

International Journal Research Publication Analysis

Page: 01-14

LEXICON BASED SENTIMENT ANALYSIS FOR INFORMATION RETRIEVAL ALGORITHM TO EXTRACT IMAGE

***¹Dr. S. Brindha, ²Sajith Balaji Sarvesh**

¹Faculty of Liberal Arts and Business Studies, SRM Institute of Science and Technology,
Kattankulathur 603 203.

²Student, Computer Science and Engineering, Karpaga Vinayaga College of Engineering and
Technology, Chengalpattu 603 203.

Article Received: 29 June 2026

Article Revised: 19 July 2026

Published on: 06 August 2026

***Corresponding Author: Dr. S. Brindha**

Faculty of Liberal Arts and Business Studies, SRM Institute of Science and Technology,
Kattankulathur 603 203.

DOI: <https://doi-doi.org/101555/ijrpa.2906>

ABSTRACT

Machine learning is a form of data analysis that employs automation to create analytical models. Machine learning, a kind of artificial intelligence, operates on the premise that robots can identify patterns in data, derive conclusions, and make decisions with minimal human intervention. The fundamental objective of pattern recognition, regardless of whether supervised or unsupervised, is classification. Among the several contexts in which pattern recognition has evolved, the statistical methodology has been the most thoroughly examined and applied in practice. Recent years have witnessed a heightened interest in neural network methodologies and strategies rooted in statistical learning theory. Key considerations in the design of a recognition system encompass the delineation of pattern classes, sensing environment, pattern representation, feature extraction and selection, cluster analysis, classifier design and learning, training and test sample selection, and performance evaluation. The overarching issue of identifying intricate patterns within random patterns of varying scale, orientation, and location persists unresolved. This research presents a text extraction pipeline designed to extract text from diverse high-quality photos sourced from social media. Following the classification of the input photographs, they are subjected to class-specific preprocessing, including text localization and illumination improvement. The structured representation developed in the prior stage facilitates the discovery of association rules, the identification of prevalent keywords, and the execution of sentiment analysis through a lexicon-based methodology that utilizes a collection of positive and negative terms. Every

tweet is assigned a score according to a scoring function.

KEYWORDS: Information Retrieval, Fuzzy logic Analysis for Text Retrieval Method, Pattern Retrieval, Pattern Mining, Lexicon based Sentiment analysis, Sentiment Image Analysis.

I. INTRODUCTION

In recent years, information retrieval (IR) has emerged as a more prominent field of research within computer science. Information retrieval systems are employed across diverse application fields, such as web search, digital library search, blog search, information filtering, recommender systems, social search, and others. The principal objective of information retrieval (IR) is to identify "relevant" documents or material pertaining to user requests as specified by a query from an extensive data corpus within a suitable time limit. Assessing the visual resemblance of images defined by feature descriptors is a method to ascertain visual content similarity. Social media constitutes a distinct category of information source. The extensive content generated by ordinary users can elucidate their wants, desires, and perspectives. Applications of this data encompass enhancing connection with consumers and residents, as well as utilizing it for personal advertising. Nevertheless, images constitute a substantial segment of the content on social networks. The previously described OCR engines are not designed to handle this type of input data. An ordinary social network image is taken with a smartphone and may exhibit subpar quality (inadequate lighting, distorted perspective), background text, etc. The accomplishment of the ILSVRC'12¹ problem demonstrates that classification challenges are the primary contributors to the success of deep learning.

Convolutional networks employing ImageNet are very efficient visual representations known as "Deep Features," and their application in content retrieval has recently produced encouraging outcomes. Information retrieval employs four sorts of models: fuzzy set models, probabilistic models, vector space models (VSM), and Boolean logic models [2]. Boolean logic models are the foundation of conventional information retrieval technology. Any question framed as a Boolean expression of terms can be constructed using the Boolean logic model, with documents transformed into sets of words [3]. When the query language comprises word groupings, the document can be obtained. The Boolean model requires an exact correspondence between the query language and the material, together with accurate input from the user.

