

INNOVATIONS IN NATURAL LANGUAGE PROCESSING FOR SMART ENTERPRISE SOLUTIONS

**^{1*}Yangjingyu Li, ²Nandar Win ³Dr. Bhairab Sarma, ⁴Shantanu Bhadra,
⁵Dr. Payel Sengupta, ⁶Dr. Jayanta Aich, ⁷Dr. Susmita Biswas**

^{1*}Lecturer, Biocomputing and Engineering, Lincoln University College, Malaysia
ORCID ID: 0009-0005-2761-690X, Email ID: lisagz@lincoln.edu.my

²Lecturer, Academic & Management, PhD in Management, Vivekananda American University, Myanmar, Email ID: nandar.phdscholar@lincoln.edu.my ORCID ID: 0009-0001-8582-6172

³Associate Professor, The Assam Royal Global University, Assam, India Email ID: sarmabhairab@gmail.com

⁴Assistant Professor, Department of Computational Science, Brainware University Email ID: shantanuihrm@gmail.com

⁵Assistant Professor, Brainware University Email ID: Payel9433@gmail.com

⁶ Associate Professor, HOD, Department of Computational Science, Brainware University, Email ID: jaich@brainwareuniversity.ac.in

⁷Associate Professor, Department of Cyber Science & Technology, Brainware University, Barasat, West Bengal, India Email ID: bi.susmita@gmail.com

Abstract

Natural Language Processing (NLP) is fast revolutionizing enterprise activities through automated processing, classification, and management of text-based information at large scale. The proposed study is concerned with the issue of multi-class text classification in the context of enterprise news by enjoying the opportunities of transformer-based modeling to obtain high precision and reliability. The methodology included the utilization of a pre-processed and stratified Kaggle news dataset comprising 10,000 instances and 22 classes, as well as the DistilBERT architecture that is effective and supports contextual learning. The experiment process involved the use of advanced tokenization, optimized hyperparameters and stringent evaluation measures such as accuracy, macro and weighted F1-scores, per-class examination and visualization of the embedding through t-SNE and PCA. The findings showed that the DistilBERT classifier performed with an overall accuracy of 80.2%, and macro F1-score of 80.0%, and that it was strong even in infrequent and overlapping categories, which were well above the traditional and earlier deep learning baselines. The model has a nuanced understanding that was confirmed by error analysis and visualization, but there were still slight difficulties with the separation of semantically close classes. Finally, this work demonstrates that transformer-based NLP models can be effective in enterprise-scale text classification and points to the potential to build on innovation by adapting to a domain, providing interpretability, and scalable deployment, heralding smarter, context-aware enterprise solutions in the emerging digital environment.

Keywords: Natural Language Processing, Transformer Models, Multi-Class Text Classification, Enterprise Automation, DistilBERT

Introduction

Natural Language Processing (NLP) is at the cutting edge of technology, and it is transforming the ability of businesses to engage with data, customers and internal operations in a revolutionary way. Driven by major advances on deep learning, transfer learning, and the creation of large language models, NLP is empowering organizations to automate their most complex tasks, extract actionable insights in unstructured data of massive sizes, and create intelligent systems with human-level linguistic capabilities. The explosion of digital communication lines, including emails and support tickets, business documents and social media, has increased the pressure on effective language technologies in the enterprise market to be robust, scalable, and context-aware [1], [2], [3].

Over the past few years, the advent of architectures like transformers has drastically enhanced the accuracy and efficiency of the NLP applications. Such models as BERT, where bidirectional representation of text is exploited, and attention-based networks have achieved new standards in natural language understanding and generation [1], [3]. Such groundbreaking progress has propelled the creation of large pre-trained models that can learn few-shot and zero-shot learning, enabling companies to roll out high-performing systems using little task-specific data [4]. Sequence-to-sequence models, with BART and other transformer-based models as their prime example, also provide such enterprises with the means to streamline processes like document summarization, translation, and information extraction faster and more accurately than ever before [5].

