



Intelligent Human Computer Interface Devices based on Neural Network

¹Ms. Yogita Soni, ²Dr. Ritesh Kumar Yadav

¹Research Scholar, ²Assistant Professor

^{1,2}Department of Computer Science

^{1,2}Sarvepalli Radha Krishnan University, Bhopal (M.P.)

Abstract

Intelligent solutions based on artificial intelligence (AI) techniques, algorithms, and sensor technologies are needed to give computers human communication abilities and to allow for organic human-computer interaction. In order to investigate trends in human-computer intelligent interaction (HCII) research, classify the available evidence, and determine possible avenues for future investigation, this study sought to identify and analyze the most advanced AI techniques, algorithms, and sensor technologies in the body of existing HCII research. We map the corpus of research on HCII in a methodical manner. Intelligent emotion, gesture, and facial expression identification, false news, and face anti-spoofing detection have been the main areas of study in the HCII domains. This research examines face anti-spoofing detection based on HCII utilizing various machine learning techniques.

Keywords- Human Computer, Intelligent Interaction, Anti-spoofing Detection, Machine Learning

I. INTRODUCTION

As science and technology advance, many technological pioneers are attempting to update HCII technology by combining text, audio, vision, and other information—a process known as multimodal information. Additionally, multimodal interaction has gained popularity in both business and academics [1]. Throughout this revolution, multimodal technology will progressively transform the whole world, not only voice and visual recognition. For instance, a number of industry-leading multimodal interaction fundamental technologies, including lip recognition, voice recognition, translation, and synthesis, have been used in a variety of sectors. By recording the motions of human hands and limbs, gesture interaction technology converts commands into a language that computers can understand. After touch screens, keyboards, and mice, it has emerged as another crucial HCII technique [2, 3]. The industry



standard for intelligent hardware is using microphone arrays to interpret data and hardware to remove noise. However, there is still a significant speech recognition bottleneck in complicated and loud environments [4]. Data-driven intelligence may be used to achieve the potential revolution of HCII, even while the next generation of revolutionary HCI technologies, like the introduction of touch technology and graphical interfaces, may not have an influence on the whole sector. The rapid advancement of artificial intelligence has significantly increased machine intelligence, and the thorough investigation of human-machine interaction has led to the creation of new technologies for automated voice recognition and gesture interaction [5].

The result of interdisciplinary research is HCII. In 1975, the idea of HCI was initially put out, and in 1980, the term was adopted as a profession. Research on human-computer interaction (HCI) is growing every day as the notion becomes more widely accepted. "The way people and machines deal with each other" [6] is the simplest definition of human-machine interaction. The development of deep learning technologies has significantly sped up HCI research. Fingertip contact between humans and robots has steadily changed from command to emotional communication, and this connection has encountered certain obstacles along the way. Voice and gesture, for instance, are becoming a part of our daily lives and are essential to the use of virtual reality as an entrance mechanism [7]. The logic of HCI is entirely different from the physical world, because both people and machines in the virtual world are not constrained by the objective rules of the physical world [8]. People may get a higher-dimensional perspective on information and a wider dimension of information reception via HCI in the virtual environment. Simultaneously, it might artificially increase our access to knowledge and experience. However, speech and gesture interaction are more important in the virtual world. The following are the primary issues at the moment. The first is how the computer can better grasp human language using machine learning and artificial intelligence technology, and how the dialog robot employed in voice recognition can successfully identify genuine interactive noises and background noise [9].

The second is that the challenge of gesture recognition is to precisely distinguish between continuous motions that are really conscious interaction actions and those that are unconscious. The third is determining which HCI system features are better suited for gesture recognition. The fourth is how deep learning technologies may enhance the precision of



interactive gesture detection and action capture. Due to these issues, in the future, conversational robots will be manufactured with machines adapting to people rather than humans adapting to machines. Artificial intelligence-based dialogue robots will progressively gain popularity [10].

