

Application of Natural Language Processing with R in clinical trials

Feifeng Zhou, Zhenchao Chen, Pfizer (China) Research and Development Co.,Ltd. China

ABSTRACT

Artificial intelligence (AI) technology has been playing pivotal role in clinical trials, where the natural language processing (NLP) is increasingly leveraged by R language to conduct text classification. NLP enables a shift from time-consuming manual and siled management of natural language data to automated and standard processes for analyzing text and pharmacopoeia codes.

In this paper, a classifier using NLP in R is developed to classify drugs into several study-specific classifications by the names of concomitant drugs collected from CRF that remains unstructured or after WHODRUG coded. The variables for drug classification and relationship to the study drug of interest are then generated and can be used for clinical review or summary tables. Overall, thanks to comprehensive NLP text classification R packages, statistical programmers can efficiently process and analyze ongoing trial data, providing strong support for all stages of clinical trials.

INTRODUCTION

In the realm of clinical trials, the analysis of vast amounts of textual data, including patient records, doctor's notes, and clinical reports, holds immense potential for enhancing the efficiency and accuracy of research outcomes. Recent advancements in natural language processing (NLP) techniques, coupled with the flexibility and power of the R programming language, have opened new doors for harnessing this data [1]. When combined with NLP techniques, R enables the extraction of meaningful insights from unstructured textual data, which can then be utilized for clinical decision-making, patient outcome prediction, and drug discovery [2].

The application of R language NLP in clinical trials involves several key steps. Firstly, the textual data is preprocessed to remove noise and standardize the format. This includes tasks such as tokenization, stemming, and stop-word removal. Next, NLP algorithms are applied to extract relevant information, such as identifying medical entities, relationships between entities, or sentiment analysis. The extracted information is then analyzed using statistical techniques to gain insights into clinical trends, patient behavior, or drug effectiveness.

This paper aims to explore the application of NLP in drug classification and its integration with R for clinical trial data analysis [3]. In the proceeding sections, we will demonstrate how to implement NLP with R to categorize drug names into specific categories by following a structured approach involving text preprocessing, feature extraction, model training, and application on new data, and provide some simple examples.

OVERVIEW OF NATURAL LANGUAGE PROCESSING (NLP)

Natural Language Processing (NLP) is a subfield of artificial intelligence (AI) and computational linguistics concerned with the interaction between computers and human languages [4]. It enables computers to understand, interpret, and generate human language in a way that is valuable and meaningful. NLP encompasses a wide range of tasks, including but not limited to:

- Text Classification: Categorizing texts into predefined categories based on their content.
- Named Entity Recognition (NER): Identifying and classifying named entities (such as drugs, diseases, or symptoms) within text.
- Sentiment Analysis: Determining the sentiment expressed in text (positive, negative, neutral).
- Information Extraction: Structuring and extracting relevant information from unstructured text data.
- Machine Translation: Translating text from one language to another.

- Question Answering: Finding precise answers to questions posed in natural language.

NLP IN HEALTHCARE AND CLINICAL TRIALS

NLP has significant applications within the healthcare sector, particularly in clinical trials, where large volumes of textual data are generated and analyzed. Here are some key applications:

1. Clinical Data Extraction and Integration:
 - Electronic Health Records (EHR): NLP can extract structured data from unstructured clinical notes in EHRs, enabling comprehensive patient profiles and facilitating clinical decision-making [4].
 - Clinical Trial Protocols: Automating the extraction of key information from clinical trial protocols, ensuring compliance and consistency in study execution.
2. Clinical Trial Recruitment and Eligibility Criteria:
 - Patient Matching: Using NLP to match patient profiles with trial eligibility criteria, thereby enhancing recruitment efficiency and reducing screening times.
 - Semantic Search: Improving the accuracy of searching and matching eligible patients based on complex criteria expressed in natural language.
3. Pharmacovigilance and Drug Safety:
 - Adverse Event Monitoring: NLP helps in analyzing adverse event reports and identifying potential safety concerns related to drugs or treatments.
 - Drug-Drug Interactions: Identifying and predicting potential interactions between drugs using NLP techniques applied to drug labels, clinical notes, and literature.
4. Drug Classification and Coding:
 - WHODRUG Coding: Automating the classification of drugs based on WHODRUG standards using NLP techniques for consistency and accuracy.
 - Concomitant Medication Identification: Identifying and classifying concomitant medications mentioned in clinical trial data to understand their impact on study outcomes.
5. Clinical Trial Data Management and Analysis:
 - Text Mining [5]: Extracting insights from unstructured text data collected during trials, such as patient-reported outcomes, clinician notes, and adverse event descriptions.
 - Real-time Data Analysis: Enabling real-time monitoring and analysis of trial data for early detection of trends or anomalies.

