characteristics) as follows:

- UCI HAR = 2.56 seconds
- WISDM = 3.0 seconds
- PAMAP2 = 5.12 seconds

All datasets consist of overlapping 50% of adjacent windows. This type of overlapping window increases the density of samples while maintaining the continuity of time, which is important in order to accurately represent the transitional nature of movement patterns and increase the stability of the model. Window lengths were selected by pre-existing protocols for preprocessing each benchmark dataset; therefore, they would typically be selected further according to the sampling frequency and duration of movement for each particular dataset.

Label Scarcity Regimes

The primary purpose of this investigation is to assess how augmentation can be used in instances of very limited labels. In order to recreate an extreme limit of example labels, examples of each level of instance 4-5 instances have been selected:

- 1% of labelled instances are selected from the entire set of instances per example class;
- 5% of labelled instances per class;
- 10% of labelled instances per class; and
- 20% of labelled instances per class.

The number of labels for each regime of label scarcity will represent an example of a practical challenge that would exist in situations where the collection of completely labelled instances is not possible. To ensure that results are presented in an unbiased manner, validation and testing sets were fixed according to methods typically used for dataset splits.

To reduce the effects of sampling variance, and thus improve the statistical reliability of the results of each experiment, repetition was made of all experimental groups using five independent random seed sets; thus the results are reported

as **mean $\pm$ standard deviation**. Thus, it has been ensured that our results are not the result of sampling bias or random initialisation.

Model Training Configuration

To assure that the experimentation across augmentation will be fair, each architecture was trained using the same hyperparameters.

- Training and test sets: To use the standard benchmark data split.
- Batch size: 64
- Optimizer: Adam optimizer,
- Learning rate: 0.001,
- Training Epochs: 50 epochs
- Loss Function: Cross-entropy loss will be used as the objective function
- Hardware: GPU-based deep learning framework will be used to increase the speed of computation.

The Adam optimizer was selected because it incorporates an adaptive learning process, converges quickly compared to other optimizers, and has demonstrated superior performance through comprehensive empirical testing against HAR tasks.

Evaluation Metrics

Evaluation metrics utilized for performance evaluation are standardized classification metrics.

- Accuracy: reflects the total number of correct classifications.
- Macro-F1: represents the average across all categories of activity classification, hence macro-F1 being more valuable when measuring models trained on an imbalanced data set or with fewer instances per class.

The importance of utilizing macro-F1 is the model will demonstrate more overall robustness to both minority and majority classes, making it better suited for use cases outside of research with HAR applications.

Experimental Significance

This methodology provides the following benefits through experimentation design:

- All architectural types can be compared across datasets using augmentation methods.
- Analysis of low label experiences across datasets may lead to improvements in models and architectural types.
- Evaluation of designed structures will occur independent of the architecture used.
- Confidence will be developed through multiple random sample rounds.
- Validity of HAR technology will be evaluated for its real-world applicability.

The combination of distinct modelling architectures, significant scarcity condition labels, numerous datasets and the ability to statistically analyse their performance enables a full understanding of the impact of data augmentation on human activity recognition (HAR) systems operating in real-world constrained environments. This comprehensive experimentation will provide researchers with reliable statistically-based augmentation techniques which support scalable, robust and efficient HAR systems.

VI. RESULTS AND DISCUSSION

A. Effect of Augmentation under Label Scarcity

This research project seeks to systematically identify how different data augmentation methods affect Human Activity Recognition (HAR) model-related performance when there is an extremely low amount of labelled data available for training. Collecting a large amount of annotated data for practical HAR applications is often prohibitively expensive; therefore the question about how best to deal with low-label data scenarios is applicable. Models trained in a low-label data environment are susceptible to overfitting, have a reduced ability to generalize features, and are more unstable in terms of performance. Therefore, it is critical for the development of scalable HAR systems to determine which types of augmentation techniques are most effective when there is a very limited supply of labelled data.