In light of the current issues, the Google academic and literature database Web of Science was used to screen almost 1000 studies using keywords associated with HCII and machine learning, such as "intelligent HCI," "speech recognition," "gesture recognition," and "natural language processing." After five years (2019–2022) of screening, over 100 papers were ultimately chosen as the study material of this work out of almost 500 studies of research methodologies. The state of intelligent HCII's machine learning applications across a range of sectors, including natural language processing, voice interaction, and gesture identification, is examined. This paper summarizes and analyzes the use of natural language processing in chatbots and search engines, as well as the comprehension of voice and gesture interaction in a Virtual Reality (VR) environment in HCII. This study focuses on using machine learning technology to increase dynamic gesture interpretation, which might serve as a model for future HCII development.

II.LITERATURE REVIEW

Jarosz et al. [1], present an adaptable human- - robot cooperation engineering that integrates feelings and temperaments to give a characteristic encounter to people. To decide the profound condition of the client, data addressing eye stare and look is joined with other logical data, for example, whether the client is seeking clarification on pressing issues or has hushed up for quite a while. In this way, a fitting robot conduct is chosen from a multi-way situation. This engineering can be handily adjusted to cooperations with non-encapsulated robots like symbols on a cell phone or a PC. We present the result of assessing an execution of our proposed engineering overall, and furthermore of its modules for distinguishing feelings and questions. Results are promising and give a premise to additional turn of events.

Prathiba et al. [2], CBVR is so famous nowadays, in view of the expanded use of video based scientific frameworks. Video based examination is very compelling than picture investigation, as a progression of activities are caught by the video. This winds up with better critical thinking skill. The CBVR frameworks assume a significant part in supporting the human-PC collaboration. This paper presents a multimodal CBVR that considers both the



visual and sound data for recovering pertinent recordings to the client. Two modules are utilized by this work to manage video and sound information. The video information is handled to identify the huge casing from shots and is accomplished by Lion Streamlining Calculation (LOA). The elements are removed from the visual information and as for the sound information, MHEC and LPCC highlights are extricated. The separated elements are grouped by Kernelized Fluffy C Mean (KFCM) calculation. At last, the element data set is framed and is used in the question matching cycle during the testing stage. The presentation of the proposed work is tried as far as accuracy, review, F-measure and time utilization rates. The proposed CBVR framework demonstrates preferable execution over the current methodologies and is apparent through accomplished results.

Ince et al. [3], this paper presents an investigation of a varying media interface-based drumming framework for multimodal human-robot communication. The intelligent multimodal drumming game is utilized related to humanoid robots to lay out a varying media intuitive connection point. This study is essential for an undertaking to plan robot and symbol collaborators for schooling and treatment, particularly for kids with unique necessities. It explicitly centers around assessing robot/virtual symbol coaches, unmistakable cooperation gadgets, and versatile mixed media gadgets inside a straightforward drumming-based intuitive music mentoring game situation. A few boundaries, including the impact of the epitome of the coach/interface and the presence of criticism and preparing instruments, were the focal point of interest. For that reason, we made an intelligent drumming game depending on turn-taking and impersonation standards, in which a human client can play drums with a humanoid robot (Mechanical Drum Mate). Three intelligent situations with various trial arrangements for humanoid robots and cell phones were created and tried. As a piece of those situations, a framework that empowers drum strokes to be consequently distinguished and perceived utilizing hear-able signals was executed and integrated into the exploratory system. We checked the materialness and viability of the proposed framework in a drum-playing game with grown-up human guineas pigs by assessing it both unbiasedly and emotionally. The outcomes showed that the actual robot guide, the criticism and preparing instruments decidedly affected the subjects' exhibition and, true to form, albeit the actual medium is liked, the virtual mechanism for drumming caused less disappointment.