POPULAR NLP LIBRARIES IN R

When discussing popular NLP libraries in R, several tools and libraries stand out for their functionality, ease of use, and community support. Here's an expanded look at some of these libraries:

Package Name	Overview	Key Features
tm (Text Mining Infrastructure in R [6])	A fundamental package for text mining in R. It provides a framework for managing text data and performing various preprocessing tasks.	• Document-term matrix creation. • Text cleaning, stemming, and stopword removal. • Integration with other R packages for statistical analysis.

Package Name	Overview	Key Features
Quanteda [7]	Quanteda is an R package for quantitative analysis of textual data, offering functionalities for text corpus management, document-feature matrix creation, and text analysis.	• Fast tokenization and manipulation of text data. • Support for complex corpus structures and text scaling. • Integration with tidyverse principles for data handling.
text2vec [8]	It focuses on efficient text vectorization, particularly suitable for large-scale text data.	• Fast vectorization methods like GloVe, word2vec, and others. • Parallel processing for scalability. • Integration with dplyr for easy data manipulation.
NLP [9]	A basic package that provides fundamental functions and classes for natural language processing tasks.	• Tokenization, stemming, and part-of-speech tagging. • Named entity recognition. • Integration with tm package for text mining tasks.
RWeka	It provides an interface to the Weka machine learning library, which includes various tools for text classification and clustering.	• Implementation of classifiers like SVM, Naive Bayes, etc., for text data. • Support for feature selection and evaluation of models. • Integration with Weka's extensive range of algorithms.
keras [10] / tensorflow [11]	While primarily known for deep learning tasks, keras and tensorflow can be utilized for advanced NLP tasks such as sequence modeling, sentiment analysis, and text generation [12].	• Implementation of state-of-the-art deep learning models (e.g., LSTM, GRU) for NLP. • Pre-trained models (e.g., BERT, GPT) available via tensorflow-hub. • Integration with GPU for accelerated computations.
tidytext	It integrates text mining and analysis into the tidyverse framework, facilitating seamless workflows [5].	• Conversion of text data into tidy data structures. • Sentiment analysis and other text mining operations. • Integration with ggplot2 for visualization.
SpaCyR [13]	It provides an interface to the Python spaCy library, offering fast and efficient natural language processing capabilities.	• Tokenization, part-of-speech tagging, and named entity recognition. • Dependency parsing and entity linking. • Integration with other R packages for further analysis.

Package Name	Overview	Key Features
textmineR	It focuses on text mining tasks, including topic modeling and document clustering.	• Implementation of LDA (Latent Dirichlet Allocation) and other topic modeling algorithms. • Automated text cleaning and preprocessing. • Integration with ggplot2 for visualization.
wordVectors	It provides functionality for working with pre-trained word embeddings, such as Word2Vec models [14].	• Loading and querying word embeddings. • Similarity calculations between words. • Integration with other text analysis workflows in R.

Table 1. Popular NLP libraries in R

Each of these libraries in R offers unique strengths and capabilities, catering to different aspects of NLP tasks from basic text processing to advanced machine learning applications. Depending on your specific needs — whether it's text cleaning, sentiment analysis, or deep learning-based tasks — leveraging the right combination of these libraries can significantly streamline your NLP projects in R.

