



SMART WAY SENTIMENT ANALYSIS IN SOCIAL MEDIA

Jeba Stanley.J¹, Maruthapandi.M², Karthickraja.P³, Sivaprakash.S⁴, Herosingh.C⁵

¹Assistant Professor, Department Of Computer Science And Engineering

^{2,3,4,5}UG Students Department of Computer Science and Engineering

N.S.N College of Engineering and Technology, Karur, Tamil Nadu, India.

ABSTRACT

Social media provides a low-tech means of communication for people to express their thoughts, feelings, and attraction. The paper's goal is to extract different sentimental behaviors that will be utilized to help decide on a course of action and classify people's feelings and affections as neutral, contradictory, or clear. Noise removal was used to pre-process the data in order to eliminate noise. The research project used a variety of methodologies. The popular classification techniques were used to extract the sentiment once the noise was removed. The data was classified using multi-layer perceptions (MLP) and convolutional neural networks (CNN). These two classification results were compared against other techniques such as support vector machines (SVM), random forests, decision trees, Naïve Bayes, etc. using the sentiment classification from Twitter data and the consumer affairs website. In the proposed work, convolutional neural networks and multi-layer perceptions exceed other machine learning classifiers.

INTRODUCTION

Together with spoken language, non-verbal indicators used in interpersonal communication to convey emotions and provide feedback include hand gestures, facial expressions, and voice tones. However, the latest developments in HCI—from the traditional mouse and keyboard to automated speech recognition software and customized interfaces for individuals with disabilities—do not fully capitalize on these important communicative skills, which frequently leads to a less than natural interaction. If computers could recognize these emotional inputs, they could give specific and appropriate help to users in ways that are more in tune with the user's needs and preferences. Over the past two decades, machine learning has become one of the mainstays of information technology and, with that, a rather central, albeit usually hidden, part of our lives. With ever-increasing amounts of data becoming available, there is good reason to believe that smart data analysis will become even more pervasive as a necessary ingredient for technological progress. Learning, like intelligence, is a multifaceted process that is difficult to define precisely. Terms like "to gain knowledge, understanding of, or skill in, by study, instruction, or experience" and "modification of a behavioural tendency by experience" are included in dictionary definitions. Psychologists and zoologists research how humans and animals learn. We concentrate on machine learning in this book. There are numerous similarities between machine learning and animals. Undoubtedly, a great deal of machine learning methodology comes from psychologists trying to use computational models to refine their theories of animal and human learning. Furthermore, it appears possible that some aspects of biological learning may be clarified by the ideas and methods machine learning researchers are investigating.

EXISTING SYSTEM

The majority of existing systems only employ one of the following—voice or face—to determine a person's emotional state. These algorithms typically don't account for the fact that people can still somewhat fake their facial expressions and neglect other body language clues that could indicate their emotional state. The fact that those devices can only identify fundamental emotions like joy, sorrow, etc. is another drawback. They only consider one of the aforementioned stimuli. They use raw image processing to understand the expressions of the person on his face, not his emotional state. As such, there are many limitations to this approach. As I already said, it's rather easy to fake or cheat with such an approach. Another disadvantage is that such devices are rather inflexible because they are designed based on pre-existing fundamentals. It doesn't take into consideration a person's cultural aspects that might change the way he behaves, and as a result, it is impractical to sample on a scale. In the existing system, support vector machine kernel identification with complete facial emotion recognition is done. This research work analyzes the performance of four different SVM kernels (Radial Basis Function, Linear Function, Quadratic Function, and Polynomial Function) for face emotion recognition. A database of 714 face emotion images was created by capturing twice the seven facial expressions of 51 people with a digital camera.



Disadvantages

- Facial emotions are defined as a positive or negative experience linked to a specific pattern of physiological activity.
- They are characterized by a coordinated set of responses that may involve verbal, physiological, behavioral, or mechanisms.
- Typically, emotions are expressed through the face, speech, or body language.

PROPOSED SYSTEM

Preprocessing of the input photos, feature extraction, training, classification, and database comprise the stages of the architecture that is suggested in this work. Preprocessing of the input photos involves cropping and face detection. Finding distinctive characteristics in the data is a process known as feature extraction, which can be carried out using certain methods such as principal component analysis, feature averaging, etc. The collected features will be fed into the neural network together with predetermined network parameters to train the network. The neural network will classify data per the network's defined targets. Face photos are extracted from a collection of expressions on the face. In the original photograph, there is also a camera model and time. The face is identified, cropped, and saved as a distinct image for improved performance. Next, features are extracted from the clipped image. The neural network will be trained to acquire knowledge using these features as input. In addition, the testing image will undergo preprocessing, from which features will be retrieved and fed into the neural network. The neural network's classifier will categorize the input test image's expression.

