

SPECIAL FEATURE

Automatic Image
Captioning

Silicon

...beyond teaching

The Science & Technology Magazine

Wired Digest

Vol. 17 • Issue 1 • March - May 2018

Our Vision: "To become a center of excellence in the fields of technical education & research and create responsible citizens"

Nuclear Pasta

The strongest material in the universe discovered till date is known as “Nuclear Pasta”. The material that exists deep inside the crust of neutron stars is found to be the strongest material by a team of scientists at McGill University and his colleagues from Indiana University and the California Institute of Technology. Neutron stars are thought to have been formed by the gravitational collapse of the remnant of a massive star after a supernova explosion, provided that the star is insufficiently massive to produce a black hole. At the surface, the pressure is low enough that conventional nuclei, such as helium and iron, can exist independently of one another, and are not crushed together due to the mutual Coulomb repulsion of their nuclei. However, at the core, the pressure is so great that this Coulomb repulsion cannot support individual nuclei, and some form of ultra-dense matter, such as the theorized quark–gluon plasma, should exist. Since the attractive nuclear force that can fuse nuclei together is short ranged, the repulsion of like positive charges must be overcome to get nuclei close enough to induce fusion. This high density material is named such, as they have a unique structure and forces between the protons and neutrons cause them to assemble in the shape of “lasagna” and “spaghetti” named after different shapes of pasta. Together, the enormous densities and strange shapes make this nuclear pasta incredibly stiff. Breaking the stuff requires 10 billion times the force needed to crack steel, for example, researchers report in an article accepted to be published in Physical Review Letters.

Theoretically nuclear pasta has different phases which exist in the inner crust of the neutron stars forming a transitional phase between conventional matter at the surface and ultra-dense matter at the core. At the top of the transitional phase it experiences high pressure; conventional nuclei will be condensed into much more massive semi-spherical collections. These formations would be unstable outside the star, due to their high neutron content and size, which can vary between tens and hundreds of nucleons. This semi-spherical phase is known as the gnocchi phase. The next phase is known as spaghetti phase where the gnocchi phase is compressed where they are crushed into long rods and may contain thousands of nucleons immersed in neutron liquid, called as spaghetti. Further compression causes the spaghetti phase rods to fuse, and form sheets of nuclear matter. This is called the lasagna phase. Scientist Matthew Caplan ran the largest computer simulations which required 2 million hours’ worth of processor time or the equivalent of 250 years on a laptop with a single good GPU; he and his colleagues were able to stretch and deform the material deep in the core of neutron stars. According to him their results are valuable for astronomers who study neutron stars and interpret astronomical observations of these stars and this study could help astrophysicists better understand gravitational waves such as those detected last year when two neutron stars collided. Their new results even suggest that lone neutron stars might generate small gravitational waves.

Dr. Lopamudra Mitra

Dept. of EEE

Voice-based Information Extraction System

Abstract : Designing intelligent expert systems capable of answering different human queries is a challenging and emerging area of research. A huge amount of web resources are available these days and majority of which are in the form of unstructured documents covering articles, corporate reports, online news, medical records, social media communication data, etc. A user in need of certain information has to assess all the relevant documents to obtain the exact answer of their queries which is a time consuming and tedious work. Also, sometimes it becomes quite difficult to obtain the exact information from a list of documents quickly as and when required. This work aims to designing an intelligent information extraction system, which access the document contents quickly and provide the relevant answers to the user queries in a structured format just like a human expert answers to the questions.

Keywords: Artificial Intelligence, Automatic Speech Recognition, Speech Processing, Natural Language Processing, Information Extraction

I. INTRODUCTION

In the last few years, there has been a major change in the technology increasing the requirement of different sophisticated information processing systems. This leads to the development of intelligent systems based on Human Computer Interaction (HCI) technology. Research in the area of speech and language processing enables machines to speak and understand natural languages, leading to the development of different essential and luxury products. The field of Artificial Intelligence added new features to the HCI technology, where the system perceives its environment and takes actions intelligently that maximize its chance of successfully achieving the goals. There have been sufficient successes today in these areas having a widespread area of applications in designing different human-computer interactive systems such as talking computer systems, question answering systems, expert systems and information retrieval systems [1].

A huge amount of data is available in the web resources. A significant part of such information is in the form of unstructured document and thus it becomes difficult to find the exact information from a list of documents at the hand set unless all the documents are read. It is more than ever a key issue for knowledge management to develop intelligent tools and methods to give access to document content and extract relevant information quickly. This resulted in development of Information Extraction (IE) System for analyzing unstructured document and discovering valuable and relevant knowledge from it in the form of structured information [2]. The Information Extraction (IE) task can be expressed as to process a collection of unstructured texts which belong to different domains and producing cleaned relevant information to the users in a structured format. It involves processing

human language texts by means of Natural Language Processing (NLP) tools that enables the machines to understand human languages and interact with the human by producing natural language utterances.

Information Extraction (IE) is a rapidly growing field as much of the information in the web is expressed in natural languages. The efficient use of this information poses a major challenge for computer scientists in the industry. There are five main activities involved in the IE process: query processing, syntactic analysis, semantic analysis, textual information access, and structured representation [3]. The task of IE systems is to extract features such as name or location, and the relationships among the features. IE is based on the existence of implicit internal structures in natural languages that convert unstructured information into specific features and values [4].

This work focuses on development of an information extraction system, where the users can ask their queries and the system answers to the queries quickly like a human expert. The input queries are given in terms of spoken utterances, the system front end converts the queries into text utterances by using the speech to text conversion technology. The user queries are then semantically processed by using the NLP tools. The system provides the answers of the user's queries by referring a collection of documents in a knowledgebase which are domain specific.

II. PROPOSED MODEL FOR VOICE BASED INFORMATION EXTRACTION

This section presents the description about the development of an intelligent voice based IE system which answers to the user queries quickly by referring

a collection of documents. Often users queries are of the 'Wh' forms like- what, when, where, why, which, who, etc. The answers of those queries contain the information about the location, place, date, time event, people, etc. In the recent era of technology, users in need of certain information may obtain the details from different web resources easily. However, the web repositories are increasing day by day and most of the documents are unstructured. For obtaining exact answer from those documents requires assessing all related documents. Also, most of the time, users require quick answers to be obtained easily instead of surfing a large repository of documents particularly for the Wh-form of queries. This requires the development of expert question answering systems. Just like asking the questions to a human expert and getting the answers back immediately, the expert system answers the queries of the users. Figure (1) presents the overview of an intelligent question answering system for answering user's queries.

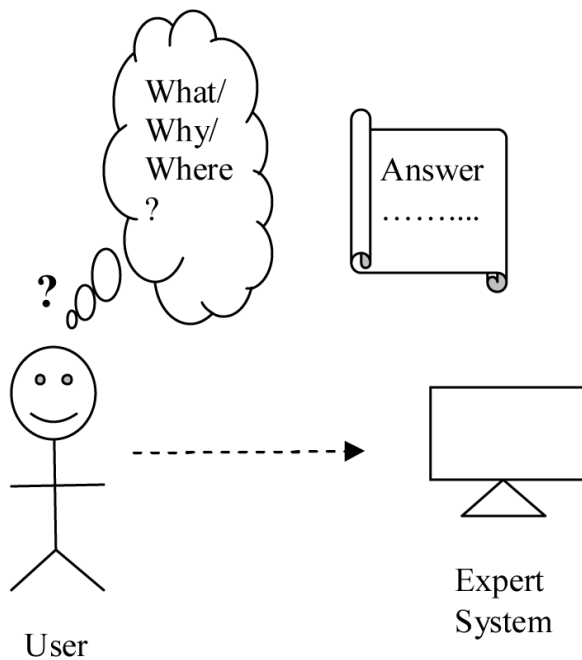


Figure-1: Overview of intelligent Question Answering System

The system is trained with a knowledgebase containing a collection of documents from different domains. There are two main components of the presented model: speech processing unit and information extraction unit. The front end speech processing unit processes the spoken user queries and converts it into the corresponding text format for further processing. The backend information extraction unit extracts the relevant information on the user's query by referring the knowledgebase and

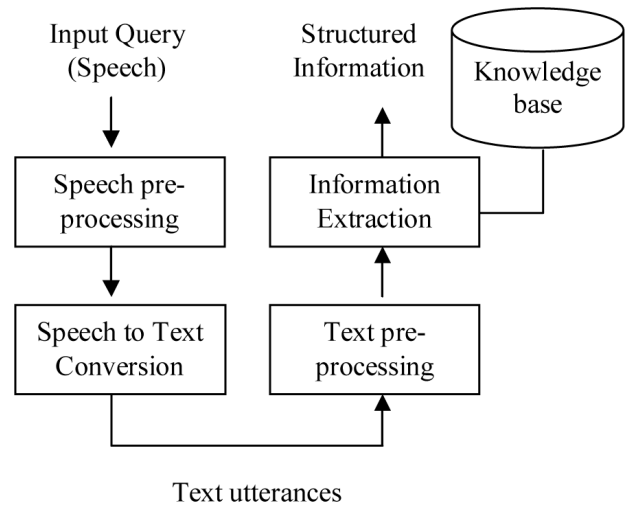


Figure-2: Phases of voice based IE System

provides the answer of the queries in a structured format. Figure (2) presents the phases involved in the voice based information extraction process.

III. SPEECH PROCESSING UNIT

The speech processing unit consists of two main phases: speech pre processing and speech to text conversion. During speech pre processing the user's queries in the form of spoken utterances are processed to remove noise and enhance the quality of the speech signal. The enhanced speech signal containing the user's query is then passed to the speech to text (STT) conversion phase. The STT module converts spoken words into respective written texts.