II. RELATED WORKS

Text mining is a technology that facilitates the extraction of relevant and structured information from imperfectly structured data patterns. Computers encounter a complex difficulty in comprehending unstructured input and converting it into structured data. The variety of language approaches available enables individuals to fulfill this obligation effortlessly. In comparison to computers, humans exhibit a slower processing speed and possess a limited capacity for information storage.

In executing these tasks, computers have considerably greater proficiency than humans. Consequently, upon comparing data mining and text mining, we may ascertain that text mining is more vital. The predominant portion of data now accessible within any organization is formatted as text. Conversely, as text mining is utilized to organize unstructured textual material, this activity is more challenging than data mining. This technique is known as Web Information Retrieval, or Web IR, and involves the application of information retrieval theories and methodologies to the World Wide Web. Within the realm of the World Wide Web, it largely addresses the technical challenges faced by Information Retrieval (IR).

Schouteten et al. [14] devised a face recognition system utilizing a one-class support vector machine, specifically designed for mobile devices operating on the Symbian platform. An EER of 7.92% and 3.95% can be attained, respectively, based on the global threshold and the individual threshold, as per the assessment of the Nokia 6680 mobile for evolution. Optical Character Recognition (OCR) entails converting photographs that include text into a machine-readable format. This procedure use algorithms and software to analyze photos, recognize characters, and transform them into editable text. The fundamental functionality depends on OCR technology, which can be integrated into other software applications and internet converters.

Sheetrit et al. [11] expressed interest in employing pattern recognition techniques for anatomical neuroimaging data. Nevertheless, relatively less research has been conducted on the extraction of visual information to optimize forecasting. This work employed a Gaussian process machine learning methodology to estimate the body mass index, age, and gender of the diverse elements within the dataset.

Shinde, R et al. [12] proposed shot face unlocks, a prevalent technique for unlocking mobile devices that utilizes data from a 180-degree panoramic shot of the device surrounding the user's head. The facial recognition technique was formulated utilizing multiple support vector machines and neural networks, and the outcomes validate the efficacy of the technology.

Shruti et al. [10] elucidate that contemporary agriculture necessitates prompt identification of

foliar diseases to prevent losses and optimize production. Conventional diagnostic procedures reliant on human examination are dependable yet time-intensive and prone to human mistake, hence necessitating automated systems utilizing image processing techniques. The current review paper discusses advanced methods for the automatic identification of leaf diseases, employing complex image-processing and machine-learning algorithms utilizing CNNs and BNNs. This methodology employs acquisition, preprocessing, segmentation, and feature extraction of images based on color, texture, and form attributes.

This research presents a sophisticated automated hydroponic farming system that incorporates image processing and meticulous fertilizer control to enhance plant growth. In hydroponics, plants develop without soil, depending on a nutrient-dense aqueous solution. The system's foundation comprises real-time picture processing to evaluate plant health and identify nutrient deficiencies. Computer vision image processing techniques are employed to evaluate plant photos, identifying critical health indicators such as leaf color, size, shape, and texture. According to these evaluations, the system can dispatch alert messages to the user and autonomously modify the nutrient content in the solution to sustain optimal plant health. Furthermore, it oversees and autonomously modifies other vital characteristics such as pH level, temperature, humidity, and light intensity, which are critical for nutrient absorption and general development.

Chen et al. (2025) elucidates that the novel development of vector graphics with high-resolution photos through Artificial Intelligence has emerged as a significant endeavor in edge extraction. This work uses Qiang embroidered images as a case study because they include intricate edges, making them particularly ideal for image processing and pattern recognition research. Initially, we employ suitable pre-processing techniques, specifically improved adaptive median filtering (IAMF), to mitigate image noise. The Xception architecture, founded on convolutional neural networks, is employed for edge detection and extraction.

III. PROPOSED WORK

Incoming photos are categorized into distinct classes, and suitable preprocessing is conducted for each group. The resultant images serve as input for the OCR engine. Additionally, certain post-processing techniques may be implemented on the output of the OCR engine. A dataset of photographs from a social network has been compiled, and experiments have been conducted using this dataset.