Wu et al. [4], the proposed feeling acknowledgment model depends on the various leveled



long-transient memory brain organization (LSTM) for video- electroencephalogram (Video-EEG) signal cooperation. The sources of info are facial-video and EEG signals from the subjects when they are watching the feeling animated video. The results are the relating feeling acknowledgment results. Facial-video includes and relating EEG highlights are separated in view of a completely associated brain organization (FC) at each time point. These highlights are combined through various leveled LSTM to anticipate the vital close to home sign edges at whenever point until the feeling acknowledgment result is determined at the last time point. Uniquely, a self- consideration instrument is applied to show the connection of the stacked LSTM at various progressive systems. In this cycle, the "specific concentration" is utilized to examine the human-close to home worldly groupings in each model, which works on the use of the key spatial EEG signals. Also, the cycle incorporates the worldly consideration component to anticipate the critical sign casing at next time point, which uses the key feeling information in transient area. The exploratory outcomes demonstrate that the characterization rate (CR) and F1- score of the proposed feeling acknowledgment model are essentially expanded by something like 2% and 0.015, individually, contrasted with different strategies.

Mosquera-DeLaCruz et al. [5], a multimodal communication point of interaction was produced for utilizing Google Chrome, Gmail and Facebook applications through gestural and verbal orders. The connection point actuates mouse and console orders from the handling of voice signs and recordings of the client's head development. The point of interaction doesn't handicap conventional console and mouse capabilities; besides, it just requires a webcam and a receiver, which are generally incorporated into convenient PCs.

Nayak et al. [6], infrared-Warm Imaging is a non-contact component for psychophysiological examination and application in Human-PC Collaboration (HCI). Continuous location of the face and following the Areas of Interest (return on initial capital investment) in the warm video during HCI is trying because of head movement ancient rarities. This paper proposes a three- stage HCI structure for registering the multivariate time-series warm video successions to perceive human inclination and gives interruption ideas. The main stage involves face, eye, and nose discovery utilizing a Quicker R-CNN (district based convolutional brain organization) design and utilized Numerous Example Learning (MIL) calculation for following the face returns on initial capital investment across the warm



video. The mean power of returns for capital invested is determined which shapes a multivariate time series (MTS) information. In the subsequent stage, the smoothed MTS information are breathed easy Twisting (DTW) calculation to order close to home states evoked by video improvement. During HCI, our proposed system gives significant ideas from a mental and actual interruption point of view in the third stage. Our proposed approach connotes better exactness in correlation with other order strategies and warm informational indexes.

Liu et al. [7], has been removed in line with the author(s) as well as supervisor. The Distributer apologizes for any burden this might cause. The full Elsevier Strategy on Article Withdrawal can be found at <http://www.elsevier.com/find/withdrawalpolicy>. Ensuing to acknowledgment of this unique issue paper by the capable Visitor Manager Sadia Noise, the uprightness and thoroughness of the friend audit process was explored and affirmed to fall underneath the exclusive expectations expected by Chip and Microsystems. There are likewise signs that a significant part of the Unique Issue incorporates unimaginative and intensely reworded content. Because of a design mistake in the publication framework, sadly the Manager in Boss didn't get these papers for endorsement according to the diary's standard work process.

Dybvik et al. [8], fostered a wearable trial sensor arrangement highlighting multimodal EEG+fNIRS neuroimaging pertinent for in situ examinations of human conduct in communication with innovation. A minimal expense electroencephalography (EEG) was coordinated with a wearable useful Close Infrared Spectroscopy (fNIRS) framework, which we present in two sections. Paper A give a thorough depiction of arrangement framework, information synchronization process, a system for utilization, including sensor application, and guaranteeing high sign quality. This paper (Paper B) exhibit the setup's convenience in three unmistakable use cases: an ordinary human-PC communication explore, an in situ driving examination where members drive a vehicle in the city and on the thruway, and an ashtanga vinyasa yoga practice in situ. Information on mental burden from profoundly environmentally legitimate test arrangements are introduced, and we examine examples learned. These incorporate satisfactory and unsatisfactory curios, information quality, and builds conceivable to research with the arrangement.