The STT module consist of 2 parts: DSP (Digital Signal Processing) interface at the front end and NLP (Natural Language Processing) interface at the back end. The input to the STT module is spoken words or sentences and the output is the corresponding text representations. The STT module relies on two important models: an acoustic model and a language model. In addition, large vocabulary systems are also used for determining the pronunciation model in any language.

IV. INFORMATION EXTRACTION UNIT

The information extraction unit processes the text utterances obtained by the STT conversion process by a text pre-processing process. Text pre processing involves removal of unnecessary words from text documents to obtain the relevant useful information for

further uses. This refers to the removals of stop words or too common words which don't contribute to the semantic meaning of the sentences. It also involves splitting of longer strings of text into smaller pieces, or tokens. Stemming is applied for eliminating the affixes (suffixed, prefixes, infixes, circumfixes) from the text in the user queries. The lemmatization process is applied on the processed text to capture canonical forms based on a word's lemma. Keywords are extracted and classified into named entities. Named Entity Recognition (NER) labels sequences of words in a text which are as per the pre-defined categories such as name of person, location, organizations, etc.

After getting the related information on the user's query, the answers to the query is produced by the IE phase in a determined feature and feature value form. The rules generated for the IE systems for narrative text are usually based on domain object recognition, syntactic analysis, and semantic grouping. The rules or patterns are hand coded and are generated from analysis of annotated training sets. The active learning methodology is used to identify a subset of the data that needs to be labeled and is performed by participation of human experts. The selective sampling method is used to annotate only the most important features in text which may contain the answers to the user's queries.

V. KNOWLEDGBASE CREATION AND FEATURE VALUE GENERATION

To create the knowledgebase for the system, a collection of articles covering different domains are considered and are represented in a standard format with respective class levels. The documents considered are medical records, social media interactions and streams, online news, government documents, corporate reports, stock exchange, gold price, petrol price etc.

Given set of documents $D = \{d_1, d_2, d_3, d_i, \dots, d_m\}$ and $S = \{s_1, s_2, \dots, s_j, \dots, s_j\}$ is the set of tokenized sentences and $W = \{w_1, w_2, \dots, w_k, \dots, w_p\}$ is the set of tokenized words of each sentences.

Algorithm (1) presents the steps involved in obtaining the features and feature values from the documents. For any new document not available in the document repository, sentence Tokenize function is applied to tokenize the whole text into list of sentences. Each sentence is then passed through word Tokenizer, where

sentences are tokenized into list of words. The tags which contains the possible answers for the Wh- questions such as date, location, etc are obtained by applying StanfordNerTagger on each tokenized sentence.

Algorithm 1:

- Step-1: For each new document not in D check availability in database
- Step-2: If the document is not available in database then apply sentenceTokenizer(di)
- Step-3: For each sentence s_j in d_i apply word Tokenizer(Sj)
- Step-4: Form word tokens wkfor each sj
- Step 5: Apply StanfordNERTagger on each wk set and find NER tags

VI. USER'S QUERY PROCESSING

On any user spoken query, the query is converted from speech to text for further processing by using the STT technology. The user's query is then tokenized into set of words using word Tokenizer. In order to obtain the answer to the query, the information about what needs to be obtained and what are the conditions which should be extracted from the query are needed to be identified. For this purpose, StanfordNerTagger is applied on the tokenized query to obtain the conditions related to the extracted features. To obtain the exact information on the user's query, the meaning of the 'wh' word is needed to be extracted. Therefore, a set of wh- form rules are used. The steps involved in feature value extraction process are presented in Algorithm (2).

Given Q be the Query and $Wt = \{w_1, w_2, \dots, w_q\}$ is the set of tokenized words in Q and answer is a list where answers are kept.

Algorithm 2:

- Step-1: On Query Q apply word Tokenizer(Q)
- Step-2: Form word tokens Wt for the query Q..
- Step-3: Apply StanfordNERTagger on Wt set and find out NER tags.
- Step-4: Apply whWordExtraction(Wt) to determine the purpose of the query.
- Step-5: if whWord is 'what' or 'when' then Combine 'DATE' with answers.

Step-6: else if whWord is 'where' then
Combine 'LOCATION' with answers

Step-7: else if whWord is 'how' then
Combine 'DEATH_TOLL' with answers

VII. STRUCTURED REPRESENTATION OF OUTPUT

After obtaining the required feature values to the user's wh-questions, the features and feature values are presented in a tabular form. The template of the structured output is determined based on the feature values obtained by the user's query. Table (1) shows a simplified example of extracted features and the features' values from text descriptions available in the web resources on organization of a conference.

Table.1:
Example of extracted features and feature values

Feature	Feature Value
Location	Bhubaneswar
Date	21 December 2017
Venue	Silicon Institute of Technology
Event	Conference

VIII. IMPLEMENTATION DETAILS AND RESULTS

The proposed model is implemented in Python, which is a general-purpose interpreted, interactive, object-oriented, and high-level programming language. Python emphasizes in code readability providing constructs that enable clear programming on both small and large scales. For applying the NLP tools for text query processing, the NLTK (Natural Language Toolkit) tool is used. NLTK, provides a collection of libraries and programs for symbolic and statistical natural language processing (NLP) for English written in the Python programming language. NLTK supports classification, tokenization, stemming, tagging, parsing, and semantic reasoning functionalities which are required for relevant information extraction from the repository. For the implementation of Named Entity Recognizer part to obtain the features from the text, the Stanford NER Tagger is used. Named Entity Recognition (NER) labels sequences of words in the text which are the names of things, such as person and company names, or gene and protein names, etc

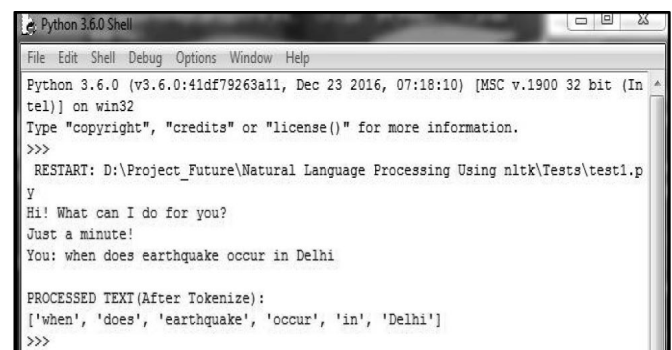
The spoken user's queries are converted to text by using the STT technology. Like any other pattern recognition technology, speech recognition cannot be 100% error free. The speech transcript accuracy is highly dependent on the speaker, the style of speech and the environmental conditions. To get accurate text transcriptions from the spoken user queries, the Google STT API is used for STT conversion process which achieves the highest accuracy rate by applying powerful neural network models in an easy-to-use API. Based upon the obtained text query, the required answer or information is extracted from the knowledge base and presented in a structured format.

The proposed model is tested on a number of user's queries asked by different speakers in spoken form and the results obtained are evaluated. In each of the tests conducted, the model successfully obtained the features and able to extract the feature values from the knowledgebase for that respective domain.

Example-1:

With document repository containing articles covering information on earthquakes in India, the obtained features are-

Date, Location, Magnitude and Death Toll which are derived as the fields in the output structured template. The speech to text conversion process on user's spoken query- "when does earthquake occurs in Delhi" is presented in Figure (3). Figure (4) presents the output of the text tokenization and NER tagging process on the same query. Figure (5) presents the output of structured template on the features and feature values for the user's query in a summary form covering multiple feature values.



```

Python 3.6.0 Shell
File Edit Shell Debug Options Window Help
Python 3.6.0 (v3.6.0:41df79263a11, Dec 23 2016, 07:18:10) [MSC v.1900 32 bit (In
tel)] on win32
Type "copyright", "credits" or "license()" for more information.
>>>
RESTART: D:\Project_Future\Natural Language Processing Using nltk\Tests\test1.p
y
Hi! What can I do for you?
Just a minute!
You: when does earthquake occur in Delhi

PROCESSED TEXT(After Tokenize):
['when', 'does', 'earthquake', 'occur', 'in', 'Delhi']
>>>

```

Figure-3: Output of speech to text conversion process

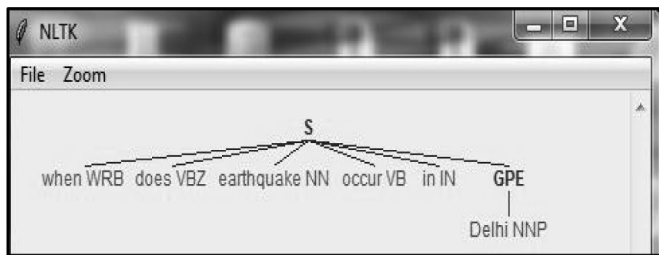


Figure-4: Output of text tokenization process

Example-2: Figure (6) presents the output of the voice based IE system on the wh- user's spoken query- "When does earthquake occurs in Bhuj" with single answer statement containing feature values generated from the system.