MODELLING AND ANALYSIS

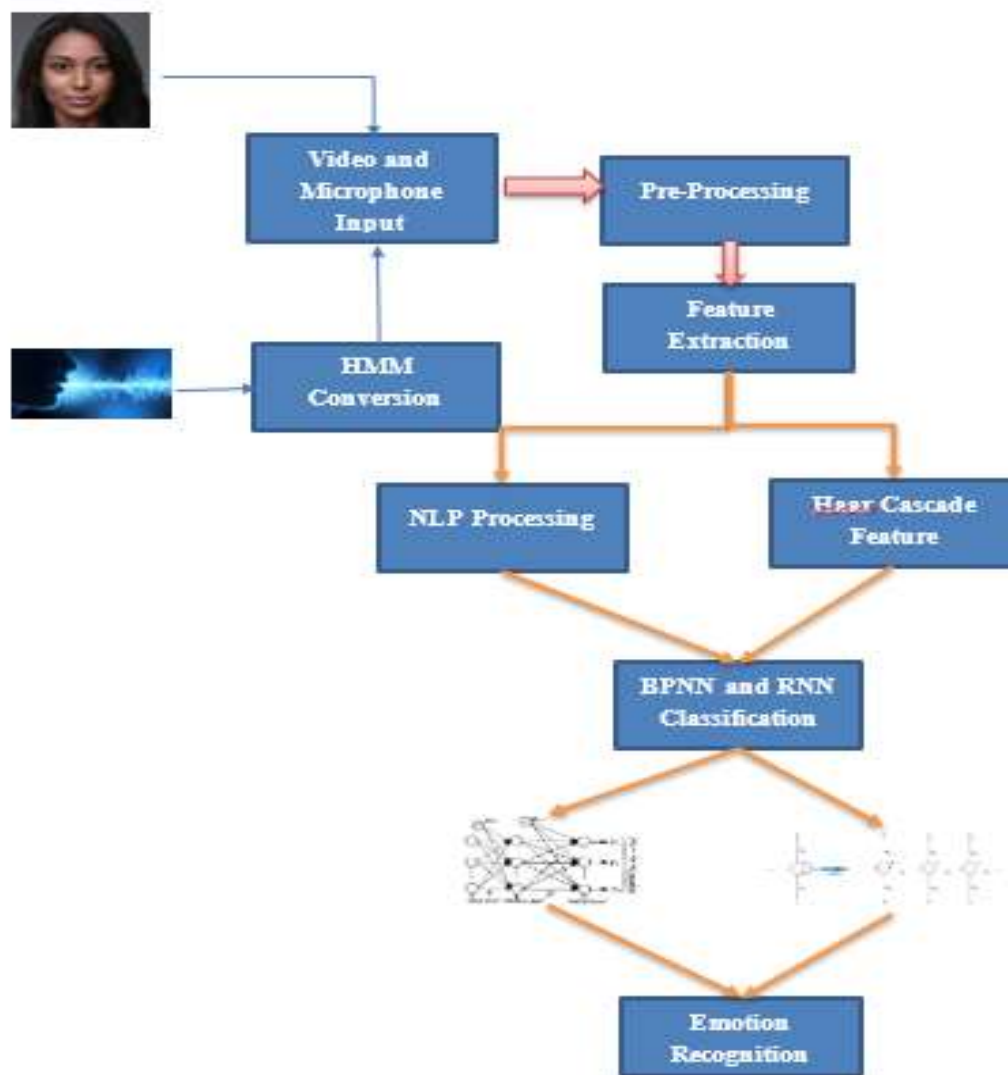
A structure made up of parts and rules describing how these parts interact with one another is called a system architecture. The foundation for visualizing a product description is a system architecture diagram, which provides an overview of the product structure. Three different kinds of system designs are recognized: mixed (which is a combination of partially distributed and partially integrated components). It is demonstrated that the kind of architecture is determined by the kind of interface. There are more interfaces in integrated systems, and they are also not well-defined. An application architecture diagram aids in the identification of applications, sub-applications, components, databases, services, etc. as well as their relationships. It offers a high-level graphical representation of the application architecture. Refer to Business and Application Services (EAM). System: An application that is bundled.

