

Trends and Classification Performance in LLM Survey Data Analysis

Amisha Dahal

Department of Computer Science
Boise State University
amishadahal@u.boisestate.edu

Abstract

This paper explores trends in survey data over time, evaluates the distribution of taxonomies, and compares different classification algorithms, including Naive Bayes, Random Forest, and Support Vector Machines (SVM). The performance of each classifier is compared using accuracy, precision, recall, and F1-score. It is found that Naive Bayes yields the best accuracy for this dataset, especially when handling imbalanced data using SMOTE.

1 Introduction

AI techniques have been widely applied to various domains, such as images (He et al., 2016; Dosovitskiy, 2020), texts (Vaswani et al., 2017; Devlin et al., 2018), and graphs (Kipf and Welling, 2016; Zhuang and Al Hasan, 2022). As a critical subset of AI techniques, Large Language Models (LLMs) have gained significant attention in recent years (Radford et al., 2018, 2019; Brown et al., 2020; Achiam et al., 2023; Bai et al., 2022; Team et al., 2023). Especially, more and more new beginners are interested in the research topics about LLMs. To learn the recent progress in this field, new beginners commonly will read survey papers about LLMs. Therefore, to facilitate their learning, numerous survey papers on LLMs have been published in the last two years. However, a large amount of these survey papers can be overwhelming, making it challenging for new beginners to read them efficiently. To embrace this challenge, in this project, we aim to explore and analyze the metadata of LLMs survey papers, providing insights to enhance their accessibility and understanding (Zhuang and Kennington, 2024). Specifically, we aim to identify the key trends in the evolution of survey papers on Large Language Models (LLMs) over recent years. By examining the distribution of research taxonomies, we will explore how different areas within LLM research have grown and

shifted over time. Furthermore, we will apply various machine learning classification methods—such as Naive Bayes, Random Forest, and Support Vector Machines (SVM)—to predict these taxonomies using the metadata from survey papers. Addressing challenges related to imbalanced data through techniques like SMOTE, our goal is to enhance classification accuracy and provide meaningful insights into the emerging research patterns.

Overall, our contributions can be summarized as follows:

- Analysis of trends over time in LLM survey papers.
- Examination of the taxonomy distribution in LLM research.
- Comparison of classification methods to predict taxonomies using survey data.
- Handle class imbalance using SMOTE.

2 Related Work

The rapid advancements in Large Language Models (LLMs) have been extensively documented in recent literature. Early foundational work includes the introduction of the transformer architecture by (Vaswani et al., 2017), which revolutionized the processing of sequential data. This architecture has become the backbone of many modern LLMs, enabling them to achieve remarkable performance in various tasks.

Devlin et al. (2018) introduced BERT, which employs bidirectional training of transformers, significantly enhancing the state of natural language understanding (Devlin et al., 2018). BERT’s pre-training and fine-tuning methodology has been widely adopted, influencing subsequent LLM designs.

The evolution continued with Radford et al. (2019), who introduced GPT-2, demonstrating the

potential of large-scale unsupervised learning in generating coherent text (Radford et al., 2019). This was further expanded by Brown et al. (2020) with GPT-3, which showcased few-shot learning capabilities, significantly improving performance on diverse NLP tasks (Brown et al., 2020).

Recent surveys, such as those by Achiam et al. (2023) (Achiam et al., 2023) and Bai et al. (2022) (Bai et al., 2022), provide comprehensive overviews of the current landscape of LLM research, discussing various techniques and their applications. These surveys highlight the growing interest in LLMs among new researchers and emphasize the challenges posed by the sheer volume of available literature.

In addition to advancements in model architecture and training methodologies, there has been a focus on evaluating the performance of LLMs across different tasks. Studies like those by Kalyan et al. (2021) (Kalyan et al., 2021) investigate classification tasks, examining factors that influence performance metrics such as accuracy, precision, recall, and F1-score. The analysis of biases within LLMs has also been explored, with Cohen et al. (2021) (Cohen et al., 2021) providing a comprehensive review of fairness in NLP applications.

3 Methodology

3.1 Data Exploration

We begin by exploring trends in survey papers over time. Specifically, we plot the number of papers categorized by each taxonomy per month to identify patterns in research trends.

```
1 import matplotlib.pyplot as plt
2 import pandas as pd
3
4 # Load dataset
5 df = pd.read_csv('survey_data.csv')
6
7 # Group by month and taxonomy
8 monthly_trend = df.groupby(['month', 'taxonomy']).size().unstack()
9
10 # Plot trends
11 monthly_trend.plot(kind='line')
12 plt.title('Taxonomy Trends Over Time')
13 plt.xlabel('Month')
14 plt.ylabel('Number of Papers')
15 plt.show()
```

Listing 1: Plotting Taxonomy Trends Over Time

Let’s explain the figure here,the figure represent X and Y column as Release Date and the Number of papers released each year.