ID	LOCATION	DATE	MAGNITUDE	DEATH_TOLL
1	Assam, India	June 12 , 1897	8.7	6,000
2	Kangra, Punjab	1905 April 1905	7.8	7
3	Nepal, Bihar, India	1934 1934 January	8.0	8
4	Maharashtra, India	1967 December 1967	6.6	1967
5	Latur, India	1993 September	6.2	6
6	Jabalpur, India	May 22nd , 1997	5.8	887
7	Chamoli, India	March 29 , 1999	6.8	103
8	Bhuj, Gujarat, India	2001 January 26	7.7	13,805
9	Sumatra, India	December 26 , 2004	9.3	227,898
10	India, Kashmir	2005 October 8	7.6	7
11	Nepal, Sikkim	2011 Sunday September	6.9	111
12	Guwahati, Assam	June 28 , 2015	5.1	2
13	India, Manipur	2016 January 4	6.7	4
14	India, Tripura	January 5 , 2017	5.7	50
15	Uttarakhand	February 6 , 2017	5.1	5
16	Delhi, Noida	July 01, 2018 Sunday	4.0	5

Figure-5:

Structured output of all earthquake data from document repository.

```

Python 3.6.0 Shell
File Edit Shell Debug Options Window Help
Python 3.6.0 (v3.6.0:41df79263a11, Dec 23 2016, 07:18:10) [MSC v.1900 32 bit (Intel)] on win32
Type "copyright", "credits" or "license()" for more information.
>>>
RESTART: D:\Project_Future\Natural Language Processing Using nltk\Tests\VIE\sqlite3test.py
opened database successfully
ID      LOCATION      DATE      MAGNITUDE      DEATH_TOLL
1      Assam, India      June 12 , 1897      8.7      6,000
2      Kangra, Punjab      1905 April 1905      7.8      7
3      Nepal, Bihar, India      1934 1934 January      8.0      8
4      Maharashtra, India      1967 December 1967      6.6      1967
5      Latur, India      1993 September      6.2      6
6      Jabalpur, India      May 22nd , 1997      5.8      887
7      Chamoli, India      March 29 , 1999      6.8      103
8      Bhuj, Gujarat, India      2001 January 26      7.7      13,805
9      Sumatra, India      December 26 , 2004      9.3      227,898
10     India, Kashmir      2005 October 8      7.6      7
11     Nepal, Sikkim      2011 Sunday September      6.9      111
12     Guwahati, Assam      June 28 , 2015      5.1      2
13     India, Manipur      2016 January 4      6.7      4
14     India, Tripura      January 5 , 2017      5.7      50
15     Uttarakhand      February 6 , 2017      5.1      5
16     Delhi, Noida      July 01, 2018 Sunday      4.0      5
Operation done successfully
>>>
  
```

Figure-6: Output of user query on 'earthquake in Bhuj'

This work presents development of a voice based information extraction system which takes voice based user query as input and provides the relevant answer

of the query by processing a collection of related documents. The user queries in speech are processed and converted to respective text utterances. The relevant information to the user's query are then extracted from the document repository and represented in a structured format. The model is tested on different text documents covering different domains. However, the model may further be enhanced to work on multi domain online document repository instead of limited domain offline repository.

ACKNOWLEDGEMENT

We express our sincere gratitude to our project guide Dr. Soumya Priyadarsini Panda of CSE department for guidance and encouragement in carrying out this research work.

REFERENCES

- [1] R. Glauber, and D. B. Claro, "A Systematic Mapping Study on Open Information Extraction", Expert Systems with Applications, Elsevier, Vol. 112, pp. 372-387, 2018.
- [2] K. Rajbabu, H. Srinivas, and S. Sudha, "Industrial Information Extraction Through Multi-Phase Classification using Ontology for Unstructured Documents", Computers in Industry, Elsevier, Vol. 100, pp. 137-147, 2018.
- [3] G. Chen, C. Wang, M. Zhang, Q. Wei, and B. Ma, "How Small Reflects Large? -Representative Information Measurement and Extraction", Information Sciences, Elsevier, Vol. 460, pp. 519-540, 2017.
- [4] Y. Wang, L. Wang, M. Rastegar-Mojarad, S. Moon, F. Shen, N. Afzal, S. Liu, Y. Zeng, S. Mehrabi, S. Sohn, and H. Liu, "Clinical Information Extraction Applications: A Literature Review", Journal of biomedical informatics, Elsevier, Vol. 77, pp. 34-49, 2017

Alloran Pradhan
Varun Behera
Abhisekh Mohanty
 7th Sem. CSE

Density Based Smart Traffic Signal Control Using Micro-Controller And GSM Module

Abstract: The project is designed to develop a density based dynamic traffic signal system. The signal timing changes automatically on sensing the traffic density at the junction. Traffic congestion is a severe problem in many major cities across the world and it has become a nightmare for the commuters in these cities. Conventional traffic light system is based on fixed time concept allotted to each side of the junction which cannot be varied as per varying traffic density. Junction timings allotted are fixed. Sometimes higher traffic density at one side of the junction demands longer green time as compared to standard allotted time. The proposed system using a microcontroller of AT-MEGA 16 series duly interfaced with sensors, changes the junction timing automatically to accommodate movement of vehicles smoothly avoiding unnecessary waiting time at the junction. The sensors used in this project are IR and photodiodes are in line of sight configuration across the roads to detect the density at the traffic signal. The density of the vehicles is measured in three zones i.e., low, medium, high based on which timings are allotted accordingly. Further the project can be enhanced by synchronizing all the traffic junctions in the city by establishing a network among them. The network can be wired or wireless. This synchronization will greatly help in reducing traffic congestion.

Keywords: GSM, ATmega16, IR LED, AVR

I. INTRODUCTION

A steady increase in metro-city population, the number of automobiles and cars increases rapidly and metro traffic is growing crowded which leads to the traffic jam problem [1, 2]. Nowadays, controlling the traffic becomes major issue because of rapid increase in automobiles and also because of large time delays between traffic lights. So, in order to rectify this problem, we will go for density based traffic lights system. This article explains you how to control the traffic based on density. In this system, we will use IR sensors to measure the traffic density. We have to arrange one IR sensor for each road; these sensors always sense the traffic on that particular road. All these sensors are interfaced to the microcontroller. Based on these sensors, controller detects the traffic and controls the traffic system [3, 4].

The main heart of this traffic system is microcontroller. IR sensors are connected to the PORT C (PC0, PC2, PC4, and PC6) of the microcontroller and traffic lights are connected to PORT A and PORT B. If there is traffic on road then that particular sensor output becomes logic 0 otherwise logic 1. By receiving these IR sensor outputs, we have to write the program to control the traffic system. If you receive logic 0 from any of these sensors, we have to give the green signal to that particular path and give red signal to all other paths. Here continuously we have to monitor the IR sensors to check for the traffic. We have to place these IR pair in such a way that when we place an obstacle in front of this IR pair, IR receiver

should be able to receive the IR rays. When we give the power, the transmitted IR rays hit the object and reflect back to the IR receiver. Instead of traffic lights, you can use LEDs (RED, GREEN, YELLOW). In normal traffic system, you have to glow the LEDs on time basis. If the traffic density is high on any particular path, then glows green LED of that particular path and glows the red LEDs for remaining paths [5, 6].

II. OBJECTIVE

During our literature survey we came across many journal papers in which traffic is controlled with the help of microcontroller. In this manuscript, we are controlling traffic signal using microcontroller. It is density based traffic signal system. Here we are utilizing the concept of IR sensor and controlling the density of traffic. In this project with the help of command we control the microcontroller [7][8][9].

III. PROJECT OVERVIEW

The overview of this project is to implement Density based traffic control system using IR technology and ATmega16 microcontroller. The project density based traffic light control is an automated way of controlling signals in accordance to the density of traffic in the roads. IR sensors are placed in the entire intersecting road at fixed distances from the signal placed in the junction. The time delay in the traffic signal is set based on the density of vehicles on the roads.

The IR sensors are used to sense the number of vehicles on the road. According to the IR count, microcontroller takes appropriate decisions as to which road is to be given the highest priority and the longest time delay for the corresponding traffic light.

The GSM module used in the project helps in escape of emergency vehicles like ambulance, police cars and VIP cars. The flow chart is shown in fig 1.

IV. FLOW CHART

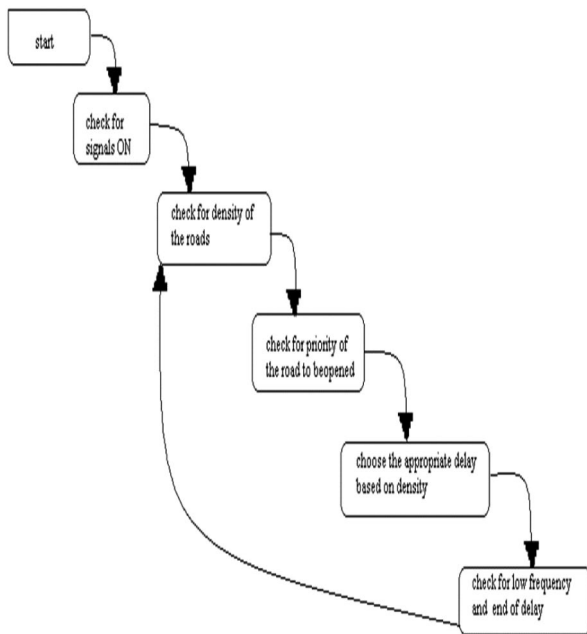


Fig.-1. Flowchart of the proposed work.

V. PROJECT DESCRIPTION

Traffic Congestion detection and Alert System use different hardware components and software to govern the system are as follows:

- Microcontroller (ATmega 16) Board
- GSM (SIM900) Module

A. Microcontroller(ATmega 16)

The ATmega16 is a low-power CMOS 8-bit microcontroller based on the AVR enhanced RISC architecture. By executing powerful instructions in a single clock cycle, the ATmega16 achieves throughputs approaching 1 MIPS per MHz allowing the system designed to optimize power consumption versus processing speed [8].

ATmega16 has 16 KB programmable flash memory, static RAM of 1 KB and EEPROM of 512 Bytes. The endurance cycle of flash memory and EEPROM is 10,000 and 100,000, respectively.

ATmega16 is a 40 pin microcontroller. There are 32 I/O (input/output) lines which are divided into four 8-bit ports designated as PORTA, PORTB, PORTC and PORTD.