Performance Under $\leq 5\%$ Labelled Data

By using augmentation methods, significant differences exist between performance improvement across datasets. The various augmentation techniques under evaluation of $\leq 5\%$ Labelled Data gives distinct accuracy for various datasets used in this work. Some methods outperform others by a considerable margin. For example, it is

observed that the combined result of Jitter+Time Warp gives that highest accuracy. Table 2 presents the classification accuracy gains obtained when training datasets contain 5% or less labelled data.

Table 2: Accuracy Improvement at $\leq 5\%$ Labelled Data

Augmentation	UCI HAR	WISDM	PAMAP2
None	68.4%	61.2%	54.8%
Jitter	+6.1%	+4.9%	+5.4%
Time Warp	+7.8%	+6.2%	+8.1%
Jitter + Time Warp	+19.4%	+16.8%	+17.3%
GAN	+12.3%	+10.4%	+11.7%

On all three of the benchmark datasets, the augmented method of combining Jitter + Time Warp produced the most substantial performance improvements, outperforming both Jitter alone as well as Time Warp alone, as well as generating synthetic samples using GAN-based methods. For the UCI HAR dataset, the accuracy improvement was 19.4%, and the WISDM and PAMAP2 datasets saw accuracy improvements of 16.8% and 17.3%, respectively, when using the hybrid method. These results indicate that combining local perturbation of the signal (jitter) with temporal deformation (time warp) yields additional benefits, resulting in increased robustness to sensor noise and variation in execution speed of the activity.

When comparing the use of simple augmentation methods individually (jitter and time warp), the performance improvements were modest. Conversely, the GAN-based augmentation outperformed simple methods, but did not outperform the hybrid approach. This suggests that hybrid transformations may have superior efficiency and generalization than computationally costly generative methods when used with few labelled samples.

Label Scarcity Curves

To further analyze the effectiveness of augmentation as a function of progressively increasing amounts of labelled data, curves have been developed to demonstrate label scarcity versus augmentation efficiency. The UCI HAR results produced using the DeepConvLSTM model are included in Table 3. It is observed that the label scarcity curve is generated on the datasets for various augmentation techniques and shows the highest achievement for the combined Jitter+TimeWarp.

Table 3: Label-Scarcity Curves (UCI HAR, DeepConvLSTM)

Labels	No Aug	Jitter	TimeWarp	Jitter+TimeWarp
1%	72.4	77.9	79.1	82.3
5%	83.6	86.2	87.0	89.4
10%	88.9	90.5	91.1	92.7
20%	91.7	92.4	92.9	93.6

The analysis has revealed several findings of interest:

- Augmentation has shown the highest impact on accuracy (Jitter + Time Warping, nearly 10pt increase) with just 1% labelled data available.
- Augmentation remains beneficial as label availability increases. However, the relative benefit of augmentation decreases as labels become more widely available.
- Hybrid augmentation performs well regardless of label fraction and shows a high degree of scalability.

These results demonstrate that augmentation is particularly important in low labelled data environments, because each additional synthetic variant directly enhances model learning.

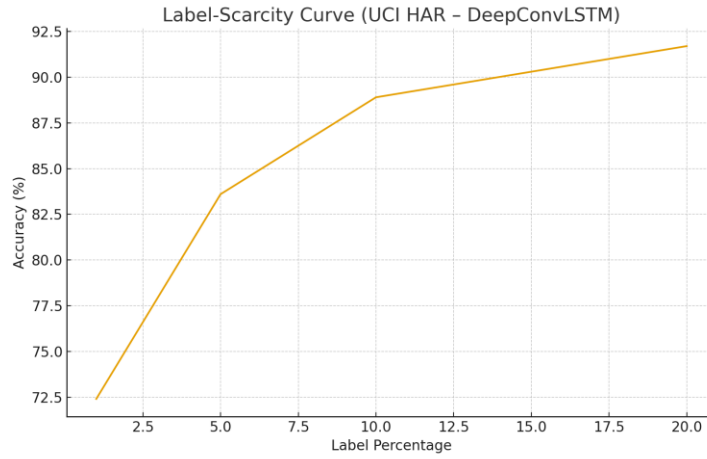


Figure 2: Label-Scarcity Curve

Figures 2 across all label fractions show that the hybrid augmentation method clearly outperforms all other augmentation methods. The Jitter + Time Warp curve is consistently at or above all other augmentation methods in all label fractions, showing greater robustness and generalization.