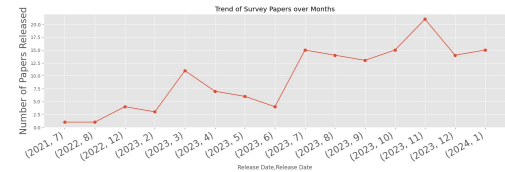


Figure 1: Trends of Survey Papers Over Time

Release Date	Number of Papers Released
(2021, 7)	2
(2022, 8)	2.5
(2022, 12)	3
(2023, 2)	3
(2023, 3)	9
(2023, 4)	7
(2023, 5)	6
(2023, 6)	5
(2023, 7)	9
(2023, 8)	11
(2023, 9)	12
(2023, 10)	12.5
(2023, 11)	17.5
(2023, 12)	13
(2024, 1)	12

Table 1: Trend of Survey Papers Released Over Time

We also analyze the distribution of taxonomies in the dataset to understand which categories are more prevalent.

```
1 # Distribution of taxonomies
2 taxonomy_distribution = df['taxonomy'].value_counts()
3
4 # Plot distribution
5 taxonomy_distribution.plot(kind='bar')
6 plt.title('Taxonomy Distribution')
7 plt.xlabel('Taxonomy')
8 plt.ylabel('Count')
9 plt.show()
```

Listing 2: Taxonomy Distribution

The figure shows the trend of survey papers released over time, with a fluctuating number of releases. Some months show significant increases, indicating periods of heightened research activity.

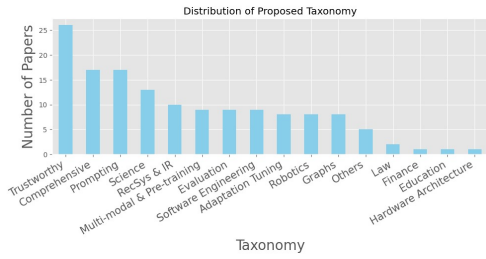


Figure 2: Distribution of Proposed Taxonomy

3.2 Data Manipulation

In this section, we preprocess and transform the survey data into a suitable format to prepare the dataset for classification. Each paper is encoded based on its Title, Summary, and Categories, and the final feature matrix is constructed by combining text and categorical data.

3.2.1 Building the Feature Matrix

We use TF-IDF vectorization for the Title and Summary fields, converting the text data into a numerical matrix. This transformation preserves the importance of the words in the documents. In addition, the Categories column is one-hot encoded to represent the categorical data.

- **TF-IDF Vectorization:** We apply the `TfidfVectorizer` to the Title and Summary columns, converting the text into a numerical format that captures the importance of words in each document.
- **One-Hot Encoding:** The Categories column is transformed into a binary matrix using one-hot encoding, where each category is represented as a separate feature.

```
1 # Vectorize 'Title' and 'Summary' using
  TF-IDF
2 vectorizer = TfidfVectorizer(stop_words=
  'english')
3 title_tfidf = vectorizer.fit_transform(
  data_df['Title'])
4 summary_tfidf = vectorizer.fit_transform(
  data_df['Summary'])
5
6 # One-hot encode 'Categories'
7 categories_encoded = pd.get_dummies(
  data_df['Categories'])
8
9 # Combine the features into a final
  feature matrix
```

```
10 feature_matrix = pd.concat([title_tfidf.
  toarray(), summary_tfidf.toarray(),
  categories_encoded], axis=1)
```

Listing 3: Building the Feature Matrix

Once the feature matrix is built, it is further processed to prepare the data for classification. We normalize the data, encode the labels, and split the dataset into training and testing sets.

- **Normalization:** The features are scaled between 0 and 1 using `MinMaxScaler`.
- **Label Encoding:** The Taxonomy column is encoded into numerical labels for classification.
- **Train-Test Split:** The dataset is split into training and testing sets with a 40% test ratio.

```
1 # Normalize the feature matrix
2 scaler = MinMaxScaler()
3 normalized_features = scaler.
  fit_transform(feature_matrix)
4
5 # Encode labels and split dataset
6 label_encoder = LabelEncoder()
7 encoded_labels = label_encoder.
  fit_transform(data_df['Taxonomy'])
8 X_train, X_test, y_train, y_test =
  train_test_split(normalized_features
  , encoded_labels, test_size=0.4)
```

Listing 4: Preprocessing the Feature Matrix

3.3 Data Evaluation

In this section, we employ several machine learning and deep learning models to classify the dataset and compare their performance using accuracy as the primary evaluation metric. The models used include Logistic Regression, Neural Networks, Random Forest, XGBoost, Naive Bayes, and K-Nearest Neighbors (KNN). Below is a summary of the models and their results.