FACE SPOOFING IN BIOMETRIC SYSTEMS A collection of characteristics from the user's face or fingerprint is extracted by a biometric identification system in order to identify authentic entries. By contrasting the extracted feature set with a collection of feature sets found in the database or statistical models, the individual's identification is determined. For biometric identification, a variety of distinctive and identified behavioral and physiological characteristics have been investigated. facial spoof detection is a commonly used countermeasure technique to differentiate real facial characteristics from fraudulent ones [8]. Its goal is to identify an individual's physiological life indicators. Fig. 1 shows the fundamental system architecture of the face authentication model-based access control system. A relevant biometric characteristic must be shown to the sensor in order to operate or evaluate the system's ability to withstand spoofing assaults. In this instance, a camera serves as a sensor. Preprocessing is done on the collected facial photos to improve system performance. The feature extraction module extracts unique facial traits that may be used to represent the pictures from the preprocessed image. A feature vector that may more effectively distinguish real photographs from fraudulent ones is the module's output. To recognize the real-time face photos, a classifier is trained [9]. For biometric authentication, only live samples will be processed; fake inputs are immediately rejected.

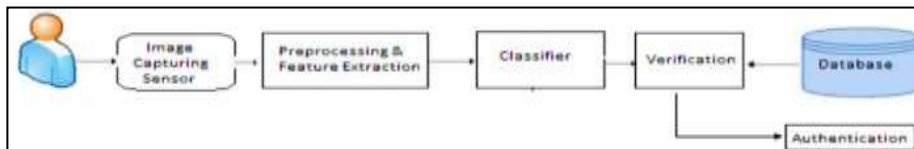


Fig. 1: Face Authentication Model

Preprocessing is done on the sensor-captured picture. Noise, illumination, attitude, and image quality variations may all affect face identification systems. A number of systems have used preprocessing to improve face detection efficiency [10]. Preprocessing often entails removing undesirable areas from the picture and, on occasion, normalizing the image in order to extract features [11].

The methods might be edge detection, scaling, sharpening, blurring, or smoothing. The feature extraction module then receives the preprocessed samples in order to extract the key characteristics that may distinguish real or living specimens from their fake equivalents. The literature presents a number of popular feature extraction methods. For many security levels, Conditional Random Field (CRF) and multi-model fusion techniques have been used. Anti-



spoofing hints are used both within and outside of a face. Spontaneous eye blinks are used as inside-face cues to prevent picture and 3D model faking [12]. Video replays are protected against spoofing by using the outside-face indications of scene context. The way the technology operates is non-intrusive. The technique described here records video snippets using a webcam. A clue to determine a person's liveliness is their blinking, which is a passive activity that doesn't need to alert the user like speaking or moving the face. In order to account for long-range contextual relationships within the observation series, the authors simulate blinking activity using CRFs. To identify facial liveliness, a novel method based on the thermal infrared spectrum was introduced. The canonical correlation analysis between the visible and thermal infrared faces was used by the model. In order to demonstrate additional correlative traits and advance live face identification capabilities, the correlation of various face portions is also examined [13, 14].

NEURAL NETWORK

A common calculation used to prepare feed-forward brain networks for administered learning is the back-proliferation calculation. There are theories of back-proliferation for different artificial neural networks (ANNs), and generally speaking, a class of computations is referred to be "back-spread" in the conventional sense. Back-propagation effectively calculates the gradient of the loss function with respect to the network weights for a single input-output example, as opposed to a naïve direct calculation of the gradient for each weight independently. Due to its effectiveness, gradient methods—of which gradient descent or its derivatives, such as stochastic gradient descent—can be used to train multilayer networks, updating weights to minimize loss. The gradient of the loss function for each weight is determined using the back-propagation method utilizing the chain rule. In order to avoid repeatedly doing the same computations for intermediate terms in the chain rule, it does this by calculating the gradient one layer at a time and iterating backward from the final layer [11].

However, the phrase "back-propagation" is sometimes used loosely to refer to the full learning method, including the application of the gradient, such as stochastic gradient descent, even though it only refers to the technique used to compute the gradient and not its application. Gradient computation is made more generic by back-propagation via the Delta rule, which is the single-layer form of back-propagation extended by automated



differentiation and in which back-propagation is a specific instance of reverse accumulation (also known as "reverse mode"). Figure 4's Back Propagation Neural (BPN) network serves as an illustration of how the word "back-propagation" is often used in neural networks.