3.1 Pre-processing

With the goal of returning only the most relevant pages, "sparse search methods" attempt to reduce the dimensionality of the indexed content. These methods work on the assumption that getting all of the documents from a collection would be pointless as only a small subset of the data is pertinent to a given query. **Data Collection from Text:** An analysis of the incoming social network image stream was conducted, and a representative group of images primarily in Russian containing potential social propaganda and information leakage were selected. The 67 photos that make up the input dataset include 64554 symbols. Four distinct groups have been established for all of the photographs.

Demotivators. Social networks frequently employ photographs of this type. Typically, each photograph has a main image (scene, people, etc.) and some text underneath it. Twenty-four percent of photographs are in this category.

Photographs of documents, certificates, identity cards, etc. are typically included in this category. Some sentences of text may be included in images in this category. There are 38% pictures in this category. Examined the Images in this category are high-quality scans of papers that include a lot of text. Twenty-four percent of photographs are in this category.

Smartphone. The photographs in this category are the same as those from the scanned category, but they were taken using a low-quality smartphone camera. There are 14% of photos in this category. Mostly 85% of the portrait galaxy essential remain engaged by the text document. Images from certificates and demotivators may include a background, a scene, and typically a few phrases of text. On the other hand, photos from scanned categories and smartphones typically contain a lot of text and are a picture of a text document.

3.2 Feature extraction

Opinion mining is increasingly recognized as a burgeoning area of research. The objective is to align specified textual sequences by a technique known as pattern matching. Preparation for demotivators. Demotivator class graphics generally use a dark background with light text positioned below the primary image. Dark text may subsequently manifest on the primary image. Our experiments, however, demonstrated that Tesseract exhibits markedly superior performance when the text is dark against a light background. Moreover, certain demotivators may feature text with markedly inconsistent font sizes, thereby diminishing identification accuracy.

Consequently, we employ three iterations to process images of this category. Upon completion of the requisite preprocessing, the image is supplied to Tesseract, which receives the text and corresponding bounding boxes at each stage. Text segments recognized at each

phase are subsequently amalgamated.

All extensive Hough lines of the image are encompassed inside the minimal rectangle that delineates the primary image. During the second iteration, the complete image undergoes bitwise inversion, and the background color is superimposed onto the primary image. All bounding boxes from previous iterations are obscured with background color and input into the OCR engine during the third iteration. Preparation of certificates. Images in this category may exhibit low resolution, and a small segment of the image may contain text. Furthermore, the presence of background elements in the primary document input image diminishes recognition accuracy. Consequently, we utilize the preprocessing outlined below.

Improvement of image resolution: It is applied to images that are below a specified threshold size. The principal document is subsequently extracted from the input image through the document localization procedure. The background color is selected from the identified bounding boxes, while the rest of the image is obscured. The preprocessing has been scanned. Text recognition algorithms perform effectively with photographs of this category. Resolution enhancement is implemented solely when the image's area falls below the threshold. In the initial phase of preprocessing for smartphones, lighting improvement is applied to the input image if deemed essential. When the image's area falls below the threshold, an enhancement in image resolution is implemented. The subsequent phase entails identifying perspective distortions and, upon their detection, implementing the appropriate perspective correction to the image.

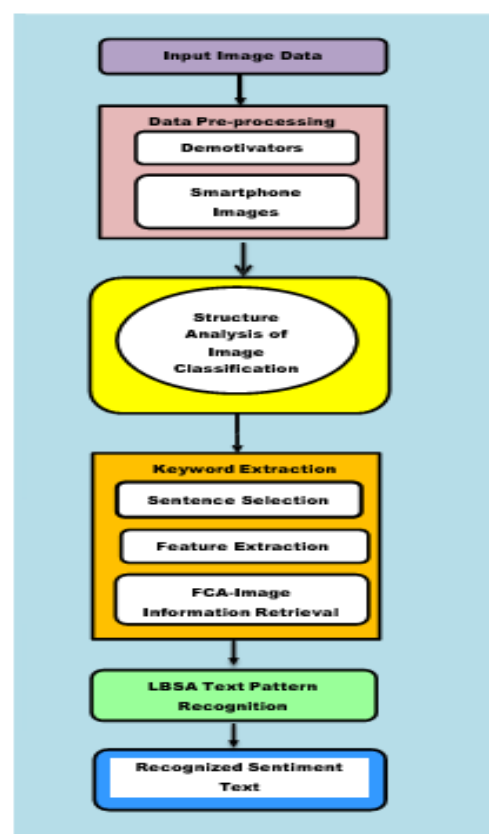


Fig. 3.1 Flow Diagram of LBSA-IR.