ATmega16 has various in-built peripherals like USART, ADC, Analog comparator, SPI, JTAG etc. Each I/O pin has an alternative task related to in-built peripherals. The following table shows the pin description of ATmega16.

B. GSM (Sim 900) modem

This is a GSM/GPRS-compatible Quad-band cell phone, which works on a frequency of 850/900/1800/1900MHz and which can be used not only to access the Internet, but also for oral communication (provided that it is connected to a microphone and a small loud speaker) and for SMSs.

Externally, it looks like a big package (0.94 inches x 0.94 inches x 0.12 inches) with L-shaped contacts on four sides so that they can be soldered both on the side and at the bottom. Internally, the module is managed by an AMR926EJ-S processor, which controls phone communication, data communication (through an integrated TCP/IP stack), and (through an UART as shown in fig 2 and a TTL serial interface) the communication with the circuit interfaced with the cell phone itself. The processor is also in charge of a SIM card (3 or 1.8 V) which needs to be attached to the outer wall of the module [9]. The circuit diagram is shown in fig 3.

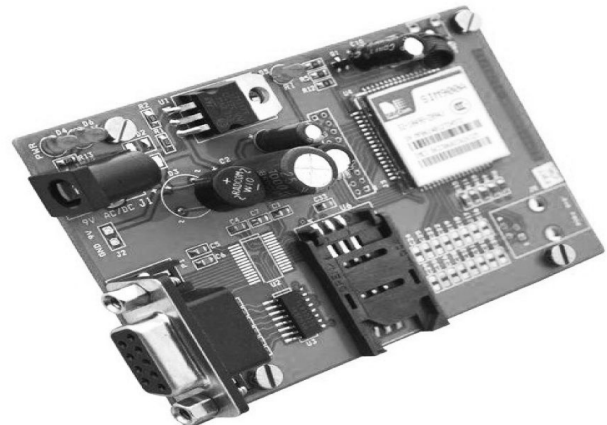


Fig. 2: GSM module

Smart Wheel Chair for Differently Abled People Using DTMF Control

Abstract–In this world a number of people are differently-abled. Their life revolves around wheels. This paper presents an approach for controlling wheelchair movement using (Dual Tone Multi Frequency) DTMF of a smart phone. The aim of DTMF controlling method is to introduce an automated ambulation tool, that helps the users to control their wheelchairs' position and move it via a smartphone to the desired destination. This application also provides the user with the ability to control the wheelchair even when sitting on it, moreover, the user can give orders to adjust the path. It even works under no internet facility. The encouraging results set a path for implementing this concept in more aided-systems.

Keywords– DTMF, Smartphone, Motor Driver, ARDUINO

I. INTRODUCTION

DTMF stands for Dual Tone Multi Frequency (DTMF). It is worked on methods of digital signal processing. Wireless-control of robots uses RF circuit that has the drawbacks of limited working range and limited control. This DTMF gives advantage over the RF. It increases the range of working and also gives good results in case of motion and direction of robot using mobile phone through microcontroller.

1.1 DTMF BASICS

DTMF, or dialing tone, is very commonly used. DTMF (Dual-tone Multi Frequency) is a tone composed of two sine waves of given frequencies. Individual frequencies are chosen so that it is quite easy to design frequency filters, and so that they can easily pass through telephone lines (where the maximum guaranteed bandwidth extends from about 300 Hz to 3.5 kHz). DTMF was not intended for data transfer; it is designed for control signals only. With standard decoders, it is possible to signal at a rate of about 10 “beeps” (=5 bytes) per second. DTMF standards specify 50ms tone and 50ms space duration. For shorter lengths, synchronization and timing becomes very tricky.

1.2 DTMF USAGE

DTMF is the basis for voice communications control. Modern telephony uses DTMF to dial numbers, configure telephone exchanges (switchboards), and so on. Occasionally, simple floating codes are transmitted using DTMF - usually via a CB transceiver (27 MHz). It is used to transfer information between radio transceivers, in voice mail applications, etc. Almost any mobile (cellular) phone is able to generate DTMF after establishing connection. If your phone can't generate

DTMF, you can use a stand-alone “dialer”. DTMF was designed so that it is possible to use acoustic transfer, and receive the codes using standard microphone as shown in fig1.

		High Frequency Group			
		1209 Hz	1336 Hz	1477 Hz	1633 Hz
Low Frequency Group	697 Hz	1	2	3	A
	770 Hz	4	5	6	B
	852 Hz	7	8	9	C
	941 Hz	*	0	#	D

Fig. 1: Representation of Keypad with associated frequency group

1.3 DTMF DECODER

DTMF signaling is used for telecommunication signaling over analog telephone lines in the voice-frequency band between telephone handsets and other communication devices and the switching centre. The DTMF system generally uses eight different frequency signals transmitted in pairs to represent sixteen different numbers, symbols and letters. When someone presses any key in the key pad of the handset, a DTMF signal is generated unique tone which consists of two different frequencies one each of higher frequency range. (>1 KHz) and lower frequency (<1 KHz) range. The resultant tone is convolution of two frequencies. The frequencies and their corresponding frequency are shown in Table 3.1. Each of these tones is composed of two pure sine waves of the low and high frequencies superimposed on each other. These two frequencies explicitly represent one of the digits on the telephone keypad.

Thus generated signal can be expressed mathematically as follows:

$$f(t) = A_H \sin(2\pi f_H t) + A_L \sin(2\pi f_L t)$$

Where:

A_H, A_L : are the amplitudes

f_H : high frequency range

f_L : low frequency range

The MT-8870 is a DTMF Receiver that integrates both band split filter and decoder functions into a single 18pin DIP package. It is manufactured using CMOS process technology. The MT8870 offers low power consumption (35 mW max) and precise data handling. Its decoder uses digital counting techniques to detect and decode all 16 DTMF tone pairs into a 4-bit code. The DTMF signal from the user mobile phone is picked up by the system mobile phone. The tip and ring of the microphone is connected to the specified

pin of MT-8870 as shown in the Fig. 2. C_1, R_1 and R_2 have been adjusted for gain control of the input signal. Resistance R_3 and capacitor C_2 has been used to set the “guard time” which is a time duration through which a valid DTMF tone must be present for its recognition. The “Q-test” signal (pin15) indicates that the valid DTMF tone has been detected. Increase the resistor

between pin2 and pin3 (not the one connects to 100nF) from 100K to 220k, 330K or 470K. This increases the input gain from 1 to 2.2, 3.3 or 4.7 to suit your input signal strength.

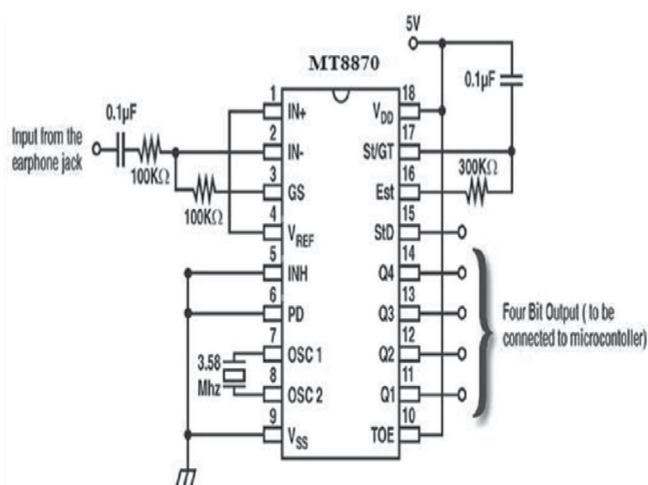


Fig.2:Pin Configuration of DTMF decoder.

1.4 OPERATION OF THE CIRCUIT:-

The DTMF signal from the user (Sender) mobile phone is picked up by the system (Receiver) mobile phone. Then the connection is established between the two phones, whatever phone key is pressed at the Sender mobile phone, the corresponding DTMF tone is heard in the ear piece of the receiver phone. Received DTMF tone is fed to the DTMF decoder. The DTMF decoder will give the corresponding BCD value of the tone. This Output is connected to Q4, Q3, Q2, Q1 pin of MT8870 Decoder IC and this output is fed to ATmega328P Microcontroller pin 4,5,6,7 respectively. Based on the equivalent binary digit of the DTMF tone received by the ATmega328P Microcontroller, a decision is made for Pins 8,9,10 and 11 regarding which pins should be high or low. These pins are fed to the L293D IC as input. Based on the Controller decision, the Pins are either high or low which activates the motors and moves the vehicle.

II. ARDUINO

2.1 INTRODUCTION

Arduino is an open source, computer hardware and software company, project, and user community that designs and manufactures microcontroller kits for building digital devices and interactive objects that can sense and control objects in the physical world.