The four components of the algorithm are as follows: Weight updates include backpropagation to the hidden layer, backpropagation to the output layer, and feed-forward calculation. When the value of the error function is sufficiently tiny, the algorithm is stopped. This is the simplest and most approximate formula for the BP algorithm. A few other definitions have been put forward by other scientists, but Rojas' description seems to be very clear-cut and correct. In the last stage, weights are updated continuously throughout the process [12].

Motivating Restoring Spread Organization

Finding a capacity that best directs a collection of data sources to the desired outcome is the aim of any directed learning computation. For example, in a classification job, the input might be a picture of an animal, and the right output would be the animal's name. The back-propagation technique was developed in order to find a way to train a multi-layered neural network so that it may learn the proper internal representations to allow it to learn any arbitrary mapping of input to output. Back-propagation aims to calculate the gradient, or partial derivative, of a loss function given any network weight [13].

Two phases make up the back-propagation learning algorithm:

1) Propagation 2) Weight revision

Phase 1: Propagation:- The following steps are involved in every back-propagation: The neural network's output value(s) are generated by the forward propagation of a training pattern's input through the network.

Phase 2: Weight update: The following steps must be taken for each weight: The weight's gradient is calculated by multiplying the weight's input activation and output delta. The weight is reduced by a ratio (percentage) of the gradient of the weight [14].

This proportion (rate) impacts the speed and nature of learning; the term for it is the learning rate.

The neuron trains more quickly if the ratio is higher, but the training is more accurate if it is lower. The indication of the angle of weight demonstrates whether the blunder fluctuates straightforwardly with, or conversely to, the weight. As a result, in order to "decline" the

gradient, the weight needs to be updated in the opposite direction. The first and second phases are repeated until the network performs satisfactorily [15].

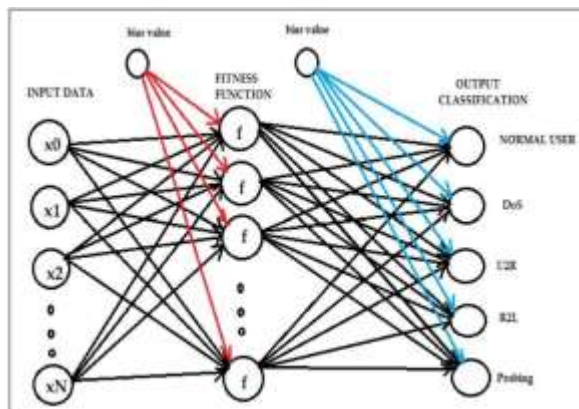


Fig. 2: Basic Diagram of Back-Propagation Neural Network

III. PROPOSED METHODOLOGY

The proposed technique is based on multilayer neural network. In this paper, explain the multilayer neural network and further data analysis using ML-NN & flow chart will explain result paper.

Multilayer Neural Network (ML-NN)

A feed-forward neural network called ML-NN is used for feature optimization in Figure 5. It is a single hidden layer network, which optimizes the features in the fewest possible characteristics by selecting them at random. By using random values for the delay and bias variables, ELM simplifies training. This is a fixed value. ML-NN's characteristic is feature reduction, which boosts learning speed and classification accuracy. The hidden layer in this ELM is given fixed random weights. In order to optimize features, I worked on an extreme learning machine here. This is a single-layered feed-forward neural network (SLFFNS), however there is no need to adjust the hidden layer (also known as feature mapping) in ELM [14]. In ML-NN, input weights, biases, and hidden node learning parameters are all randomly given and do not need tuning. ML-NN is outperforming both the conventional method and ML-NN. SLFFN, or single layer feed-forward neural networks, in this ML-NN algorithm are used for classification, regression, clustering, sparse approximation, compression, and feature learning with one or more layers of hidden nodes. It is not necessary to adjust the parameters of the hidden nodes (just the weights connecting inputs to hidden nodes). These hidden nodes may either be inherited from their predecessors unaltered

or they can be randomly assigned and never updated (i.e., they are random projections with nonlinear transformations). Typically, hidden node output weights are learnt in a single step, which is equivalent to learning a linear model.