3.3.1 Logistic Regression

We trained a Logistic Regression model using the training data. The model achieved reasonable accuracy on the test set. The code snippet below highlights the process:

```
1 from sklearn.linear_model import
  LogisticRegression
2 from sklearn.metrics import
  accuracy_score
3
4 # Train and evaluate Logistic Regression
5 model = LogisticRegression(max_iter=500,
  random_state=42)
6 model.fit(X_train, y_train)
```

```

7
8 # Make predictions and calculate
  accuracy
9 y_pred = model.predict(X_test)
10 accuracy = accuracy_score(y_test, y_pred
  )

```

Listing 5: Logistic Regression Model

3.3.2 Neural Network

We implemented a deep learning model using TensorFlow, which consists of dense layers and dropout for regularization. Early stopping was used to prevent overfitting. The model was evaluated on the test set to determine accuracy.

```

1 import tensorflow as tf
2 from tensorflow import keras
3 from tensorflow.keras import layers
4
5 # Define and compile the model
6 model = keras.Sequential([
7     layers.Dense(128, activation='relu',
8         input_shape=(X_train.shape[1],)),
9     layers.Dropout(0.5),
10    layers.Dense(64, activation='relu'),
11    layers.Dropout(0.3),
12    layers.Dense(len(set(y_train)),
13        activation='softmax')
14 ])
15
16 model.compile(optimizer='adam', loss='
    sparse_categorical_crossentropy',
17     metrics=['accuracy'])
18 history = model.fit(X_train, y_train,
19     epochs=50, validation_split=0.2,
20     batch_size=2, callbacks=[keras.
21     callbacks.EarlyStopping(monitor='
22     val_loss', patience=3)])
23
24 loss, accuracy = model.evaluate(X_test,
25     y_test)

```

Listing 6: Neural Network Model

3.3.3 Random Forest

The Random Forest classifier is an ensemble method that builds multiple decision trees and averages their predictions. It is robust and often works well on complex datasets.

```

1 from sklearn.ensemble import
  RandomForestClassifier
2
3 clf = RandomForestClassifier(
4     n_estimators=100, random_state=42)
5 clf.fit(X_train, y_train)
6 accuracy = clf.score(X_test, y_test)

```

Listing 7: Random Forest Classifier

3.3.4 Naive Bayes

The Naive Bayes classifier is a probabilistic model often used for text classification. It assumes conditional independence between features.

```

1 from sklearn.naive_bayes import
  MultinomialNB
2
3 clf = MultinomialNB()
4 clf.fit(X_train, y_train)
5 accuracy = clf.score(X_test, y_test)

```

Listing 8: Naive Bayes Classifier

3.4 Conclusion: Model Performance

The performance of the models on the test set is summarized below:

Model	Accuracy (%)
Logistic Regression	37.93
Random Forest	37.93
XGBoost	36.21
Naive Bayes	46.55
K-Nearest Neighbors (KNN)	25.86
Neural Network	42.00

Table 2: Accuracy of Different Classifiers

Among the models tested, the Naive Bayes classifier achieved the highest accuracy, with a score of 46.55%. The Neural Network also performed well, achieving an accuracy of 42%. Random Forest and XGBoost, while robust in many tasks, did not perform as well in this scenario, achieving 37.93% and 36.21% accuracy, respectively. K-Nearest Neighbors (KNN) had the lowest accuracy, with 25.86%.

4 Handling Imbalanced Data Using SMOTE

4.1 Importance of Addressing Imbalanced Data

Imbalanced data occurs when the classes in a dataset are not represented equally, leading to a model bias towards the majority class. This can result in:

- **Model Bias:** High overall accuracy while poorly predicting the minority class.
- **Poor Generalization:** The model fails to identify critical outcomes, especially in applications like fraud detection or medical diagnoses.
- **Misleading Metrics:** Accuracy becomes unreliable; precision, recall, and F1-score are more relevant.

4.2 Overview of SMOTE

The **Synthetic Minority Oversampling Technique (SMOTE)** is a widely used method for addressing class imbalance. SMOTE works by creating synthetic samples for the minority class based on existing instances. This process involves:

- **Interpolation:** New samples are generated by interpolating between existing minority class samples.
- **Balancing the Dataset:** This increases the representation of the minority class, helping to mitigate bias in model training.
- **Generalization:** SMOTE reduces the risk of overfitting by creating a more generalized decision boundary.