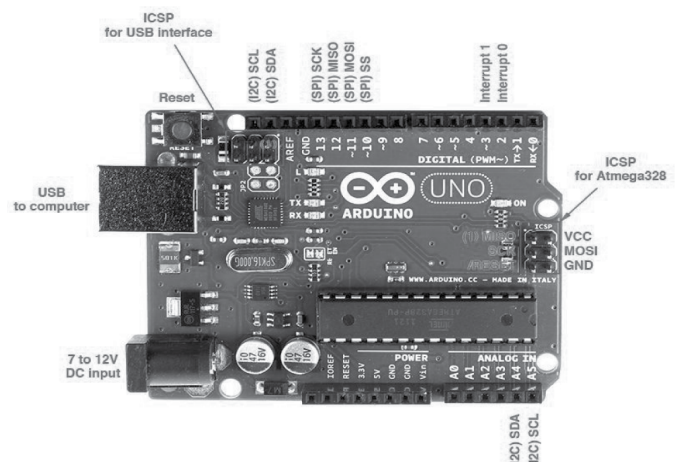


Fig. 3: Pin Diagram of Arduino

2.3 PIN DESCRIPTION

1. A typical example of Arduino board is Arduino Uno. It consists of ATmega328p- a 28 pin microcontroller.

2. Arduino Uno consists of 14 digital input/output pins (of which 6 can be used as PWM outputs), 6 analog inputs, a 16 MHz crystal oscillator, a USB connection, a power jack, an ICSP header, and a reset button
3. Power Jack: Arduino can be powered either from the PC through a USB or through an external source like an adaptor or a battery. It can operate on an external supply of 7 to 12V. Power can be applied externally through the pin V_{in} or by giving voltage reference through the I/O Ref pin.
4. Digital Inputs: It consists of 14 digital input/output pins, each of which provides or takes up 40mA current. Some of them have special functions like pins 0 and 1, which act as Rx and Tx respectively, for serial communication, pins 2 and 3, which are external interrupts, pins 3, 5, 6, 9, 11, which provide PWM output and pin 13 where LED is connected.
5. Analog inputs: It has 6 analog input/output pins, each providing a resolution of 10 bits.
6. ARef: It provides reference to the analog inputs
7. Reset: It resets the microcontroller when low. The pin description shown in fig 3.

Low Frequency Group	High Frequency Group	Digit	O E	D 3	D 2	D 1	D 0
697	1209	1	H	0	0	0	1
697	1336	2	H	0	0	1	0
697	1477	3	H	0	0	1	1
770	1209	4	H	0	1	0	0
770	1336	5	H	0	1	0	1
770	1477	6	H	0	1	1	0
852	1209	7	H	0	1	1	1
852	1336	8	H	1	0	0	0
852	1477	9	H	1	0	0	1
941	1209	*	H	1	0	1	0
941	1336	0	H	1	0	1	1
941	1477	#	H	1	1	0	0

2.4 MOTOR DRIVER (L298N)

The L298 is an integrated monolithic circuit in a 15-lead Multi watt and PowerSO20 packages. It is a high

voltage, high current dual full-bridge driver designed to accept standard TTL logic levels and drive inductive loads such as relays, solenoids, DC and stepping motors. Two enable inputs are provided to enable or disable the device independently of the input signals. The emitters of the lower transistors of each bridge are connected together and the corresponding external terminal can be used for the connection of an external sensing resistor. An additional supply input is provided so that the logic works at a lower voltage. The pin configuration is shown in fig 4.

Truth Table:

A	B	Description
0	0	Motor stops
0	1	Motor runs clockwise
1	0	Motor runs Anti-clockwise
1	1	Motor stops

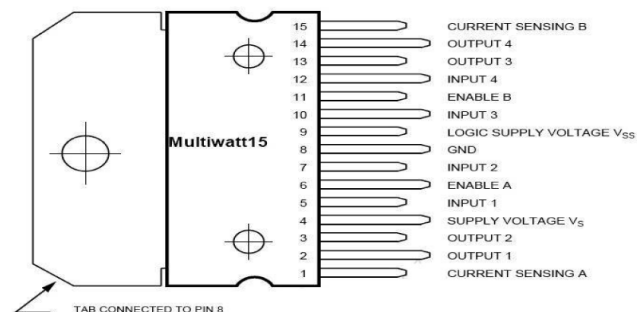


Fig.4: Pin configuration of L298N.

2.5 INSTRUMENTS REQUIRED

- DTMF Decoder module (MT-8870)
- Smartphone
- Aux wire
- Arduino UNO
- Battery Connector
- Wheelchair chassis
- 12v battery
- L298N DC Motor Driver
- 2 DC Motors

The circuit diagram of DTMF control is shown in fig 5. And the prototype is shown in fig 6.

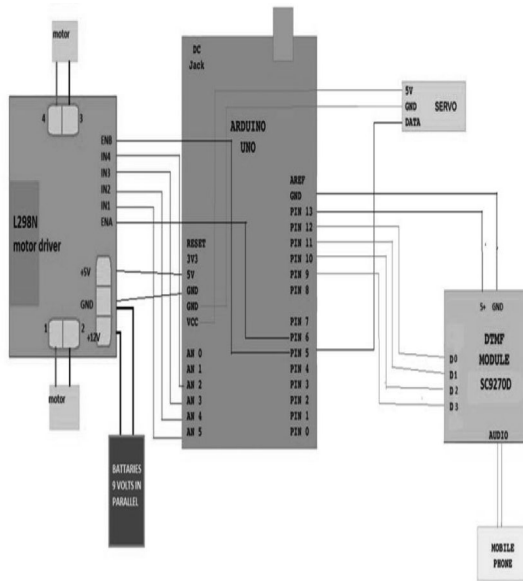


Fig.5: Circuit Diagram of DTMF Controlled Wheelchair



Fig.6: Prototype of the Smart Wheel Chair

CONCLUSIONS:

In this paper, an automated wheelchair has been developed which works according to dial pad of a smartphone. The robot moves wirelessly according to the button pressed on the cell phone. By developing

DTMF based robotic vehicle, we have overcome the drawbacks of RF communication which have a limited range whereas this wheelchair can be controlled from anywhere just using this DTMF technology. In this paper with the use of a mobile phone for robotic control can overcome these limitations. It provides the advantages of robust control, working range as large as the coverage area of the service provider, no interference with other controllers and up to twelve controls. Although the appearance and capabilities of wheelchair vary vastly, all wheelchairs share the features of a mechanical, movable structure under some form of control. This work is a step to make a smart wheelchair that can be controlled using a smartphone. It is dedicated to make the life of differently abled people simpler in wheels.

ACKNOWLEDGEMENT

We express our sincere gratitude to our project guide Dr. Lopamudra Mitra of EEE department for guidance and encouragement in carrying out this research work.

REFERENCES:

- [1] M. A. Goodrich and A. C. Schultz, "Human-robot interaction: a survey," *Foundations and Trends in Human-Computer Interaction*, vol. 1, no. 3, pp. 203–275, 2007.
- [2] F. Karray, M. Alemzadeh, J. A. Saleh, and M. N. Arab, "Human-Computer Interaction: overview on state of the art," *International Journal on Smart Sensing and Intelligent Systems*, vol. 1, no. 1, p. 137159, 2008.
- [3] F. Boniardi, A. Valada, W. Burgard, and G. D. Tipaldi, "Autonomous indoor robot navigation using a sketch interface for drawing maps and routes," in *Proceedings of the 2016 IEEE International Conference on Robotics and Automation (ICRA)*, Stockholm, Sweden, May 2016.
- [4] D. Wigdor and D. Wixon, *Brave NUI World: designing natural user interfaces for touch and gesture*. San Francisco, CA, USA: Morgan Kaufmann Publishers Inc., 2011.

Sanjeev Kumar Das
Shakti Mahapatra
Sourav
Satya Swarup Nayak
Ranjana Bharati

Dept. of EEE

Automatic Image Captioning

Abstract: Image captioning involves generating a human readable textual description, given an image, such as a photograph. Automatic Image captioning is the process of generating captions for images where no prior information is available. Image captioning is a field that has gained increased interest due to increase in processing power, improvement in neural networks and the advent of deep learning concepts. Image captioning has massive implications because the vast majority of images that are uploaded online or on mobile phones has no contextual information associated with the images. Moreover, real time caption generation combined with computer vision can be path breaking in terms of self-driving cars, robots, drones or security cameras. Image captioning show cases the true power of deep neural networks as it manages to generate descriptive captions with minimal information provided. The automatic image captioning task is divided into two sub-problems: object detection and caption generation. In this work, the Convolutional Neural Network (CNN) is used to encode the images and Recurrent Neural Network (RNN) is used to generate the captions.

Keywords: Deep Neural Network, Convolutional Neural Network, Recurrent Neural Network, Image captioning, Object detection

I. INTRODUCTION

Machine Learning (ML) has made many things possible that we couldn't have thought about just a decade ago. One of such applications of ML is in Image Caption Generation. Image Caption Generation is the process of generating the captions describing an image, given the image as input [1]. Caption Generation can be described as an application of two fields of computer science: Computer Vision and Natural Language Processing (NLP) [2]. Often with Computer Vision comes the complexity of detecting the various objects present in the image. We as humans have developed the eyes to see and the brains to understand the objects in our surroundings over an evolution of millions of years. Trying to replicate the same results from a computer system presents itself with many difficulties [3].

The greatest difficulty is the detection of objects present in the image. The objects come in all shape, size and color [4]. To be able to differentiate different objects or group similar objects comes with great difficulty. Being able to extract useful features from the objects is a really challenging task. This is where Deep Neural Network (DNN) come into play [5]. DNN specializes in complex function approximation [6]. DNNs can combine the features present in the input and use the combined features to detect complex features that human couldn't possibly fathom. How DNNs visualize the input is something beyond understanding: its complex and varies from input to input. Mapping the exact pathway using which a DNN takes the input and produces an output is not possible [7]. CNNs are specialized DNN architectures, especially built for processing visual and speech data. CNNs utilize convolution to produce a compact representation of the input [8]. By layering many convolution unit one after another, a really efficient and compact representation

of the input can be produced [9]. The second part of the problem, natural language processing is achieved using another ANN architecture call Recurrent Neural Networks (RNNs). RNNs process sequential data i.e. if the input is a sequence, where one input depends on the previous inputs, RNNs are used [10].

II. PROPOSED APPROACH

The problem is divided into two tasks: object detection and natural language resolution. To tackle the task of identifying objects inside an image, the Google Inception engine V3 is used. The inception engine is a 22 layered CNN that has been trained on millions of images to identify common object classes. The inception engine allows creating encoded versions of images which are essentially multidimensional vector. Each value of the vector is associated with some feature. The inception engine was chosen due to fact that not only it is accurate but it works well with low resolution images. The encoded images then form the basic building blocks of the model. The encoded images are used in conjunction with encoded captions to train a Recurrent Neural Network.