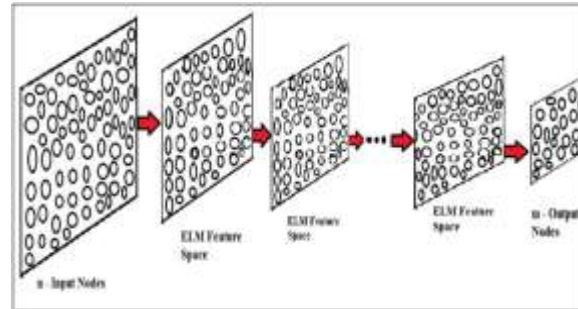


Fig. 3: Feature optimization in the ML-NN working model

IV. SIMULATION RESULT

Step 1	Importing the libraries and packages
Step 2	Initializing the parameters: Batch size, Number of training epochs, Number of filters, Size of filter, Pool size
Step 3	Reading the path of input files and initialize the output folder for data after pre-processing
Step 4	Pre-processing the images for giving them as the input to the model
Step 5	Converting the images to matrix form; flattening each image into an array vector and storing them in the common image matrix
Step 6	Assigning the labels to the image classes
Step 7	Shuffling the data to prevent overfitting and generalization of training
Step 8	Separating the train data and test data
Step 9	Normalizing the data
Step 10	Defining a model and its respective layers
Step 11	Compiling the model
Step 12	Fitting the data into the compiled model, i.e., training the model using the initially defined parameters
Step 13	Plotting the Loss and Accuracy curves of the training process
Step 14	Print the Classification Report and Confusion Matrix of the training process

Real image with confidence



Fig. 4: Image_1 Real Confidence

Fake image with confidence

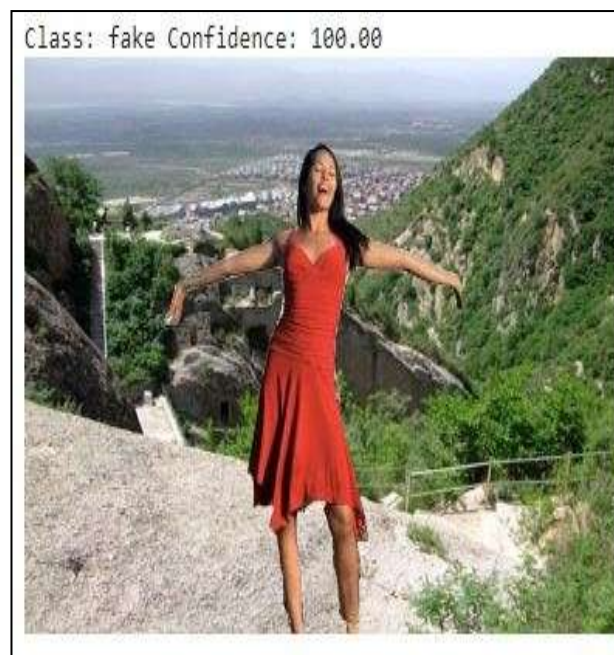


Fig. 5: Image_1 Fake Confidence

Real image with confidence



Fig. 6: Image_2 Real Confidence
Fake image with confidence



Fig. 7: Image_2 Fake Confidence
Real image with confidence



Fig. 8: Image_3 Real Confidence
Fake image with confidence



Fig. 9: Image_3 Fake Confidence

V. CONCLUSION

This research reviews a neural network-based approach for detecting fake faces. The

machine learning framework receives the picture as input. It is simple and effective for the system to distinguish between a genuine and fake picture when an input image is provided since we trained the fake and real face datasets independently. The primary cause of the improved performance in identifying the fake pictures is the neural layers' adaptability.