If you restrict yourself to four layers, they may be determine as

- 1) Algorithm
- 2) Programming language or compiler
- 3) processor/memory

I/O, Other abstraction definitions may contain three layers

- 1) Application
- 2) System software
- 3) Hardware

**Fig 1.1 : System Architecture**

Video and Microphone Input

Transducers, like microphones, transform sound waves into electrical signals. Telephones, hearing aids, public address systems for concerts and public gatherings, film production, live and recorded audio engineering, sound recording, two-way radios, megaphones, and radio and television transmission are just a few of the numerous uses for microphones. They are also utilized by computers for knock sensors, ultrasonic sensors, speech recognition, and voice recording. Today, a variety of microphone types are in use, each using a unique technique to translate a sound wave's fluctuations in air pressure into an electrical signal. The most popular types are the condenser microphone, which employs a vibrating diaphragm as its microphone, and the dynamic microphone, which uses a coil of wire hung in a magnetic field. Two types of piezoelectric crystals are used in contact microphones: a capacitor plate and a crystal. To capture or reproduce a signal, microphones usually require a pre-amplifier to be connected to them.

HMM Conversion

A group of finite states joined by transitions is known as a hidden Markov model. Two sets of probabilities define each state: a transition probability and an output probability density function that, depending on the state, defines the condition probability of emitting each output symbol from a finite alphabet or a continuous random vector. The Hidden Markov Model (HMM) is utilized in this project's execution to process speech system input and convert it to text. Here, the user's details will be used to convert the text pattern to detect the signal. The state space of the hidden variables in the hidden Markov models discussed above is discrete, but the observations themselves are not. Can be either continuous (usually from a Gaussian distribution) or discrete (usually created

from a categorical distribution). or continuous; such distributions are usually Gaussian. It is also possible to expand hidden Markov models to include continuous state spaces. These models include those in which all hidden and observable variables have a Gaussian distribution and the Markov process over hidden variables is a linear dynamical system with a linear link among related variables.

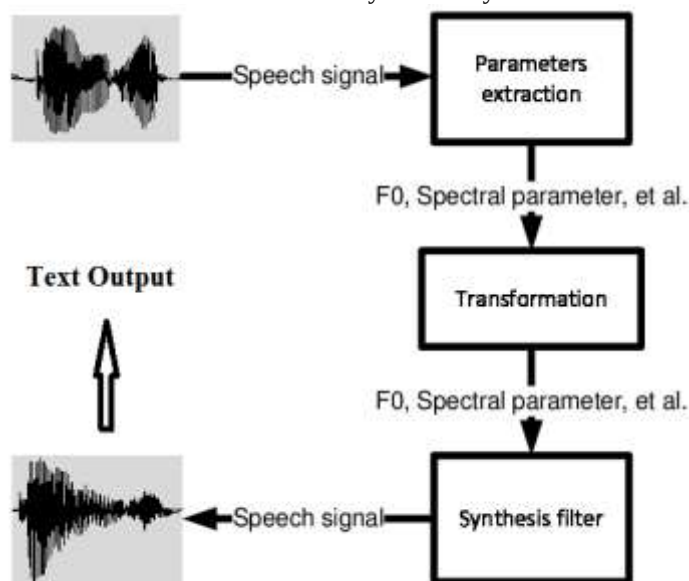


Fig 1.2: Speech Processing

Emotion Recognition

Emotion recognition has become an important field of research in the human-computer interaction domain. The latest advancements in the field show that combining visual and audio information leads to better results if compared to using a single source of information separately. From a visual point of view, a human emotion can be recognized by analyzing the facial expression of the person. More precisely, human emotion can be described through a combination of several facial action units.