A. PRE PROCESSING

Pre-processing is the first step in the training process as, each image is of different resolution and raw image has too many values to be practically fed into a CNN. So, the first step is to resize the image to 299x299. The reason for choosing this particular resolution is that VGG16 started with this standard and this has been the accepted value for inception engine V3. The reason for re-sizing is that it reduces the amount of redundant information as an image of the size 299x299 can represent the same information as 1024x720 without any significant loss.

Moreover, images generally 3 channels representing RGB values, and this means that an image of resolution 1024x720 has 1024x720x3 values which is cumbersome. The other benefit is that consistent feature vectors for all images are obtained regardless of the images size.

B. ENCODING WITH INCEPTION

As raw images cannot be used as features due to multiple reasons, encoding is the process of extracting the relevant features (or combination of features) of images such that the vital information contained in the image is presented in a systematic usable form. The process of encoding is done with the help of Google's Inception engine. The inception engine, as it is called, is a multi-layered CNN which was developed at Google to provide state of the art performance on the ImageNet Large-Scale Visual Recognition Challenge and to be more computationally efficient than its competitor architectures. The main idea of the Inception architecture is based on finding out how an optimal local sparse structure in a convolutional vision network can be approximated and covered by readily available dense components.

C. PRE PROCESSING CAPTIONS

Each image is associated with five captions in the dataset. The preprocessing of the captions is done by the following steps-

- Step-1: <start> and <end> tags are attached with each caption to mark the boundaries of the caption.
- Step-2: Any special symbol, if present, is removed
- Step-3: <start> and <end> and another tag <unk> (to denote unknown words) are added to the corpus.
- Step-4: The captions were tokenized and word-index and index-word dictionaries were created. Word-index maps the word to their corresponding index value in the feature vector and vice-versa for index-word.
- Step-5: The captions were padded with <pad> at the end to maintain equal lengths.

D. TRAINING OF NEURAL NETWORK

The training process is carried by passing the image though the encoder network and then associating the image with appropriate captions. During the forward pass, the text from decoder network are sampled and based on the generated text and actual text, the cost is calculated using the softmax cross-entropy loss function given in

equation (1). The delta rule used for weight update for softmax cross entropy during backpropagation is given in equation (2). Where y_i is the one-hot encoding and p_i is the softmax probability.

$$L = - \sum_i^n y_i \log(p_i) \dots \dots \dots (1)$$

$$\frac{\partial L}{\partial o_i} = p_i - y_i \dots \dots \dots (2)$$

III. ATTENTION MODEL

The previous section detailed an approach towards captioning an images but a question still remains, how exactly a RNN makes sense of the objects and their relationships. This is where the attention model is used. Attention Mechanisms in Neural Networks are (very) loosely based on the visual attention mechanism found in humans. Human visual attention is well-studied and while there exist different models, all of them essentially come down to being able to focus on a certain region of an image with high resolution while perceiving the surrounding image in low resolution, and then adjusting the focal point over time.

If the source sentence is 50 words long, the first word of the English translation is probably highly correlated with the first word of the source sentence. But that means decoder has to consider information from 50 steps ago, and that information needs to be somehow encoded in the vector. Recurrent Neural Networks are known to have problems dealing with such long-range dependencies. In theory, architectures like LSTMs should be able to deal with this, but in practice long-range dependencies are still problematic. For example, researchers have found that reversing the source sequence (feeding it backwards into the encoder) produces significantly better results because it shortens the path from the decoder to the relevant parts of the encoder. Similarly, feeding an input sequence twice also seems to help a network to better memorize things.

With an attention mechanism encoding the full source sentence into a fixed-length vector is not required. Rather, the decoder attends to different parts of the source sentence at each step of the output generation. Importantly, the model learn what to attend to based on the input sentence and what it has produced so far. So, in languages that are pretty well aligned (like English and German) the decoder would probably choose to attend to things sequentially. Attending to the first word when producing the first English word, and soon. Figure(1) shows the attention mechanism.

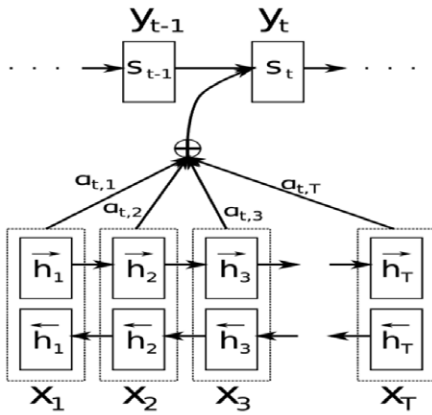


Fig.1: Illustration of Attention Mechanism

Here, the y_t s are our translated words produced by the decoder, and the x_i s are our source sentence words. The above illustration uses a bidirectional recurrent network. The important part is that each decoder output word y_t depends on a weighted combination of all the input states, not just the last state. The a s are weights that define in how much of each input state should be considered for each output. So, if $a_{3,2}$ is a large number, this would mean that the decoder pays a lot of attention to the second state in the source sentence while producing the third word of the target sentence. The a_t s are typically normalized to sum to 1 (so they are a distribution over the input states).

IV. BAHDANAU ATTENTION

For this work the Bahdanau Attention is used, which was used for Neural Machine Translation [9]. The attention model used is described in equation (3). The current hidden state s_i is calculated by an RNN f with the last hidden state s_{i-1} , last decoder output value y_{i-1} , and context vector c_i as given in equation (4). In the code, the RNN will be a nn. GRU layer, the hidden state s_i will be called hidden, the output y_i called output, and context c_i is called context. The context vector c_i is a weighted sum of all encoder outputs, where each weight a_{ij} is the amount of “attention” paid to the corresponding encoder output h_j as given in equation(5). Where each weight a_{ij} is a normalized attention “energy”. Each attention energy is calculated with some function a (such as another linear layer) using the last hidden state s_{i-1} and that particular encoder output h_j is given in equation (7). Each decoder output is conditioned on the previous outputs and some x , where x consists of the current hidden state (which takes into account previous outputs) and the attention “context”, which is calculated below. The function g is a fully-connected layer with a nonlinear activation, which takes as input the values y_{i-1} , s_i , and c_i concatenated.

$$p(y_i | \{y_1, \dots, y_{i-1}\}, x) = g(y_{i-1}, s_i, c_i) \dots \dots \dots (3)$$

$$s_i = f(s_{i-1}, y_{i-1}, c_i) \dots \dots \dots (4)$$

$$c_i = \sum_{j=1}^{T_x} a_{ij} h_j \dots \dots \dots (5)$$

$$a_{ij} = \frac{\exp(e_{ij})}{\sum_{k=1}^T \exp(e_{ik})} \dots \dots \dots (6)$$

$$e_{ij} = a(s_{i-1}, h_j) \dots \dots \dots (7)$$

V. TESTING THE MODEL

To test the model, the following steps are used:

- Step-1: Reset the hidden states of the network.
- Step-2: Pre-process the image and pass through the encoder network and get the attention weights.
- Step-3: Initiate the decoder network by <start> tag.
- Step-4: Stop sentence sampling once <end> tag is generated.
- Step-5: The sentence from <start> to <end> is the generated sentence.

V. RESULTS

On giving the image shown in (Figure.2) as input to the trained model, the predicted caption was “little boy suspended over a blue playground play equipment” against the actual caption “a boy playing on a playground”. The attention plot for the same is shown in (Figure.3) Similarly, for test image-2 shown in (Figure. 4), the attention plot is presented in (Figure.5). The training loss decreased with the number of epochs. The graph for is shown in (Figure.6). The trained model was able to predict captions and the average BLEU score for the predicted captions was 0.6526. Better results can be obtained and the training loss can further be decreased if trained on better hardware.



Fig. 2: Test image-1 given to the trained model

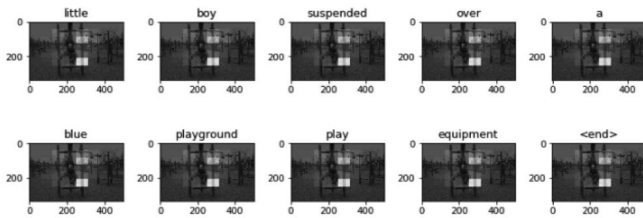


Fig. 3: Attention plot of the test image-l



Fig. 4: Original Caption: "two identical dogs bound across a lush green meadow". Predicted Caption: "two black dogs running in a grassy field"

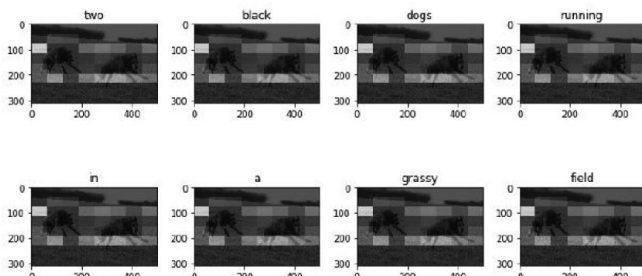


Fig. 5: Attention heatmap

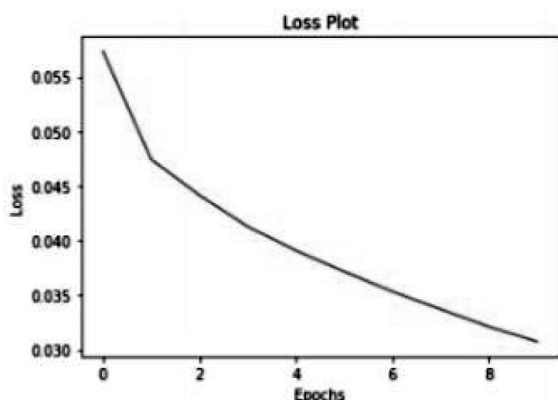


Fig. 7: Training Loss vs Number of Epochs

VI. CONCLUSIONS

In this work an automatic image caption process is discussed. Due to our hardware limitations, even training for 10 epochs took around 5 hours. The major problem

through the course of this work has been the setup of a suitable environment and even trying the first model. Also, the limitation of hardware available to us further constrained the amount of training and fine tuning that can be performed. Nevertheless, using pre-trained "imagenet" weights for the inception engine that was used for encoding the images, and by just training the model for 10 epochs, the model is able to yield results that are somewhat satisfactory.