There are, however, numerous challenges of incorporating NLP into smart enterprise solutions. The off-the-shelf language models could be not as correct when used on the enterprise data that is unstructured and heterogeneous, as well as often contains domain-specific jargon, abbreviations and inconsistencies [6], [7], [8]. Businesses also have to deal with changing data privacy laws and the need to ensure that sensitive data is safeguarded, particularly in cases where companies use cloud-based NLP services. Technical complexity of certifiable privacy in NLP pipelines has been shown to be not only the technical aspect of certification, but also an increasingly important need that extends beyond legal requirements [9]. Moreover, with increased complexity, deep learning models have been deemed as black boxes and discussions on interpretability, accountability and bias are becoming a reality especially in areas such as healthcare, law and finance where the transparency of decisions is paramount [10], [11]. Explainable AI techniques are hence critical to promote trust, regulatory compliance, and adoption by the stakeholders within the enterprise environment.

The other consistent challenge is associated with the scalability and flexibility of NLP solutions to different business activities and organizational levels. Although bigger companies may possess the resources to annotate a large amount of data or tune complex models, small and medium enterprises (SMEs) can find themselves limited by financial and technological resources to adopt the most state-of-the-art NLP [12]. This disparity in the ability to reach digital transformation may lead to disparity in the scope of innovation and competitiveness in the wider business community. The significance of domain adaptation and hybrid systems, that is, combining the rule-based and machine learning models, in achieving the bridge between generic language understanding and the idiosyncrasies of the most particular business processes, is also becoming increasingly apparent [13], [14].

This study is limited to recent developments in the field of NLP that are paving the way to the new generation of smart enterprise solutions. The analysis includes the developments in the pre-trained language models, transfer learning techniques, explainable AI, and business process automation tools based on NLP. Although the consideration covers both the technical and practical perspective, the review mostly focuses on text-based NLP, and multimodal solutions (related to audio or vision) are mentioned only when it is contextually applicable. The work views enterprise adoption as technology-based because the adoption is influenced by organizational, cultural, or economic reasons, which are not in the direct scope of the study. Moreover, the reviewed literature will be restricted to significant breakthroughs and conducted studies published no earlier than five years ago to guarantee that synthesis will display the state of the art [15].

What makes this research important is that it could help explain how NLP innovations are transforming the way enterprises work. The rise of business digitalization, and industry 4.0 concepts have all made NLP a key component of the automation of knowledge work, operational efficiency and customer experience enhancement [12], [16]. Processing of summarized information, smart chatbots, virtual assistants, and semantic search engines are becoming a constituent part of the enterprise ecosystem, which makes information flow and decision-making a seamless process. The potential of faster extraction of entities and relationships of business processes in an unstructured text with in-context learning and pre-trained models is speeding up automation of processes that were previously manual and error prone [17]. In addition, NLP combined with sentiment analysis, opinion

mining, and contract review is delivering real-time analytics with actionable insights to organizations that are competitive advantage drivers [16], [18].

In spite of these achievements, there still are considerable impediments. Businesses have to overcome the issues connected with the field specific customization, data confidentiality, explanation of models, and the ever-changing nature of language in business. Especially SMEs can only engage in digital innovation by having easy to use and affordable NLP frameworks [12], [19]. The need to have more transparent and accountable AI models is also driving the explainability and fairness research, so that NLP technologies can be used beneficially to all stakeholders in a fair and responsible way [10], [11]. The capacity to readily adapt to business process management tasks, automate some of the parts of speech tagging and offer fine grained entity extraction processes, as NLP continues to evolve will prove to be essential to the future of enterprise intelligence [20].

In this light, the research will be framed by the following research objectives:

- To search and examine the latest advances of natural language processing technologies, which can be used to create smart solutions in the enterprise.
- To review the major pitfalls and best strategies related to the implementation of NLP systems in enterprise settings and concentrate on data diversity, privacy, explainability, and flexibility.
- To define new trends, unanswered research questions, and possible directions in the way of NLP and AI integration in business process automation and digital transformation.

With these objectives, the paper will give technology leaders, enterprise architects, and AI practitioners a clear idea on how NLP is the next smart enterprise innovation. Finally, the synthesis also aims at filling the intellectual canyon between the fast-paced scholarly work and the sensible, responsible application of NLP to the environment of various businesses today [21].

2. Methodology

The approach to the research is to combine state-of-the-art transformer-based natural language processing (NLP) models and maintain the best pre-processing, training, and evaluation practices to build an enterprise-ready multi-class text classification system. It consisted of a broad goal to categorize news reports into a set of 22 fixed categories with high accuracy and reliability, which represented the real-world enterprise requirements in terms of content recommendation, market intelligence and knowledge management. The pipeline includes the following six major steps: data acquisition, preprocessing, selection of the model, training set up, evaluation metrics, and interpretability metrics.