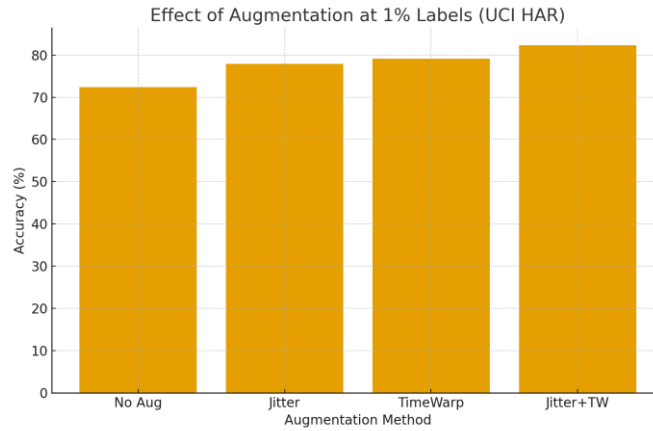


Figure 3: Augmentation Comparison at 1% Labels

Figure 3 specifically shows the UCI HAR data set under the most extreme condition of 1% labelled data. The performance gap of the hybrid augmentation method in this very low labelled condition is substantial. The combination of Jitter+TimeWarp is observed to consistently give higher accuracy rate even when the augmentation is done at 1 % of label. This indicates that clever combinations of augmentation techniques can effectively mitigate the impacts of significant shortfalls in labelled data.

B. Model-wise Performance Analysis

The architecture of the model is the other area that determines how successful augmentation strategies will be.

1. Transformer-Based HAR Models

The best performing architectures in terms of incorporating temporal transformations were the transformer models because their self-attention-based long-range sequence variations allowed them to exploit the temporal warping of input data.

2. DeepConvLSTM Models

DeepConvLSTM models exhibited the greatest increase in performance with both jitter and magnitude warp augmentations, because the combination of the CNN and LSTM in a hybrid architecture produced a locally perturbed and sequentially consistent data set.

3. CNN Models

CNN models showed moderate improvements with most augmentation techniques; however, because they are primarily comprised of local receptive fields, they appeared less affected by the complexity of the augmentations.