ACKNOWLEDGEMENT

We express our sincere gratitude to our project guide Dr. Pradyumna Kumar Tripathy of CSE department for guidance and encouragement in carrying out this research work.

REFERENCES

- [1] L. Cun, "Gradient Based Learning Applied to Document Recognition", <http://yann.lecun.com/exdb/publis/pdf/lecun-01a.pdf>
- [2] L. Cun, Bengio, "Deep Learning", <https://www.cs.toronto.edu/~hinton/absps/NatureDeepReview.pdf>
- [3] Martn A., Paul B., Jianmin C., Zhifeng C., Andy D. "TensorFlow: a system for large-scale machine learning", <https://tensorflow.org>
- [4] Szegedy et al, "Going Deeper with Convolutions", <https://arxiv.org/pdf/1409.4842v1.pdf>
- [5] Szegedy et al, Rethinking the Inception Architecture for Computer Vision, <https://arxiv.org/pdf/1512.00567v3.pdf>
- [6] Min Lin et al, "Network in Network", <https://arxiv.org/pdf/1312.4400.pdf>
- [7] Arora et al, "Provable bounds for learning some deep representations", <https://arxiv.org/abs/1310.6343>
- [8] Xu et al, "Show, Attend and Tell: Neural Image Caption Generation with Visual Attention", <https://arxiv.org/pdf/1502.03044.pdf>
- [9] Bahdanau et al, "Neural Machine Translation by Jointly Learning to Align and Translate", <https://arxiv.org/pdf/1409.0473.pdf>
- [10] Tom M. Mitchell, "Machine Learning", McGraw Hill, ISBN 978-0-07-042807-2.

**Amit Kumar,
Debasish Patnaik,
Sudhanshu Ranjan**

7th Sem. CSE

Biosensors

The use of biosensors has got the utmost importance in the field of discovery of drug. 'Biosensors' refers to a powerful and innovative analytical device which involves the biological sensing element having wide range of applications like biomedicine, environment monitoring and many more.

Clark and Lyons invented the first biosensor in the year 1962 that measured glucose in the biological sample. The sensors utilized the electrochemical detection of oxygen or hydrogen peroxide. With time the performance of the biosensors has changed from classical electrochemical to visual/optical, glass, silica, nano-material and polymers and improved the sensitivity detection and also the selectivity. On the basis of label based and label free detection the biosensors work technically. In recent times the label free detection is being used for various detection. These sensors are widely used in the fields of environmental science and medicine. The first electromechanical sensor, "Glucometer" which was discovered in the year 1962 was made using glucose oxidase-based biosensors. These sensors play a vital role in the clinical diagnosis of diseases, as it helped in the monitoring and early diagnosis of the disease. The Scheme of steps involved in the preparation and functioning of a biosensors is shown in figure 1.

These biochemicals are also important aspect in Environmental monitoring. These were used for identifying the pesticidal residues rapidly in order to prevent health hazards. Silicon nanomaterials are used in biosensors due to abundance, biocompatibility and other properties. Also, silica or quartz or glass materials have many unique features hence they are used for making

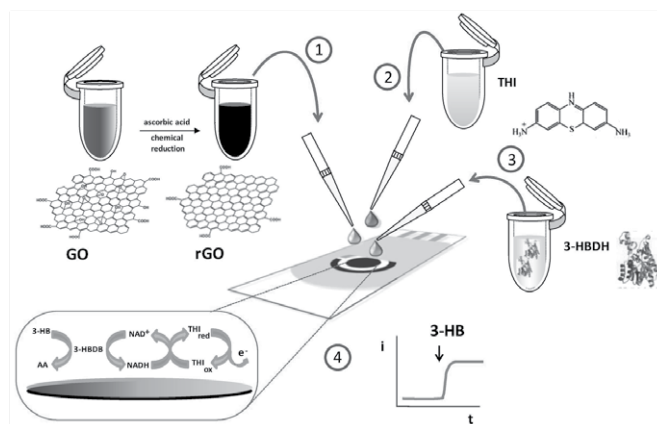
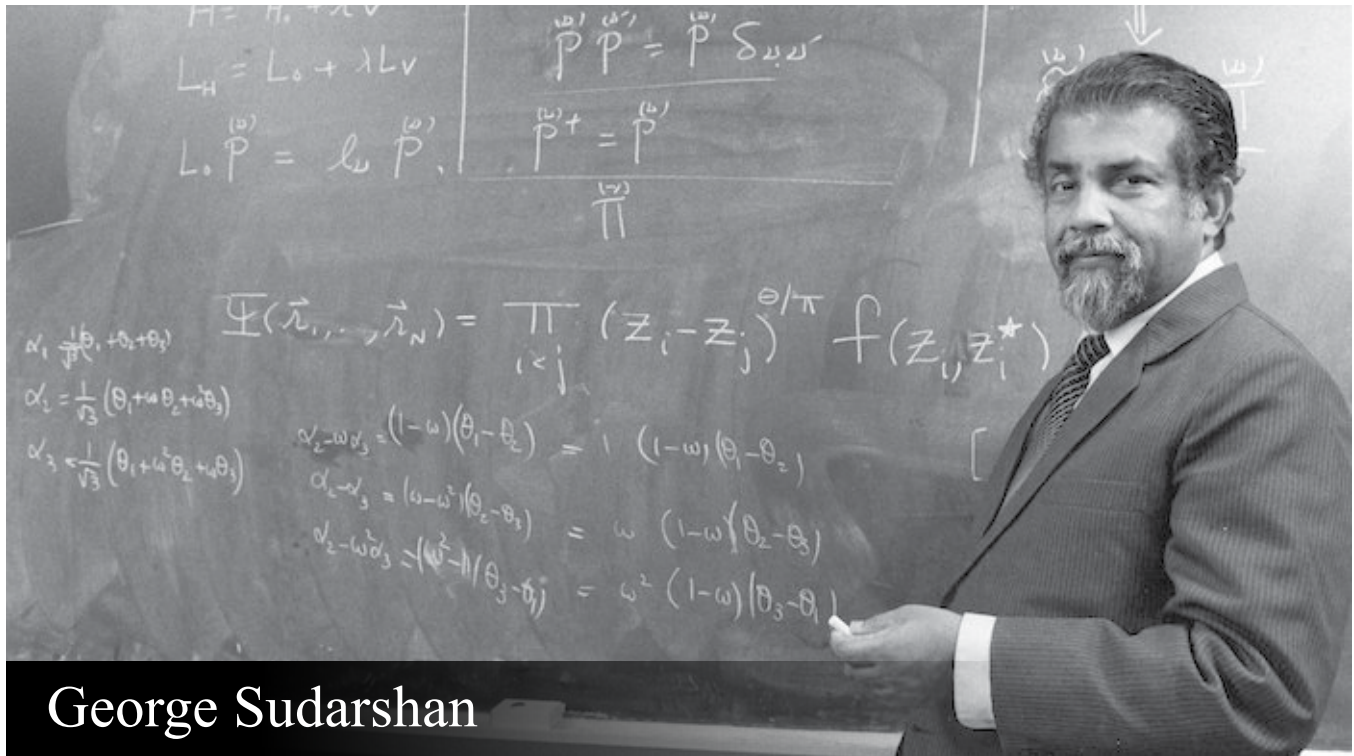


Figure 1: Scheme of steps involved in the preparation and functioning of a biosensors.

new biosensors for bioinstrumentation improvement. Also, nanomaterials like gold, silicon, silver and copper are used for making biosensors. The platinum-based nanoparticles are used for electrochemical amplification with a single level response for detection of low concentration of DNA.

Integrated strategies that use multiple technologies ranging from electromechanical, electrochemical and fluorescence-optical-based biosensors are modern methods in the field of biosensor discoveries. Some of these biosensors have tremendous application prospects in disease diagnosis and medicine. 2D and 3D detection are being worked upon in the recent times. The advancement in the development of microbial biosensors will contribute largely for environment monitoring and energy demand.

Sweta Patnaik
7th Sem., ECE



George Sudarshan

Indian physicist George Sudarshan was born on 16th September 1931, in Kerala. He was part of a Syrian Christian family. Ennackal Chandy George Sudarshan, popularly known as George Sudarshan spent his early life in Kottayam, where he went to CMS college, Kottayam for his early education. In the year 1951, he graduated from the Madras Christian College with honours. He secured a masters degree from the University of Madras, in the very following year. George got an opportunity to work along with Homi Bhabha, a popular nuclear physicist in the Tata Institute of Fundamental Research (TIFR). As a graduate student, he went to the University of Rochester, New York for further education. There he met Robert Marshak, who later on became the president of the American Physical Society. In the year 1958, he was awarded a Ph.D from the University of Rochester. After this, he went to Harvard University as a post-doctoral fellow alongside Julian Schwinger.