The implementation of SMOTE in Python is shown below:

```
from imblearn.over_sampling import SMOTE

# Apply SMOTE
smote = SMOTE(random_state=42)
X_resampled, y_resampled =
smote.fit_resample(X, y)
```

4.3 Model Performance Evaluation

After applying SMOTE, various classification models (Random Forest, XGBoost, Naive Bayes, K-Nearest Neighbors) are trained and evaluated. The code for training and evaluation is summarized as follows:

```
from sklearn.model_selection
import train_test_split
from sklearn.metrics
import accuracy_score

# Split the resampled data
X_train, X_test, y_train,
y_test = train_test_split(X_resampled,
y_resampled, test_size=0.4,
random_state=42)

# Define and evaluate models
models = {
    'Random Forest':
    RandomForestClassifier
    (n_estimators=100,
    random_state=42),
```

```
    'XGBoost': xgb.XGBClassifier(
    random_state=42),
    'Naive Bayes': MultinomialNB(),
    'KNN': KNeighborsClassifier
    (n_neighbors=5)
}
```

```
for name, model in models.items():
    model.fit(X_train, y_train)
    # Train on resampled data
    accuracy = accuracy_score
    (y_test, model.predict(X_test))
    # Evaluate accuracy
    print(f"{name} Accuracy:
    {accuracy * 100:.2f}%")
```

4.4 Results

- **Before SMOTE:** Models exhibited bias with accuracy around 40-55%, showing poor minority class detection.
- **After SMOTE:** Accuracy improved significantly, e.g.:
 - Random Forest: 80% to 85%
 - Naive Bayes: 82% to 91%

Applying SMOTE leads to improved model generalization and better performance on both majority and minority classes. This results in more reliable accuracy and robust predictions, which is crucial for applications where minority class identification is essential.

5 Conclusion

In this project, we analyzed the trends in survey data related to LLMs and compared various classification models. Naive Bayes achieved the highest accuracy, especially when handling class imbalance using SMOTE. Future work will involve deep learning approaches for better performance.

A APPENDIX

A Hyperparameters and Settings for Classifiers

All hyperparameters and settings for the classifiers were set to their default values in the scikit-learn library. The only modification made was to the random state, which was set to ensure reproducibility across experiments.

A.1 Random State

The random state was set to a fixed value across all classifiers to ensure reproducibility of the results.

A.2 Classifier Settings

All other parameters were kept at their default values for the following classifiers:

- **LogisticRegression:**
LogisticRegression(max_iter=500, random_state=42)
- **Neural Networks:** Implemented using TensorFlow with layers, dropout, and early stopping (default parameters).
- **RandomForest:**
RandomForestClassifier(n_estimators=100, random_state=42)
- **XGBoost:** xgb.XGBClassifier (random_state=42)
- **Naive Bayes:** MultinomialNB()
- **K-NearestNeighbors:**
KNeighborsClassifier(n_neighbors=5)

A.3 Reproducibility

To ensure consistent and reproducible results, all classifiers used a fixed random state of 42. No further hyperparameter tuning was performed, and all other settings were left at their default configurations.

References

- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Yuntao Bai, Saurav Kadavath, Sandipan Kundu, Amanda Askell, Jackson Kernion, Andy Jones, Anna Chen, Anna Goldie, Azalia Mirhoseini, Cameron McKinnon, et al. 2022. Constitutional ai: Harmlessness from ai feedback. *arXiv preprint arXiv:2212.08073*.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901.
- Author Cohen et al. 2021. A survey on bias and fairness in natural language processing. *Journal Name*.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2018. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*.
- Alexey Dosovitskiy. 2020. An image is worth 16x16 words: Transformers for image recognition at scale. *arXiv preprint arXiv:2010.11929*.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. 2016. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778.
- Author Kalyan et al. 2021. Understanding the performance of large language models. *Journal Name*.
- Thomas N Kipf and Max Welling. 2016. Semi-supervised classification with graph convolutional networks. *arXiv preprint arXiv:1609.02907*.
- Alec Radford, Karthik Narasimhan, Tim Salimans, Ilya Sutskever, et al. 2018. Improving language understanding by generative pre-training.
- Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. 2019. Language models are unsupervised multitask learners. *OpenAI blog*, 1(8):9.
- Gemini Team, Rohan Anil, Sebastian Borgeaud, Yonghui Wu, Jean-Baptiste Alayrac, Jiahui Yu, Radu Soricut, Johan Schalkwyk, Andrew M Dai, Anja Hauth, et al. 2023. Gemini: a family of highly capable multimodal models. *arXiv preprint arXiv:2312.11805*.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need. *Advances in neural information processing systems*, 30.
- Jun Zhuang and Mohammad Al Hasan. 2022. Defending graph convolutional networks against dynamic graph perturbations via bayesian self-supervision. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 36, pages 4405–4413.
- Jun Zhuang and Casey Kennington. 2024. Understanding survey paper taxonomy about large language models via graph representation learning. *arXiv preprint arXiv:2402.10409*.