2.1 Data Acquisition and Sampling

The dataset we have used in this paper is the publicly accessible dataset, News Articles Classification Dataset for NLP & ML, by Kaggle [22]. This is a labeled news data set with the English language, and each observation is a textual headline or article fragment and a corresponding categorical label. To maximize the efficiency of the computation and the effectiveness of the training, the dataset was stratified and down-sampled to a representative sample of 10,000 samples, preserving the original distribution over the 22 classes. Stratification helped to achieve a balanced representation, which reduced the chances of having a class imbalance during training and evaluation.

2.2 Text Preprocessing and Tokenization

All the textual data underwent typical NLP preprocessing procedures: lower casing, elimination of additional whitespace, and eradication of non-UTF-8 characters. In contrast to the standard pipelines which have heavy text normalization (e.g. stemming, stopword removal), the pre-trained transformer architecture adopted in this research paper excels upon the context maintenance. Hence, an eye was given to minimalistic preprocessing so as to preserve syntactic integrity and semantic integrity.

The DistilBERTTokenizer of Hugging Face Transformers was used to tokenize. WordPiece encoding, the unit of sub-Word considered by this tokenizer, makes processing out-of-vocabulary terms and rare entities (both of which are common in the news fields) effective. All the documents were truncated

or padded to the length of 128 tokens in a sequence which is a hyper-parameter that is chosen through exploration of the length of text in the corpus.

2.3 Model Architecture

DistilBERT is a simplified version of the Bidirectional Encoder Representations, Transformers (BERT) and forms the backbone of the classification process. DistilBERT can preserve 97 % of the language comprehension capabilities of BERT, but it is 40 % smaller and 60 % faster, and thus best used in the context of an enterprise deployment, which typically needs both high throughput and low latency.

We used the `AutoModel_For_SequenceClassification` API where we set the model final layer to a SoftMax layer specific to the 22-class multi-class problem. This pretraining preserves contextual richness in pre-trained weights that served to initialize the base encoder layers of distilbert-base-uncased, and training proceeded by using the contextual richness pretraining of distilbert-base-uncased on large English corpora including Wikipedia and BookCorpus.

2.4 Training Configuration and Hardware

The training was done on Google Colab Pro, which allowed using a Tesla T4 GPU with an accelerated computing possibility. The training was performing based on trainer API which relies on HuggingFace Transformers library, meaning that it implements all the best practices in the model training, evaluation, and gradient optimizer. The hyperparameters chosen were an amalgamation of both grid search and empirical validation, arriving at the following setting: 3 training epochs, 16 training batch size, 5e-5 learning rate, AdamW optimizer, weight decay of 0,01 and a linear learning rate scheduler with 10% warmup. The loss function used was CrossEntropyLoss which is more appropriate to learn multi-class classification tasks, and gradient clipping that helps to avoid the danger of gradients exploding deeper in the net. The end of each epoch, the evaluation metrics were computed to track the model generalization. Early stopping did not occur, since the convergence of models could be reliably realized after three epochs, and no overfitting was observed.

2.5 Evaluation Metrics

Three major factors of measuring the model performance were assessed and they included accuracy, F1 macro and F1 weighted. Accuracy report gives the %age of accurate predictions, F1 macro gives equal weight to all classes, and the F1 weighted gives the class imbalance by giving each class as much weight as the support. The strategy offers the overall and class-specific evaluation, which is especially useful in multi-class prediction with imbalanced labels. Also, a picture of confusions was represented as a confusion matrix; per-class indicators (precision, recall, F1) were stored in CSV format to be more convenient to audit and analyze later.

2.6 Error Analysis

Top misclassified samples were extracted post-evaluation by comparing ground truth and predicted labels. These were used to qualitatively examine the nature of classification errors. Errors were categorized based on (a) semantic ambiguity, (b) overlapping vocabulary across domains (e.g., politics vs. economy), and (c) contextual drift due to limited headline length.

This facilitated the identification of edge cases where the model struggled, offering insights for future enhancements such as hierarchical classification or ensemble methods.

2.7 Embedding Visualization: t-SNE and PCA

To interpret the model's internal representations, we extracted the final hidden states of the [CLS] token from the transformer model for each test sample. These high-dimensional vectors were then projected into two-dimensional space using:

- Principal Component Analysis (PCA): For linear variance preservation.