REFERENCES

- [1] Jarosz, M.; Nawrocki, P.; Śnieżyński, B.; Indurkha, B. Multi-Platform Intelligent System for Multimodal Human- Computer Interaction. *Comput. Inform.* 2021, 40, 83–103.
- [2] Prathiba, T.; Kumari, R.S.S. Content based video retrieval system based on multimodal feature grouping by KFCM clustering algorithm to promote human–computer interaction. *J. Ambient. Intell. Humaniz. Comput.* 2021, 12, 6215–6229.
- [3] Ince, G.; Yorganci, R.; Ozkul, A.; Duman, T.B.; Köse, H. An audiovisual interface-based drumming system for multimodal human–robot interaction. *J. Multimodal User Interfaces* 2020, 15, 413–428.
- [4] Wu, D.; Zhang, J.; Zhao, Q. Multimodal Fused Emotion Recognition About Expression-EEG Interaction and Collaboration Using Deep Learning. *IEEE Access* 2020, 8, 133180–133189.
- [5] Mosquera-DeLaCruz, J.H.; Loaiza-Correa, H.; Nope- Rodríguez, S.E.; Restrepo-Girón, A.D. Human-computer multimodal interface to internet navigation. *Disabil. Rehabil. Assist. Technol.* 2021, 16, 807–820.
- [6] Nayak, S.; Nagesh, B.; Routray, A.; Sarma, M. A Human– Computer Interaction framework for emotion recognition through time-series thermal video sequences. *Comput. Electr. Eng.* 2021, 93, 107280.
- [7] Liu, X.; Zhang, L. Design and Implementation of Human- Computer Interaction Adjustment in Nuclear Power Monitoring System. *Microprocessors and Microsystems. Microprocess. Microsyst.* 2021, 104096.
- [8] Dybvik, H.; Erichsen, C.K.; Steinert, M. Demonstrating the feasibility of multimodal neuroimaging data capture with a wearable electroencephalography + functional near-infrared spectroscopy (eeg+fnirs) in situ. *Proc. Des. Soc.* 2021, 1, 901–910.
- [9] Iio, T.; Yoshikawa, Y.; Chiba, M.; Asami, T.; Isoda, Y.; Ishiguro, H. Twin-Robot Dialogue System with Robustness against Speech Recognition Failure in Human- Robot Dialogue



with Elderly People. *Appl. Sci.* 2020, *10*, 1522.

- [10] Robert, L.P., Jr.; Bansal, G.; Lütge, C. ICIS 2019 SIGHCI workshop panel report: Human–computer interaction challenges and opportunities for fair, trustworthy and ethical artificial intelligence. *AIS Trans. Hum.-Comput. Interact.* 2020, *12*, 96–108.
- [11] Cao, Y.; Geddes, T.A.; Yang, J.Y.H.; Yang, P. Ensemble deep learning in bioinformatics. *Nat. Mach. Intell.* 2020, *2*, 500–508.
- [12] Panwar, H.; Gupta, P.; Siddiqui, M.K.; Morales- Menendez, R.; Singh, V. Application of deep learning for fast detection of COVID-19 in X-rays using nCOVnet. *Chaos Solitons Fractals* 2020, *138*, 109944.
- [13] Shen, C.-W.; Luong, T.-H.; Ho, J.-T.; Djailani, I. Social media marketing of IT service companies: Analysis using a concept-linking mining approach. *Ind. Mark. Manag.* 2019, *90*, 593–604.
- [14] Duan, H.; Sun, Y.; Cheng, W.; Jiang, D.; Yun, J.; Liu, Y.; Liu, Y.; Zhou, D. Gesture recognition based on multi- modal feature weight. *Concurr. Comput. Pract. Exp.* 2021, *33*, e5991.
- [15] Gurcan, F.; Cagiltay, N.E.; Cagiltay, K. Mapping Human– Computer Interaction Research Themes and Trends from Its Existence to Today: A Topic Modeling-Based Review of past 60 Years. *Int. J. Hum.-Comput. Interact.* 2021, *37*, 267–280.