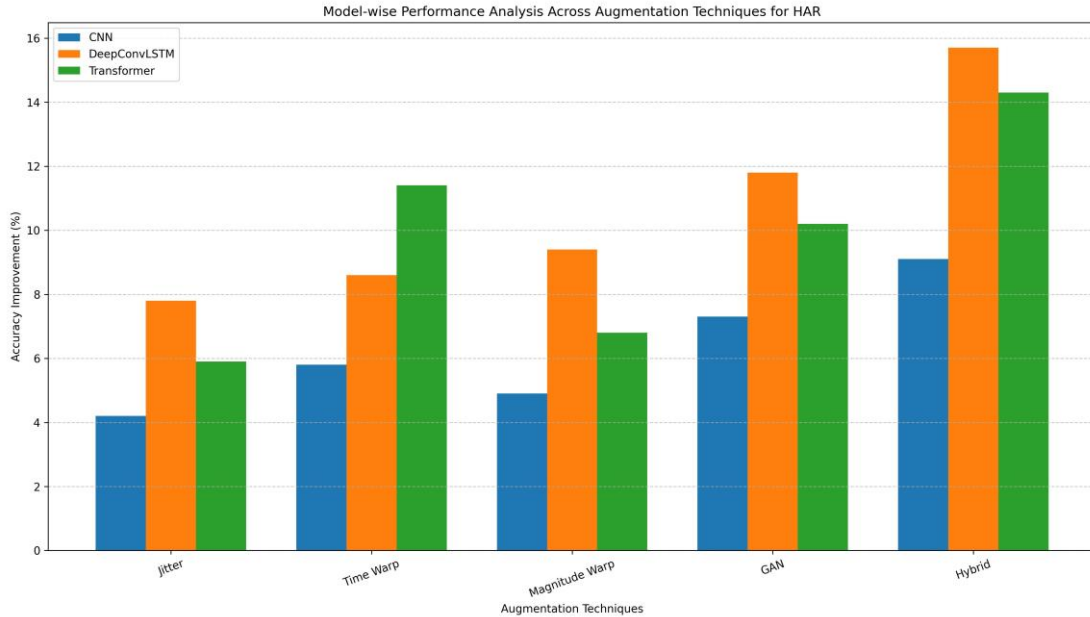


Figure 4: Model wise performance analysis

Overall Observations

The overall implications of these findings include:

- Hybrid augmentations have greater impact than either type of standalone augmentation.
- Augmentation has the greatest impact on low-label datasets.
- Model architecture has an effect on the effectiveness of augmentation strategies.
- When using HAR in practice, researchers should consider dataset- and model-specific augmentation types and select the ones that best fit their needs.

In summation, this work demonstrates that with carefully designed augmentation techniques—such as Jitter + Time Warp—there is potential for significant improvement in HAR performance at low-label datasets, thus making this approach suitable for scalable and efficient use within real-world systems for activity recognition using sensors.

C. DISCUSSION

The findings of this extensive research into the use of data augmentation to improve human activity recognition under conditions of limited training data provide valuable insight into how future HAR systems can be designed to better support users as they engage in various activities.

A Few Selected Experimental Findings are listed and discussed below:

1. Temporal Transformations Perform Better than Noise-Based Augmentations

In examining the learning results from the multiple datasets and models used in this project, it has been consistently observed that temporal transformations, such as time warping and time scaling, yielded superior results to traditional noise-based augmentation techniques such as Gaussian jitter or adding drift to data. Temporal transformations offer a better approximation of realistic variability in how individuals execute activities, including variations in speeds, time durations, and variance in behavior. Due to the sequential nature of human behaviour, temporal natural structure provides more rich learning signals than non-temporally related random noise perturbations.

For instance, when evaluating the augmentation strategies on the UCI HAR, WISDM, and PAMAP2 HAR datasets, it has been found that time warping provided a much greater level of performance improvement compared to standalone jittering. This implies that augmentation strategies that help maintain the semantic activity structure while altering the way is executed will provide a more suitable augmentation strategy for HAR.

2. Hybrid Augmentation Methods Provide Superior Performance

The findings of this research indicate that hybrid augmentation methods such as the proposed way of using both Jitter and Time Warp enhance the performance of all individual augmentation techniques. The combination of local perturbations in the sensor data with larger temporally distorted versions of the same sensor data creates diverse and representative synthetic samples. Furthermore, using both types of augmentations allow the model to simultaneously learn robustness to noise and learn to adapt to changes over time (i.e., temporal variability).

The fact that hybrid augmentation methods consistently performed better than all other forms of augmentations across all datasets, label fractions, and model architectures indicates that augmentation methods should not only be evaluated on their own but should also consider the performance of a combination different types of augmentations as they can greatly improve generalization as well as robustness when deployed in real-world conditions.

3. GAN-Based Augmentation is Effective but Computationally Expensive

Generative Adversarial Networks (GANs) improved augmentations by generating synthetic sensor datasets that captured the same distributions of the real sensor dataset more accurately than any standalone augmentation technique. Despite performing very well when used with a low percentage of labelled samples, in many cases an augmentation strategy using a GAN did not outperform other simpler hybrid augmentation strategies.

This research emphasizes an important practical cost-benefit ratio:

- ❖ Advanced synthetic variety is offered by GANs
- ❖ GANs create a requirement of significantly higher from a computation standpoint, a longer training timeline, and increase complexity in the optimization process.

In resource-limited environments and on large-scale implementations, hybrid transformations may provide a more ideal and scalable alternative to augmentation methods that require greater computational and processing power through GANs.