The suggested pipeline facilitates the extraction of textual information from social media, when a substantial volume of data is conveyed through low-quality photos or intermixed with non-textual content. The OCR engine utilizes the input images to generate the output text. If the input image corresponds to the demotivator class, supplementary rounds of the second and third stages are implemented. Words exhibiting low confidence are excluded from the final text if the overall confidence of the document falls below a specified level (an extra parameter). Predefined classes are utilized to classify input photographs. Subsequently, the vibrant image is converted to grayscale and undergoes class-specific preprocessing. Additional preprocessing and recognition iterations are required for the "demotivator" category. The final text file is generated following the post-processing of the Tesseract engine's output to eliminate certain extraneous symbols.

3.3 Fuzzy Concept Analysis for Information Retrieval Method (FCA-IR)

By concentrating on the allocation of specific classes instead of categories throughout the document's trajectory, text classification enhances its organization in a coherent and user-friendly format. In the supervised method, the model is developed using the training set. The categories are predefined, and the documents in the preparation dataset are manually

annotated with one or more category labels. Upon training on the earnings dataset, the classifier can subsequently predict the category of a new document. The classifier's strength, contingent upon the employed classification algorithm or approach, yields a confidence estimate reflecting the certainty of the organization's label's accuracy.

The impact of the stipulations in the document. Any statement possessing semantic significance is classified as a term, and the document's meaning is augmented by the aggregation of these words. Terminological methods address the contemporary issues of polysemy and synonymy. A term with several meanings is termed polysemic; a term with various meanings is referred to as synonymous. The semantic interpretations of numerous newly identified terms are vague and do not meet user expectations. The two principal categories of text mining techniques are phrase-based and term-based methods.

Stemming

Stemming is the procedure of eliminating derivational and inflectional affixes from words to revert them to their base form. Stemming is frequently employed in occupations related to information retrieval.

Step 1:	<i>Input image</i>
Step 2:	<i>If ($W_n == 0$) Return final Sentiment (F_p, eP)</i>
Step 3:	<i>Else if $*W_p == 0$ Return final Sentiment (F_n, E_n)</i>
Step 4:	<i>Else { If ($F_p - F_n > 0.1$) Return final sentiment (F_r, E_r)}</i>
Step 5:	<i>If ($F_p + F_n > 0$) Return final sentiment (F_r, e_r)</i>
Step 6:	<i>converted plain image to text considered to be a meaningful English sentence but it captures the original sentiment and context.</i>
Step 7:	<i>Sentiment polarity classification system M is a quadruple $M = \{\alpha, \lambda, \delta, \phi\}$</i>
Step 8:	<i>Function to calculate the similarity score of two messages.</i>
Step 9:	<i>Using accuracy, determine the k nearest neighbours that are the heuristically optimal number.</i>
Step 10:	<i>Assess the reaction to text messaging on social media and other web channels.</i>

Numerous research indicate that stemming enhances the efficacy of information retrieval systems. The Porter stemmer is the preeminent algorithm for English stemming. Nevertheless, this stemming technique possesses numerous limitations as English morphology cannot be adequately elucidated by its basic principles. Stemming is the process of reducing an inflected term to its root form. The names "programming," "programmer," and "programs" can all be abbreviated to the shared root "program." In other terms, "program" is synonymous with the three preceding inflection words. This text-processing method is frequently advantageous for mitigating sparsity and/or standardizing terminology. The similarity in meaning between inflected words and their stems minimizes redundancy and enables natural language processing (NLP) models to discern linkages between them, hence enhancing the model's comprehension of their usage in related settings.