As a physicist, he has made a tremendous contribution in many areas of physics. He was the co-founder of V-A theory of the weak force. This theory later on set the path for the Electroweak theory. He was also involved in the development of quantum representation of coherent light. His most important work was in the field of quantum optics. His theory proves the equivalence of classical wave optics to quantum optics. One of his finest work is the suggestion of Tachyons, the essential

particles that travel faster than light. His Dynamical Maps was something that led to the aids of the theory of open quantum system. He also proposed the idea of the Quantum Zero Effect, with the support of Biddyanait Misra.

He has also taught at many different institutions such as, Tata Institute of Fundamental Research (TIFR), Syracuse University, University of Rochester and Harvard University. Since the year 1969, he was with the University of Texas (Austin) in the capacity of a professor of Physics and was also a Senior Professor at the Indian Institute of Science. He also worked as the Director of the Institute of Mathematical Sciences in Chennai, during the 1980's.

For his vast work in the field of physics, he has received many awards and honours including C.V Raman Award (1970), Padma Bhushan (1976), Bose Medal (1977), First Prize in Physics, Third World Academy of Science (1985), Majorana Prize (2006), Padma Vibhushan (2007), Dirac Medal of the ICTP (2010) and the Kerala Sastra Puraskaram, the state award for a lifetime accomplishment in Science (2013). At the age of 86, he passed away in the city of Austin, United States.

Abhilash Mishra

7th Sem. ECE



A New Airplane Uses Charged Molecules, not Propellers or Turbines, to Fly

A newly designed airplane prototype does away with noisy propellers and turbines. Instead, it's powered by ionic wind: charged molecules, or ions, flowing in one direction and pushing the plane in the other. That setup makes the aircraft nearly silent. Such stealth planes could be useful for monitoring environmental conditions or capturing aerial imagery without disturbing natural habitats below. Most planes rely on spinning parts to move forward. In some, an engine turns a propeller that pushes the plane forward. Or a turbine sucks in air with a spinning fan, and then shoots out jets of gas that propel the plane forward. Ionic wind is instead generated by a high-voltage electric field around a positively charged wire, called an emitter. The electricity, often supplied by batteries, makes electrons in the air collide with atoms and molecules, which then release other electrons. That creates a swarm of positively charged air molecules around the emitter, which are drawn to a negatively charged wire. The movement of molecules between the two wires, the ionic wind, can push a plane forward. The current design uses four sets of these wires. Moving ions have helped other things to fly through the air, such as tiny airborne robots. But conventional wisdom said that

using the approach to move something through the air as big as an airplane wasn't possible, because adding enough battery power to propel a plane this way would make it too heavy to stay aloft. (The ion thrusters that propel spacecraft through the vacuum of space work in a very different way and aren't functional in air.) Attempts to build ion-propelled aircraft in the 1960s weren't very successful.

MIT aeronautics researcher Steven Barrett thought differently. With the right aircraft design and light enough batteries, flight might be possible, his initial calculations suggested. So he and his team used mathematical equations to optimize various features of the airplane — its shape, materials, power supply — and to predict how each version would fly. Then the researchers built prototypes of promising designs and tested the planes at the MIT indoor track, launching them via a bungee system. "The models and the reality of construction don't always match up perfectly," Barrett says, so finding the right design took a lot of tries. But in the new study, he and his collaborators report success: 10 flights of the aircraft, which has a 5-meter wingspan and weighs just under 2.5 kilograms.

Source: Science News, Magazine of the Society for Science and the Public

Mathematical Techniques in Tomography

The importance of mathematics in the study of problems arising from the real world, and the increasing success with which it has been used to model situations ranging from the purely deterministic to stochastic is well defined. This consists of finding some unknown properties of an object or a medium observing the response of the object or the medium to a signal. One can find the points on the surface of an object in a non-invasive way. This could be done by analyzing the reflected/transmitted ultrasonic sound waves or x-rays. This is fraught with the probability of added noise which could introduce an error in the estimation of shape and size of an obstacle like gall stone or kidney stone.

In this thesis reconstruction of the image has been done from the scattered data disturbed by white noise. White noise can be effectively filtered as its expectation is

zero, however its variance is non-zero. This gives an idea how much absolute error could have been there in determination of the shape and size of the object. Image reconstruction has been done from the scattered data disturbed by a noise which arises out of a stable process. In this type noise converges in mean to the original shape with an additive noise again represented by a stochastic integral.

Experimentally shown that the shape of a cylindrical object looks almost like the original shape if the amplitude of the noise is small. But for large amplitude the distortion is wild. The original shape is unrecognizable. In sequel the reconstruction of the three dimensional compact object has been done which is disturbed by noise using spherical harmonics.

Manasa Dash

Dept. of BSH

Some interesting Apps on Android

IFTTT: It is easily one of the most useful apps ever. It's an app that creates commands to carry out a set of basic tasks automatically. What's great about the app is the sheer number of services, products, and other apps that have IFTTT supports. You can have it turn on your smart lights in your home, save images from Instagram and upload them to Dropbox, and there is even some Google Assistant and Amazon Alexa stuff available. It doesn't take very long to learn.

Google Translate: It is the go-to translation app available on any platform. It has received a number of updates over the years, including the ability to use your camera to point at something and have it translated in real-time. There is also a neural network powering the platform that helps make translation even more accurate. It has a slew of additional features as well, including the ability to translate a two way conversation in real-time. Travelers already know how useful this app is.

Simple Habit: Smartphones aren't known for their mental health benefits, but apps like Simple Habit make a compelling case. Simple Habit helps reduce stress in your life by presenting you with meditation guides

tailored for your life. Select your current location and how much time you have to spare, and let your mind relax from the hustle of everyday life.

Dashlane: The seemingly endless number of passwords needed to sign into your digital apps and services is dizzying, which makes password security a priority. Alongside the equally well-regarded 1Password, Dashlane password manager will help you generate secure passwords, then encrypts and stores them on your device so you don't lose track.

Relax melodies: Choose from more than 50 sounds to mix into tailored sleep audio tracks. It also includes guided meditations that specifically focus on helping you fall and stay asleep. When you find the right tunes that help you get some Zzz's, don't forget to share them with friends, as the app lets you tell the world about your insomnia cure.

Evernote is a mobile for note taking, organizing, task lists, and archiving. The users to create notes, which can be formatted text, web pages or web page excerpts, photographs, voice memos, or handwritten "ink" notes.

Source: Excerpted from readily available web sources



BIOPLASTS : Will they pave the way during phasing out of plastics?

Globally, as many as 160,000 plastic bags are used every second and currently, only 1-3 % of them are recycled. Even when disposed off properly, they take many years to decompose and breakdown, generating large amounts of garbage over long periods of time. If not disposed properly the bags can pollute waterways, clog sewers and have been found in oceans affecting the habitat of animals and marine creatures. Research shows the average operating “lifespan” of a plastic bag to be approximately 20 minutes. Plastic bags can last in landfill – an anaerobic environment – for up to 1000 years. In India, cows regularly die from plastic ingestion. There have been claims that chemicals in plastics can leach into food or drink and cause cancer. In particular, there have been rumours about chemicals called Bisphenol A (BPA) and Dioxins. Governments all over the world have taken action to ban the sale of lightweight bags, charge customers for lightweight bags and/or generate taxes from the stores who sell them. In many countries of the world, there has been a phase-out of lightweight plastic bags. Single-use shopping bag made up of plastic is commonly made from low density polyethylene (LDPE) plastic. Because of various ecological and environmental impacts of these plastics, scientists have been trying to find biodegradable solutions from natural resources such as starch based plastics, cellulose based plastics & protein based plastics etc and from petrochemicals as well. Biodegradable plastics take three to six months

to decompose fully. That’s much quicker than synthetic counterparts that take several hundred years.

Many companies and entrepreneurs have come up with concepts of biobased plastics bags also known as organic bags. A French company Sphere PTL based in Seine-Maritime in association with the German green chemistry laboratory Biotec have researched on potato starch’s potential to create organic bags. It actually has several advantages which could be useful in those fields: flexibility, non-toxicity, memory foam etc. In India, an entrepreneur named Kevin Kumala invented a bag using cassava, a widespread root in the country. It’s combustion will not release any GHG into the atmosphere. Watered down in hot water, it can be drunk by both animals and humans without being dangerous health-wise, for there isn’t any toxic substance in its composition. A start-up Envi Green owned by Ashwath Hedge, in Bangalore, have mixed twelve natural ingredients stem from fruits and vegetable rich in starch or in fibers (corn, tapioca, potatoes, bananas), as well as vegetable oil and organic waste by-products to make a biodegradable and edible bag. In nutshell, phasing out of plastics will need a befitting replacement for convenience on consumers which should be eco-friendly, cost effective and accessible. Bioplasts is a solution which should come to the purview of government, citizens and all commercial agencies dealing with plastics.

Publication Cell

Tel: 99372 89499 / 8260333609

Email: publication@silicon.ac.in

www.silicon.ac.in

The Science & Technology Magazine

Digital Digest

Contents

Editorial	2
DD Feature	3
Profile of a Scientist	20
PhD Synopsis	22
Environmental Awareness & Concerns	23

Editorial Team

Dr. Jaideep Talukdar
Pamela Chaudhury
Dr. Lopamudra Mitra

Members

Bhagyalaxmi Jena
Nalini Singh
Dr. Soumya Priyadarshini Panda
Dr. Priyanka Kar

Student Members

Abhilash Mishra
Jagannath Gouda

Media Services

G. Madhusudan

Circulation

Sujit Kumar Jena

Make your submissions to:
publication@silicon.ac.in

Silicon Institute of Technology

सिलिकन प्रौद्योगिकी संस्थान



Bhubaneswar Campus
An Autonomous Institute
Silicon Hills, Patia
Bhubaneswar - 751024



Sambalpur Campus
An Affiliated Institute
Silicon West, Sason
Sambalpur - 768200