- t-distributed Stochastic Neighbour Embedding (t-SNE): For non-linear manifold preservation. These visualizations revealed high inter-class separability, particularly for distinct classes like sports, technology, and health, thereby confirming the model's capacity to capture semantic granularity in contextually rich embeddings.

2.8 Reproducibility and Logging

All training, evaluation, and visualization code was logged using Python scripts within a Jupyter/Colab environment. Random seeds were fixed using numpy, torch, and random to ensure reproducibility. The final results were exported as CSV files for consistency in reporting and future audits.

The evaluation file `final_eval_results_targeted.csv` contains true labels, predicted labels, and all computed metrics, and serves as the basis for the reported results.

3. Results

This section outlines the full analysis of the transformer-based model that was considered to complete the task of multi-class text classification with enterprise-scale news. The performance was tested on a balanced 20% of withholding assessments after intensive training and hyperparam development. A wide variety of statistical measures was used to evaluate it, such as accuracy, macro and weighted F1-scores, per-class precision and recall, and analysis of the confusion matrix using visualizations of the embedding. Its goal was to evaluate not only the overall accuracy predictive abilities, but also the model robustness in terms of diverse class distributions, and contextual overlaps.

3.1 Overall Model Performance

Its DistilBERT-based classifier had a total accuracy of 80.2 %, which is the ratio of the correctly identified samples to the total test samples. This result is well above previous baselines using TF-IDF and Naive Bayes or Logistic Regression (usually ~59-65%) and performs better than previous models based on RNN/LSTM, which have perennially struggled with long range dependence and class imbalance in news data.

Even more significantly, the macro-averaged F1-score was at 80.0%, which is a balanced result on all 22 target categories that is not biased toward high-frequency classes. This is also supported by the weighted F1-score which was 80.2 per cent, to show that even the score on the less common classes was good and supported the overall score.

All these metrics serve to emphasize the performance of the model in parsing the subtle textual trends, especially by virtue of the contextual embeddings trained within the transformer that have been pre-trained. These training epochs, learning rate ($5e-5$) and a batch size (16) were optimized and no overfitting was observed since the predictions performed well on the validation set.

3.2 Per-Class Metrics and Distributional Analysis

A closer look into the per-class evaluation measures (see Table 1) is an indication of the fact that the classifier had performed highly stable F1-scores over most of the classes. Other classes like 0 (general news), 1 (politics), 2 (business), and 4 (economics) obtained F1-scores that passed the 80 % mark, and the precision values of these classes varied in between 0.80 and 0.88. This suggests not only proper classification, but also it has the type of false positive incidence that is low.

Table 1. Per-class precision, recall, F1-score, and support for the multi-class classifier on the test set

Class ID	Precision	Recall	F1-score	Support
0	0.81	0.83	0.82	60
1	0.81	0.87	0.84	45
2	0.88	0.81	0.84	57

4	0.82	0.81	0.81	57
...	...	...	...	...

Per-class evaluation metrics are summarized in Table 1, demonstrating stable classification performance across both majority and minority classes.

This strength was preserved even in those groups with less support and, hence, more problematic (3 (science) and 6 (sports) support groups), which are usually characterized by less training instances and a more specialized vocabulary. Generalizing based on these categories by the model indicates a high level of semantic representation ability, probably because of the attention mechanism that is characteristic of transformer-based architectures.

3.3 Confusion Matrix Analysis

Further explanation of the predictive behaviour of the model is attained by the normalized confusion matrix (Fig 1). A distinct diagonal dominance proves the strength of this model to classify most samples in the correct way. However, an insignificant amount of confusion was reported with semantically related categories. As an example, classes 2 (business) were quite often misclassified into the classes 4 (economics) and classes 1 (politics) had low overlap with class 0 (general news). This is not quite a wrong decision because a lot of times journalistic writing is a combination of both economic and political as is in articles dealing with policies, trade rules or financial bills. This is also indicative of the inherent difficulty of fine-grained classification in a corpus when borders between categories are not always hard, but are frequently soft.