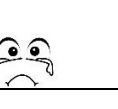
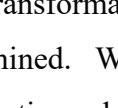
Stemming algorithms truncate the beginning or end of a word utilizing a compilation of common prefixes and suffixes present in inflected terms. We can ascertain that while this approach offers advantages, it also presents drawbacks, as it sometimes produces word stems that are not legitimate words. This technique yields optimal model parameters, including network design, learning rate, and regularization, which enhance the accuracy of blood group prediction. The output results indicate that the integration of AI, image processing, and FA significantly enhances both the accuracy and efficiency of Bgroup categorization.

3.5 Stop Words

Stop words are integral components of the lexicon commonly utilized across all languages, not exclusively English. Everyday vocabulary in any language, not alone English, inherently encompasses. Terms frequently employed across various languages, not exclusively English, are referred to as stop phrases.

The stop phrase incentive is effective for several items as it enables the emphasis on distinctive terms while eliminating relatively common ones in a comprehensible manner. Many individuals perceive stop words as a singular collection of phrases. It can undoubtedly present a diverse array of subjects in multiple circumstances. Examples of coordinating conjunctions that unite words, phrases, and clauses include for, and, nor, yet, or, but, and so.

The temporal significances of prepositions are disparate.

Social Media Conversation	Sentiment	Images
<i>We appreciate and respect your sacrifices</i>	Positive	
<i>If you carry your childhood with you, you never become older</i>	Positive	
<i>Joyful, Happy, Excellent, Brilliant, Good Food, Good Day.</i>	Positive	
<i>Bore Sunday, Lazy morning, Head ache, Fever, Not well.</i>	Negative	
<i>Unhealthy, Sick, not well</i>	Negative	

Every cluster of objects undergoes pattern mining, which is dependent on the transformation process. Two potential transformations weighted and Boolean are examined. When generating a set of transaction databases $\{D_1, D_2, \dots, D_k\}$ in both transformations, k are clustered so that the i -th transaction database D_i ($1 \leq i < k$) corresponds to the i -th cluster G_i . A collection of transactions $\{D_{i1}, D_{i2}, \dots, D_{i|G_i|}\}$ is contained in a transactional database D_i , where the j -th transaction D_{ij} denotes the j -th item ij of the cluster G_i .

3.6 Pattern Matching using TF-IDF

Text mining and statistical analysis frequently employ TF-IDF as a weighting mechanism. The expense of tf-idf escalates in direct correlation to the frequency of a phrase's occurrences inside the collection, excluding counter-occurrences resulting from the phrase's prevalence in the corpus. This may facilitate the organization of data, resulting in specific terms appearing more frequently in files. TF-IDF is an effective pre-word filtering technique for text categorization and summarization across several on social media, where they can enhance or condense textual meaning, researchers have examined their semantic correlation with words.

IV. RESULTS AND DISCUSSIONS

Data mining includes several technologies for pre-processing, analysis, and interpretation. These strategies primarily categorize into two types: pattern recognition and machine learning. The objective of pattern recognition is to identify discernible evidence of verified entities and relationships, or to detect patterns within the incoming data. These processes are

mostly linked to image analysis, although it is not the principal application type. Most machine learning techniques focus on deriving generalized insights from data, including images, which will then be utilized for predictive challenges. Relevant papers are those that assist the user in identifying the solution to an inquiry.

$$\text{Character accuracy} = \frac{\text{Number of Characters errors}}{\text{Number of Characters}} \quad \text{--- (1)}$$

As in our dataset classes are not balanced two kinds of accuracy are used. In the Table 4.1 Accuracy represents usual character accuracy and Accuracy represents equally weighted class accuracy.

$$\text{Precision} = \frac{RL}{\frac{\text{Retrieved Documents}}{\text{Relevant Document Retrieved}}} \quad \text{--- (2)}$$

where "errors" is a minimum number of edit operations (character insertions, deletions, and substitutions) needed to fully correct the text and "Number_of_characters" is a number of characters in document.