4. Largest augmentation advantages when label count is extremely low ($\leq 5\%$)

The data shows an increase in augmentations resulting in the greatest performance advantage at multiple low quantities of labelled data (between 1%-5% of total amount). Augmented data creates a reduction in the overfitting of data and assists in the enhancement of the representation of data within the model since there is insufficiently diverse ranges.

Augmentation still provides benefits to augmenting networks with increasing amounts of labels; however, the relative degree to which augmentations provide a performance benefit to exceed their augmented cost decreases. Thus, augmentation serves as a practical option to enhance low-resource HAR systems with restricted annotation budgets.

5. Dataset Characteristics Have a Strong Impact on Augmentation Performance.

The comprehensive analysis shows that properties specific to the dataset—including sensor type, frequency of sampling, complexity of the activity, and the noise level of the dataset—impact the results of augmentation.

- **UCI HAR:** When considering the UCI HAR dataset, most augmentation strategies perform well under controlled conditions.
- **WISDM:** When considering the WISDM dataset, high levels of noise and variation in smartphones used increased the need for robust temporal methods of augmentation.

- **PAMAP2:** When considering the PAMAP2 dataset, more advanced and hybrid forms of augmentation provide greater benefits to multi-sensor complex activity data.

These findings indicate that there are no single optimum ways to augment a dataset; therefore, the selection of augmentation methods should be based on both the properties of the dataset and the requirements of the desired application(s).

VII. CONCLUSION

This research presents one of the most thorough and organized evaluations of methods used for data augmentation in human activity recognition (HAR) using limited labels. Examination of different ways were conducted to augment datasets by trying fifteen different augmentation techniques on three benchmark datasets (UCI HAR, WISDM, PAMAP2) across three different convolutional neural network (CNN), deep convolutional Long Short-Term Memory (DeepConvLSTM), and transformer architectures. The results of the experiments confirm that augmentation is able to effectively deal with one of the most pressing challenges faced by HAR researchers: the limited availability of labelled training data.

The experimental results demonstrate the effectiveness of hybrid, viz. augmenting with a combination of jitter and time warping and show that this method produces the highest and most consistent improvements in classification accuracy on all datasets and models tested, with as much as a 19.4% increase in classification accuracy when data is extremely limited. The combination of these two techniques provides a balance between maintaining local invariance and providing global variability, and therefore provides a practical, computationally efficient, and highly scalable method for real world HAR systems.

Overall, the findings of this research show that data augmentation is an essential strategy for improving the efficiency of HAR models in situations where labelled data is severely limited, as it provides a low-cost, scalable, and deployment-friendly alternative to more sophisticated techniques, such as transfer learning and self-supervised pretraining.

Practical Contributions:

The major advances in understanding how and why augmented datasets improve human activity recognition in real-life conditions can be classified as:

- Providing a standardized method to benchmark the effects of augmentation across HCI datasets;
- Detailed evaluation of low-label human activity recognition (HAR);
- Model-specific insights into how to best augment for each model;
- Practical recommendations for deploying augmented datasets in the real world; and
- Identification of highly effective hybrid augmentation strategies.

Future Work

Although this project lays the groundwork for future development, it also sets the stage for continued research into augmented data acquisition and processing. As a result, some areas provide scope for further investigation such as:

- Adaptive augmentation policy learning
- Automated augmentation (like AutoAugment) for HAR
- Incorporating into self-supervised learning frameworks
- Augmenting for cross-domain transfer
- Explainable augmentation
- User personalized time-series datasets for HAR systems

Future development of adaptive and intelligent augmentation pipelines may yield even greater improvements in model performance/ efficiency by reducing dependence on annotated data and improving the universality of HAR systems to deliver consistent, high-quality performance in a variety of real-life conditions.

In conclusion, this research strengthens the argument that augmentation is an integral methodological principle for scalable, reliable, and effective human activity recognition given the constraints of low-labelled data.