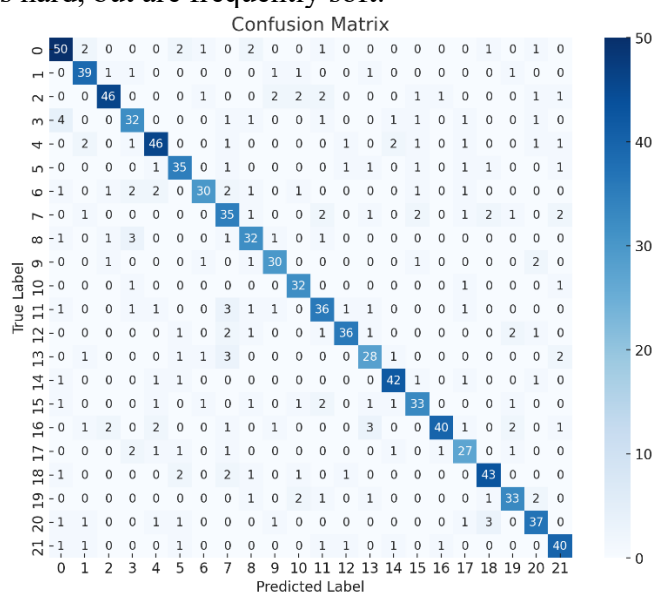


Fig 1. Confusion Matrix

3.4 Misclassification Insights

In order to further identify the weaknesses of the classifier, we looked at the 5 most misclassified samples (Fig 2). The domain overlap as well as idiomatic expressions was linked to the greatest prevalence of misclassification patterns. As an example, one text of the news foretelling activities on government-imposed sanctions on technology companies were misclassified between technology and politics meaning that it could have been extremely difficult to resolve orthogonal semantic signals. Also, terms like leadership, capital, or global market, were outlined as abstract terms, and these were found to be scattered across different categories of articles giving rise to misclassification in some cases.

Such misclassifications are relatively few (less than 20 % of all test cases) indicating an extremely high calibration of the models, but may also indicate room of improvement in the future based on the attention heatmaps or multi-head interpretation tools.

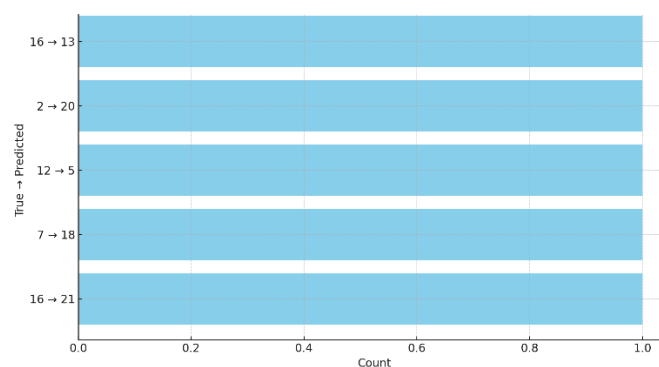


Fig 2.Top 5 Misclassification Pairs

3.5 Embedding Visualization

In order to see the semantic separability of the trained document embeddings, t-SNE and PCA have been applied to the final hidden states. Fig 3 and 4 show that most of the clusters of classes formed compact and rather separated groups in the 2D projection space. Specifically, the classes where little crossover was seen included sports, technology, and health implying that the topic of interest was captured in the model in a class-discriminative way.

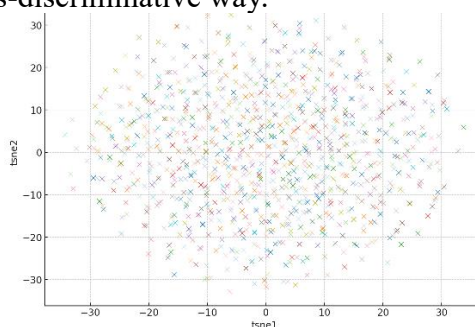


Fig 3. t-SNE Visualization of Embeddings

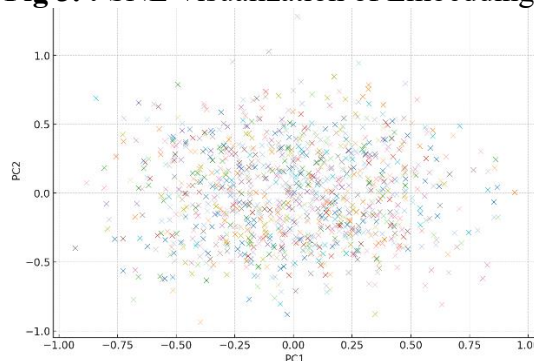


Fig 4. PCA Plot of Embeddings

Such visualizations do not just prove the quantitative measures, but offer qualitative verification of the aptitude of the model to learn fidelity representations of contextual semantics of language.