Recall is the second metric. It is the percentage of papers that have been located and are relevant to the query. When a binary classification test correctly detects or excludes a condition, accuracy is used as a statistical measure problem domains. The TF-IDF object utilizes two primary metrics: term frequency and inverse document frequency. Due to the rapid proliferation of emoticons.

$$\text{Accuracy} = \frac{TP+TN+FP+FN}{TP+TN} \quad \text{----- (3)}$$

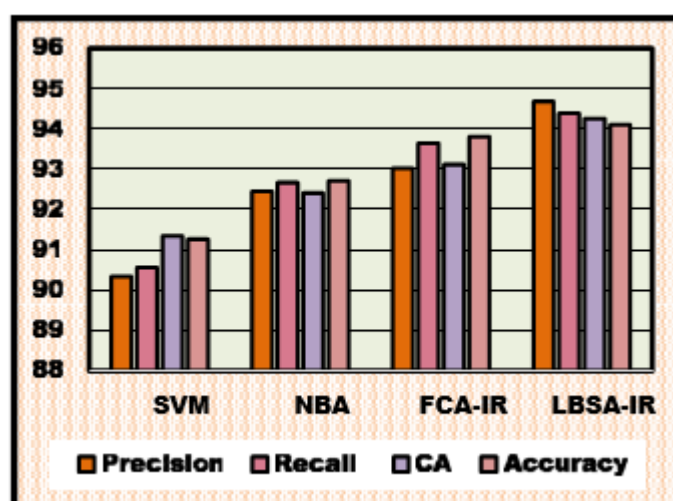
Table 4.1 gives the performance evaluation of LBSA-IR and NBA are analysed by means of existing method like K-means, Support Vector Machine, Naïve Bayes, and proposed methods are compared with proposed automated term based Fuzzy logic analysis for information retrieval (FCA-IR) and lexicon based sentiment analysis for identifying the image into suitable text format.

Table 4.1 Performance Evaluation of LBSA-IR.

Methods	Precision	Recall	CA	Accuracy
SVM	90.34	90.56	91.34	91.26
NBA	92.45	92.67	92.41	92.71
FCA-IR	93.02	93.65	93.12	93.81
LBSA-IR	94.69	94.39	94.26	94.11

Key considerations in the design of a recognition system encompass the delineation of pattern classes, sensing environment, pattern representation, feature extraction and selection, cluster analysis, classifier design and learning, training and test sample selection, and performance evaluation. The fundamental issue of identifying intricate patterns amidst random patterns of varying scale, orientation, and location persists unresolved. Following the classification of the input photographs, they undergo preprocessing tailored to specific classes. The structured representation established in the prior stage facilitates the discovery of association rules, the identification of prevalent keywords, and the execution of sentiment analysis through a lexicon-based methodology that utilizes a collection of positive and negative terms.

From figure 4.2 gives the performance evaluation of LBSA-IR and NBA are analysed by means of existing method like K-means, Support Vector Machine, Naïve Bayes, and proposed methods are compared with proposed automated term based Fuzzy logic analysis for information retrieval (FCA-IR) and proposed method lexicon based sentiment analysis for identifying the image into suitable text format. The proposed method LBSA-IR is helps to converted plain text may not be a meaningful English sentence but it captures the original sentiment and context.

**Fig.4.2 Comparison chart for performance evaluation of LBSA-IR.**

This method employs statistical analysis and cross-validation to ascertain the system's accuracy, utilizing a formula to compute the values. Identifying evolving, poorly delineated research domains across disciplines in academic literature presents a significant data collection challenge. Any information that fails to provide practical knowledge is considered irrelevant. Each object may be either attainable or inaccessible, contingent upon the user's inquiry. This section delineates the results of text extraction experiments performed on the aforementioned dataset. We employ the character accuracy metric for assessment.

V. CONCLUSION

This research presents a text extraction pipeline designed to extract text from various high-quality photos obtained from social media. Every tweet is assigned a score according to a scoring function. This paper presents the LBSA-IR recommended method for gathering relevant data without misinterpreting textual documents. Image recognition identifies the textual information associated with the corresponding pattern and extracts the individual document. To categorize the attitudes of electronic word-of-mouth communications according on their polarity. The principal obstacle in employing NLP technologies to comprehend social media interactions is mitigated by the two-step process, comprising Feature Extraction and Machine Learning. A centralized sentiment analysis system enables firms to assess consumer feedback and grievances with greater precision, thereby acquiring deeper insights.