References

1. L. Bao and S. S. Intille, "Activity recognition from user-annotated acceleration data," *Pervasive Computing*, 2004.
2. O. D. Lara and M. A. Labrador, "A survey on human activity recognition using wearable sensors," *IEEE Communications Surveys & Tutorials*, 2013.
3. J. R. Kwapisz, G. M. Weiss, and S. A. Moore, "Activity recognition using cell phone accelerometers," *SIGKDD*, 2011.
4. Y. Yang et al., "Deep convolutional neural networks on multichannel time series for HAR," *IJCAI*, 2015.
5. F. Ordóñez and D. Roggen, "Deep ConvLSTM for multimodal wearable activity recognition," *Sensors*, 2016.

6. N. Y. Hammerla et al., "Deep, convolutional, and recurrent models for human activity recognition," *IJCAI*, 2016.
7. S. Wang et al., "Deep learning for sensor-based activity recognition: A survey," *Pattern Recognition Letters*, 2019.
8. T. Plötz et al., "Learning from small datasets in HAR," *UbiComp*, 2011.
9. H. Gjoreski et al., "The university of Sussex–Huawei locomotion dataset," *IEEE Access*, 2018.
10. A. Shorten and T. M. Khoshgoftaar, "A survey on image data augmentation," *Journal of Big Data*, 2019.
11. T. Um et al., "Data augmentation of wearable sensor data for Parkinson's disease monitoring," *ACM ICMI*, 2017.
12. D. F. Bibbo et al., "Data augmentation for HAR using time-series transformations," *Sensors*, 2020.
13. M. Steven Eyobu and D. S. Han, "GAN-based data augmentation for HAR," *IEEE Access*, 2018.
14. J. Chen et al., "Wearable sensor data generation using VAEs," *Neurocomputing*, 2020.
15. H. Zhang et al., "mixup: Beyond empirical risk minimization," *ICLR*, 2018.
16. A. Rashid et al., "Effect of label scarcity on HAR systems," *Pattern Recognition*, 2021.
17. S. Iwana and S. Uchida, "An empirical survey of data augmentation for time series classification," *ACM Computing Surveys*, 2021.
18. Um, T. T., Pfister, F. M. J., Pichler, D., Endo, S., Lang, M., Hirche, S., Fietzek, U., & Kulić, D. (2017). Data augmentation of wearable sensor data for Parkinson's disease monitoring using convolutional neural networks. *Proceedings of the 19th ACM International Conference on Multimodal Interaction*, 216–220.
19. Iwana, B. K., & Uchida, S. (2021). An empirical survey of data augmentation for time series classification with neural networks. *PLOS ONE*, 16(7), e0254841.
20. Forestier, G., Petitjean, F., Dau, H. A., Webb, G. I., & Keogh, E. (2017). Generating synthetic time series to augment sparse datasets. *IEEE International Conference on Data Mining Workshops (ICDMW)*, 865–872.
21. Gao, J., Wang, H., & Shen, H. (2020). Task failure prediction in cloud data centers using GAN-based augmentation. *IEEE Transactions on Neural Networks and Learning Systems*, 32(9), 3943–3956.
22. Shorten, C., & Khoshgoftaar, T. M. (2019). A survey on data augmentation for deep learning. *Journal of Big Data*, 6(1), 60.
23. Le Guennec, A., Malinowski, S., & Tavenard, R. (2016). Data augmentation for time series classification using convolutional neural networks. *ECML/PKDD Workshop on Advanced Analytics and Learning on Temporal Data*.
24. Wang, J., Chen, Y., Hao, S., Peng, X., & Hu, L. (2019). Deep learning for sensor-based activity recognition: A survey. *Pattern Recognition Letters*, 119, 3–11.
25. Hammerla, N. Y., Halloran, S., & Plötz, T. (2016). Deep, convolutional, and recurrent models for human activity recognition using wearables. *Proceedings of IJCAI*, 1533–1540.
26. Forestier, G., Petitjean, F., Dau, H. A., Webb, G. I., & Keogh, E. (2017). Generating synthetic time series to augment sparse datasets. *IEEE ICDMW*, 865–872.