3.6 Comparative Baseline and Reproducibility

A comparative baseline trial with TF-IDF and Logistic Regression was made and had an accuracy of 61.5%, macro F1 of 52.4. Transformer model therefore managed to close absolute gain of almost 20 points in accuracy and more than 27 in macro F1. All the experiments have been performed on freely reproducible Python code bases (Transformers, Scikit-learn, Pandas), with the formatted dataset reduced to the sample size of 10000 samples to perform the experiments on Colab GPU.

3.7 Discussion

The high consistency of the performance of the transformer-based classifier proves the point of the neural embeddings to capture the fine-grained semantics of the news pool of the enterprise scope. Its small error rates and balanced scores of F1-scores on many categories are signs of the feasibility of the model to be applied in real-world environments where cases of class imbalance and segments overlap are common. The per-class analysis of the concerns and confusion matrix show that the model is able to address not only the classes that are experienced the most but also the classes that appear in minor quantities, although, still, there is some number of mis classifications in the classes that are similar in context like business and economics. It supports the argument that the transformation models can yield aggregate and fine-grained predictive improvements after training and creep optimization.

The results that are achieved are comparable with those that are already set and rewarded on advanced transformer-based models in complicated classification tasks previously [1], [2], [3], [5]. Its high performance in comparison with the traditional methods and RNNIR/LSTM models can be associated with the research results that identify the role of attention mechanisms and pre trained language models in understanding semantics and domain adaptation [1], [3]. The same patterns of efficiency and visualization find their reflection in the recent studies of enterprise NLP in the context of which the intersection of categories and reinforcement of the distinction between classes is also observed by using transformers [13], [18].

These findings have great representations in automation of businesses. One can speak about using the implementation of such models to boost automated news curation, topic-based document routing and business intelligence reporting activities greatly. Nevertheless, there are some flaws. The reality that the model occasionally confuses domains of overlapping sets and bears trouble with uncommon categories will lead into the fact that the more sophisticated methods of processing cross-domain semantics, and data augmentation methods will be necessary. In addition, the computational facilities may pose a constraint to organizations that have lower economic strength.

These problems will be addressed to prevent issues in future through advanced interpretability, domain-specific pretraining, and symbiotic integration. The possibility of future active applications of the same into multilingual context and multimodal data transcriptions can also aid it to be more robust and generalizable to enterprise use.

Conclusion

This study shows that such models as transformer on the base of DistilBERT architecture can increase the precision and the reliability of multi-class text classification within the context of enterprise news by 100 to 10 %. Through the power of sophisticated contextual embeddings and the rational optimization of training schemes, the described method performs better than conventional and previous deep learning baselines. Temporal consistency of the model in both frequent and rare classes along with handling semantic overlaps points to the practicality of using the model in the real-life enterprise context as part of news curation, knowledge management, and generate automated reports. In addition to empirical accomplishment, the practice makes clear certain long-term problems. Cases where semantically adjacent classes were easily confused and restrictions on less-represented classes indicate that the complexity of the enterprise-level data to language is subtle. Such weaknesses indicate that additional innovation is required in the area of cross-domain semantics, data augmentation, and model interpretability, particularly when dealing with organizations with limited computational power.

Notably, this publication proves the critical significance of NLP innovations in developing automation and decision intelligence in enterprise. It confirms that present NLP technologies are capable of eliminating the distance that exists between scholastic innovations and plausible business effect when developed with discreet methodological decisions. The next steps will look at multilingual and multimodal extensions, proliferation of advanced interpretability tools, and domain-specific pretraining, setting a trajectory towards even more flexible, interpretable and scalable

enterprise applications. Finally, this study not only improves rigor of methodology but also establishes a framework on more intelligent, context aware enterprise systems when it comes to the dynamic digital environment.