RESULTS AND DISCUSSION

Since humans are multimodal learners, it is expected that automatic multimodal systems will perform better than automatic unimodal systems. The findings presented in this work support this hypothesis, as the bimodal approach produced an almost five percent (absolute) improvement over the facial expression recognition system's performance. Pairs of emotions that performed poorly in one modality were easily classified in the other; for instance, the facial expression emotion classifier was able to distinguish between happiness and anger, which were typically misclassified in the acoustic domain. Consequently, these emotions were precisely categorized when these two modalities were combined at the feature level. Sadly, sadness performed poorly since it is mistaken for a neutral mood in both areas. An examination of the confusion matrices for both feature-level and decision-level bimodal classifiers shows that, despite their similar overall performance, the recognition rates for each type of emotion were very different. Compared to the best unimodal identification system—the facial expression classifier—every emotion's recognition rate increased in the decision-level bimodal classifier (except happiness, which declined by 2%). Anger and neutral states were recognized at much higher rates in the feature-level bimodal classifier. Still, there was a 9% decline in the happy recognition rate. As a result, the application will determine the optimal method for combining the modalities. The study's findings show that, despite the audio-based system's lower performance compared to the facial expression emotion classifier, its characteristics nevertheless contain insightful emotional information that cannot be gleaned from the visual data.

These findings concur with those of Chen et al., who demonstrated that complementing information is presented by facial expression and auditory data. However, it makes sense to assume that some distinguishing emotional patterns can be identified by using either visual or auditory cues. When one of the models' attributes is incorrectly acquired, this redundant information is extremely helpful in enhancing the effectiveness of the emotion identification system. For instance, a person's facial expressions will be extracted with a high degree of error if they wear spectacles, have a beard, or both. Then, the visual information's limitations can be solved by using auditory features.

CONCLUSION

We investigate the distribution laws of several emotional signal features in the experiment by analyzing and contrasting time, amplitude energy, basic frequency, and formant feature parameters under various emotional states. We categorize five different emotional states based on this: calm, grief, happiness, surprise, and rage. The identification outcomes demonstrate that we can initially identify fundamental emotional categories using this basic prosodic information, and we can then apply these categories to



the emotion recognition system, which has a limited store and computation capacity and a loose recognition accuracy. In the meantime, the prosodic characteristics identify emotional categories with bi-modality and a greater identification rate by integrating with facial expression data. The rate of recognition hasn't increased much, despite an improvement in the emotional recognition performance when combined with facial expressions. This is mostly because the facial expression changes in the adjacent video frames have a similar association when it comes to acquiring emotional information. However, when we captured the instant face image for independent analysis, we did not account for this correlation. Conversely, the position and form of the organs on the face will alter in response to changes in the expression. While the Gaussian mixture algorithm-based image analysis method in this study has a greater identification rate for face contours, it does not provide a precise characterization of changes in the eyes, nose, mouth, and other facial organs. Based on the aforementioned two arguments, we must create a correction model linked to the expressions that contain a range of rules and use the model to alter the picture recognition results in order to genuinely improve the system's performance. Furthermore, the effectiveness of the recognition algorithm is a crucial component in real-time applications, in addition to augmenting the system's robustness and accuracy. Techniques like data compression and codebook reduction can also successfully raise the recognition rate. Future developments in human-computer interaction will inevitably lead to the creation of multi-modal recognition systems that integrate with voice, images, and other emotional information.

REFERENCES

1. Dellaert, F., Polzin, T., Waibel, A. Recognizing emotion in speech. *Spoken Language*, 1996. ICSLP 96. Proceedings. Fourth International Conference on, Volume: 3, 3-6 Oct. 1996. Pages: 1970 - 1973 vol.3.
2. De Silva, L. C., Miyasato, T., and Nakatsu, R. Facial Emotion Recognition Using Multimodal Information. In *Proc. IEEE Int. Conf. on Information, Communications and Signal Processing (ICICS'97)*, Singapore, pp. 397-401, Sept. 1997.
3. Chen, L.S., Huang, T.S. Emotional expressions in audiovisual human computer interaction. *Multimedia and Expo*, 2000. ICME 2000. 2000 IEEE International Conference on, Volume: 1, 30 July-2 Aug. 2000. Pages: 423 - 426 vol.1.
4. Cowie, R., Douglas-Cowie, E., Tsapatsoulis, N., Votsis, G., Kollias, S., Fellenz, W., Taylor, J.G. Emotion recognition in human-computer interaction. *Signal Processing Magazine, IEEE*, Volume: 18, Issue: 1, Jan 2001, Pages: 32 - 80.
5. R. Jafri, H. R. Arabnia, "A Survey of Face Recognition Techniques", *Journal of Information Processing Systems*, Vol.5, No.2, June 2009.