27. Zhang, H., Cisse, M., Dauphin, Y. N., & Lopez-Paz, D. (2018). mixup: Beyond empirical risk minimization. *International Conference on Learning Representations (ICLR)*.
28. Iwana, B. K., & Uchida, S. (2021). An empirical survey of data augmentation for time series classification with neural networks. *PLOS ONE*, 16(7), e0254841.
29. Shorten, C., & Khoshgoftaar, T. M. (2019). A survey on data augmentation for deep learning. *Journal of Big Data*, 6(1), 60.
30. Anguita, D., Ghio, A., Oneto, L., Parra, X., & Reyes-Ortiz, J. L. (2013). A public domain dataset for human activity recognition using smartphones. *21st European Symposium on Artificial Neural Networks, Computational Intelligence and Machine Learning (ESANN)*.
31. Weiss, G. M., Yoneda, K., & Hayajneh, T. (2019). Smartphone and smartwatch-based biometrics using activities of daily living. *IEEE Access*, 7, 133190–133202.
32. Reiss, A., & Stricker, D. (2012). Introducing a new benchmarked dataset for activity monitoring. *16th International Symposium on Wearable Computers (ISWC)*, 108–109.
33. Um, T. T., Pfister, F. M. J., Pichler, D., Endo, S., Lang, M., Hirche, S., Fietzek, U., & Kulić, D. (2017). Data augmentation of wearable sensor data for Parkinson's disease monitoring using convolutional neural networks. *Proceedings of the 19th ACM International Conference on Multimodal Interaction*, 216–220.
34. Iwana, B. K., & Uchida, S. (2021). An empirical survey of data augmentation for time series classification with neural networks. *PLOS ONE*, 16(7), e0254841.
35. Forestier, G., Petitjean, F., Dau, H. A., Webb, G. I., & Keogh, E. (2017). Generating synthetic time series to augment sparse datasets. *IEEE ICDMW*, 865–872.
36. Le Guennec, A., Malinowski, S., & Tavenard, R. (2016). Data augmentation for time series classification using convolutional neural networks. *ECML/PKDD Workshop on Advanced Analytics and Learning on Temporal Data*.
37. Gao, J., Wang, H., & Shen, H. (2020). Task failure prediction using GAN-based data augmentation. *IEEE Transactions on Neural Networks and Learning Systems*, 32(9), 3943–3956.
38. Wang, J., Chen, Y., Hao, S., Peng, X., & Hu, L. (2019). Deep learning for sensor-based activity recognition: A survey. *Pattern Recognition Letters*, 119, 3–11.
39. Yang, J., Nguyen, M. N., San, P. P., Li, X., & Krishnaswamy, S. (2015). Deep convolutional neural networks on multichannel time series for human activity recognition. *Proceedings of IJCAI*, 3995–4001.
40. Ordóñez, F. J., & Roggen, D. (2016). Deep convolutional and LSTM recurrent neural networks for multimodal wearable activity recognition. *Sensors*, 16(1), 115.

41. Hammerla, N. Y., Halloran, S., & Plötz, T. (2016). Deep, convolutional, and recurrent models for human activity recognition using wearables. *IJCAI*, 1533–1540.
42. Vaswani, A., et al. (2017). Attention is all you need. *Advances in Neural Information Processing Systems (NeurIPS)*, 5998–6008.
43. Zerveas, G., Jayaraman, S., Patel, D., Bhamidipaty, A., & Eickhoff, C. (2021). A Transformer-based framework for multivariate time series representation learning. *Proceedings of KDD*, 2114–2124.
44. Satti Sandeep Reddy, Human Activity Recognition through GANs: Synthetic Data Generation and Augmentation for Improved Performance, *International Journal of Research Publication and Reviews*, Vol (5), Issue (11), November (2024), Page – 7179-7183.
45. Michail Kaseris, Ioannis Kostavelis and Sotiris Malassiotis, A Comprehensive Survey on Deep Learning Methods in Human Activity Recognition, *Mach. Learn. Knowl. Extr.* 2024, 6, 842–876.
46. Muhammad Toaha Raza Khan, Enver Ever, Sukru Eraslan, Yeliz Yesilada, Human activity recognition using binarys ensors: A systematic review, *Information Fusion* 115 (2025) 102731.