REFERENCES

1. Bhamare DP, Suryawanshi P. Review on Reliable Pattern Recognition with Machine Learning Techniques. *Fuzzy Information and Engineering*, 2018; 10(3):362–377.
<https://doi.org/10.1080/16168658.2019.161103>.
2. Dr.S.Brindha, Dr.S.Sukumaran, Enhanced Pattern Classification Methods for Text Categorization Mining, *International Journal of All Research Education and Scientific Methods (IJARESM)*, ISSN: 2455-6211 Volume 10, Issue 10, October-2022, Impact Factor: 7.429, Available online at: www.ijaresm.com.
3. Chen, A., Chen, H., Peng, Y., Hu, J., Deng, L., Wen, X., & Wang, X. (2025). Generate vector graphics of fine-grained pattern based on the Xception edge detection. 36.
<https://doi.org/10.1117/12.3055748>
4. Liu, J., Guo, Y., kaihong, C., Zhuo, S., Li, Y., & Li, Z. (2025). *Research on the feature model of kumquat oil gland based on OCT technology image processing algorithm*. 45.
<https://doi.org/10.1117/12.3058049>

5. Iriani A, Hendry, Manongga DHF, Chen R-C. Mining public opinion on radicalism in social media via sentiment analysis. *International Journal of Innovative Computing, Information and Control*. 2020;16(5):1787-1800
6. Lee H and Kang P. Identifying core topics in technology and innovation management studies: a topic model approach. *J Technol Transf* 2018; 43: Pp:1291–1317, Published Year 2018.
7. Langlois, P., and Titah, R.(2020) "Utilisation et impact des outils d'intelligence artificielle dans des contextes de cyberjustice." Doctoral dissertation, HEC Montréal, Published Year 2020.
8. Swapna, N., Venkatesh, R., Manoj Kumar, A., & Anjali, E. (2025). Automatic Identification of License Plate Numbers for Non-Helmet Riders Using Image Processing. *Social Science Research Network*. <https://doi.org/10.2139/ssrn.5088906>.
9. Shen, X.; Dai, Q.; Mao, S.; Chung, F.-L.; Choi, K.-S. Network together: Node classification via cross-network deep network embedding. *IEEE Trans. Neural Netw. Learn. Syst.* 2020, 32, 1935–1948. [Google Scholar] [CrossRef] [PubMed]
10. Shruthi Thakur, S., Khobragade, P., Saraf, P., Turankar, A., & Lokulwar, P. (2025). Review Paper on Plant Leaf Disease Detection Software using Image Processing. <https://doi.org/10.2139/ssrn.5017648>.
11. Sheerit E, Kurland O (2019) Cluster-based focused retrieval. In: *Proceedings of the 28th ACM International Conference on Information and Knowledge Management*, pp 2305–2308. Sheerit E, Shtok A, Kurland O (2020) A passage-based approach to learning to rank documents. *Information Retrieval Journal*, 1–28.
12. Yu, L.; Li, Y.; Weng, S.; Tian, H.; Liu, J. Adaptive multi-teacher softened relational knowledge distillation framework for payload mismatch in image steganalysis.
13. *J. Vis. Commun Image Represent.* 2023, 95, 103900. [Google Scholar] [CrossRef]
14. Zhao, B.; Cui, Q.; Song, R.; Qiu, Y.; Liang, J. Decoupled knowledge distillation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, New Orleans, LA, USA, 18–24 June 2022; pp. 11953–11962. [Google Scholar]
15. Schouteten, J. J., Llobell, F., Roigard, C. M., Jin, D., and Jaeger, S. R. (2022). Toward a valence arousal circumplex-inspired emotion questionnaire (CEQ) based on emoji and comparison with the word-pair variant. *Food Qual. Prefer.* 99:104541. doi: 10.1016/j.foodqual.2022.104541.
16. Shinde, R. (2025). Automatic Hydroponic Farming System with Image Processing based on Nutrients System. *Indian Scientific Journal of Research In Engineering And Management*, 09(01), 1–9. <https://doi.org/10.55041/ijrsrem40805>.