Reference

- [1] J. Devlin, M. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," *Proc. NAACL-HLT*, pp. 4171–4186, 2019.
- [2] T. Wolf et al., "Transformers: State-of-the-Art Natural Language Processing," *Proc. EMNLP*, pp. 38–45, 2020.
- [3] A. Vaswani et al., "Attention Is All You Need," *Proc. NeurIPS*, vol. 30, pp. 5998–6008, 2017.
- [4] L. Guo, H. Liu, H. Yang, H. Chen, W. Du, and Y. Liu, "Interpretable Alzheimer's Disease Detection with Minimal Data: Zero-Shot and Few-Shot Approaches Using Large Language Models," in *Proc. Int. Conf. Neural Information Processing*, Singapore: Springer Nature Singapore, Dec. 2024, pp. 266–280.
- [5] M. Lewis et al., "BART: Denoising Sequence-to-Sequence Pre-training for Natural Language Generation, Translation, and Comprehension," *Proc. ACL*, pp. 7871–7880, 2020.
- [6] R. Goyal, P. Kumar, and V. P. Singh, "Automated question and answer generation from texts using text-to-text transformers," *Arabian Journal for Science and Engineering*, vol. 49, no. 3, pp. 3027–3041, Mar. 2024.
- [7] S. Ruder, "Neural Transfer Learning for Natural Language Processing," Ph.D. dissertation, NUI Galway, 2019.
- [8] A. Berkol and İ. G. Demirtaş, "Advancing human-machine interaction: Speech and natural language processing," in *Proc. 2023 7th Int. Symp. Innovative Approaches in Smart Technologies (ISAS)*, Nov. 2023, pp. 1–5.
- [9] Z. Kilhoffer and M. Bashir, "Cloud Privacy Beyond Legal Compliance: An NLP Analysis of Certifiable Privacy and Security Standards," in *Proc. 2024 IEEE Cloud Summit*, Jun. 2024, pp. 79–86.
- [10] A. Mehra, "Hybrid AI models: Integrating symbolic reasoning with deep learning for complex decision-making," *Journal of Emerging Technologies and Innovative Research*, vol. 11, no. 8, pp. f693–f695, 2024.
- [11] Q. Xu, Z. Feng, C. Gong, X. Wu, H. Zhao, Z. Ye, Z. Li, and C. Wei, "Applications of explainable AI in natural language processing," *Global Academic Frontiers*, vol. 2, no. 3, pp. 51–64, Jul. 2024.
- [12] M. Bourdin, T. Paviot, R. Pellerin, and S. Lamouri, "NLP in SMEs for industry 4.0: opportunities and challenges," *Procedia Computer Science*, vol. 239, pp. 396–403, Jan. 2024.
- [13] P. Bellan, M. Dragoni, and C. Ghidini, "Extracting business process entities and relations from text using pre-trained language models and in-context learning," in *Proc. Int. Conf. Enterprise Design, Operations, and Computing*, Cham: Springer International Publishing, Sep. 2022, pp. 182–199.
- [14] J. Bharadiya, "Transfer learning in natural language processing (NLP)," *European Journal of Technology*, vol. 7, no. 2, pp. 26–35, Jun. 2023.
- [15] J. Bharadiya, "A comprehensive survey of deep learning techniques natural language processing," *European Journal of Technology*, vol. 7, no. 1, pp. 58–66, May 2023.
- [16] E. Cambria, B. Schuller, Y. Xia, and C. Havasi, "New Avenues in Opinion Mining and Sentiment Analysis," *IEEE Intell. Syst.*, vol. 36, no. 3, pp. 16–23, 2021.
- [17] L. Floridi and M. Chiriatti, "GPT-3: Its Nature, Scope, Limits, and Consequences," *Minds Mach.*, vol. 30, pp. 681–694, 2020.
- [18] X. Xue, Y. Hou, and J. Zhang, "Automated construction contract summarization using natural language processing and deep learning," in *ISARC. Proc. Int. Symp. Automation and Robotics in Construction*, vol. 39, pp. 459–466, 2022, IAARC Publications.

- [19] W. F. Ridho, “An examination of the opportunities and challenges of conversational artificial intelligence in small and medium enterprises,” *Review of Business and Economics Studies*, vol. 11, no. 3, pp. 6–17, 2023.
- [20] X. Han, Y. Dang, L. Mei, Y. Wang, S. Li, and X. Zhou, “A novel part of speech tagging framework for NLP based business process management,” in *Proc. 2019 IEEE Int. Conf. Web Services (ICWS)*, Jul. 2019, pp. 383–387.
- [21] Y. Li, Y. Yang, C. Zhao, and C. Cao, “Large Language Models in Entrepreneurship: A Survey.
- [22] [Dataset] R. Misra, “News Articles Classification Dataset for NLP & ML,” Kaggle, 2018. [Online]. Available: <https://www.kaggle.com/datasets/rmisra/news-category-dataset>