CASE STUDY: IMPLEMENTING NLP WITH R IN DRUG CLASSIFICATION

DATA COLLECTION AND TEXT PREPROCESSING

Assuming you have drug names and corresponding categories in a dataframe (df), where drug_name is the column containing drug names and category is the column containing corresponding labelled categories. The example below will use the following sample data:

```
# LOAD NECESSARY LIBRARIES
# install.packages(c("tm", "textstem", "caret", "e1071"))
library(tm)
library(textstem)
library(caret)
library(e1071)

# Example dataframe
df <- data.frame(drugname = c("Drug A", "Drug B", "Drug C", "Drug D"), category
= c("Category1", "Category2", "Category1", "Category3"))
```

Text preprocessing in Natural Language Processing (NLP) is a crucial step that involves cleaning and transforming text data into a format which could be applied to machine learning algorithms [5]. For example, we can clean raw data by removing special characters, digits, punctuation marks, converting to lowercase for uniformity, and handling abbreviations and synonyms (e.g., mapping "acetaminophen" to "paracetamol"). Here's a simple example to clean and prepare the drug names for modeling:

```
# Create a corpus from drug names
corpus <- Corpus(VectorSource(df$drug_name))

# Preprocessing steps
# Convert to lowercase like treat "Drug" and "drug" as the same token
corpus <- tm_map(corpus, content_transformer(tolower))
# Remove punctuation
corpus <- tm_map(corpus, removePunctuation)
# Remove numbers
corpus <- tm_map(corpus, removeNumbers)
# Remove English stopwords
```

```
corpus <- tm_map(corpus, removeWords, stopwords("en"))
# Stemming
corpus <- tm_map(corpus, stemDocument)
# Strip extra whitespace
corpus <- tm_map(corpus, stripWhitespace)

# Convert back to a dataframe with preprocessed drug names
df$clean_drugname <- sapply(corpus, as.character)
```

TOKENIZATION AND LEMMATIZATION

Tokenization is the process of breaking a text (in this case, drug names) into smaller units called tokens. These tokens are usually words or subwords that are meaningful for analysis [5]. For example, consider the drug name "Acetaminophen 500mg", tokenization would break this into tokens such as ["Acetaminophen", "500", "mg"].

Tokenization is important for uniform representation as it ensures that each drug name is represented consistently, making it easier for classification algorithms to process [5]. Also, tokens serve as the basic features for subsequent steps in NLP, such as feature vectorization or embedding. By implementing tokenization, we can use packages like **tokenizers** or regular expressions (e.g., **stringr**, **text**, etc.).

It is worth mentioning in this step that we need to pay attention to the following aspects:

- Abbreviations and Symbols: Drug names often include abbreviations (e.g., "mg", "ml") and symbols (e.g., "-", "/", "(", ")") that require careful handling to avoid splitting tokens incorrectly.
- Multi-word Names: Some drug names consist of multiple words (e.g., "Erythromycin Ethylsuccinate"). Tokenization should handle these as single entities when necessary.
- Non-standard Naming Conventions: Occasionally, drug names may not follow standard grammar rules or conventions, requiring special rules or pre-processing steps.
- After tokenization, validate the tokens against a known dictionary or vocabulary to ensure accuracy and completeness.

Here's a simple example to clean and prepare the drug names for modeling:

```
# Example code for tokenization and lemmatization
library(tokenizers)
text <- "acetaminophen and ibuprofen"
tokens <- tokenize_words(text)
print(tokens)

["acetaminophen" "and" "ibuprofen"]
```

Output 1. Output from an example of tokenization

In summary, tokenization of drug names in the context of NLP involves breaking down text into manageable units (tokens) to facilitate subsequent analysis such as classification into study-specific categories. It's a critical preprocessing step that ensures the accuracy and effectiveness of NLP applications in pharmaceutical and healthcare research.

FEATURE EXTRACTION

Feature extraction involves transforming preprocessed text data into numerical or vector representations that programmers could apply machine learning algorithms effectively afterwards. In other words, we need to convert the preprocessed drug names into a matrix of features suitable for modeling:

Matrix Type	Definition	Key Steps
Bag-of-Words (BoW) [5]	A simple and effective method for feature extraction where the presence of words (tokens) in a document is used as features.	- Vocabulary Creation: Construct a vocabulary of unique words (tokens) across all drug names in the dataset. - Feature Vectorization: Represent each drug name as a vector where each dimension corresponds to a word in the vocabulary, and the value represents the frequency (count) of the word in the drug name.
TF-IDF (Term Frequency-Inverse Document Frequency) [5]	A statistical measure used to evaluate the importance of a word in a document relative to a collection of documents.	- Term Frequency (TF): Calculate how frequently a term (word) appears in a document (drug name). - Inverse Document Frequency (IDF): Calculate the logarithmically scaled inverse fraction of documents that contain the term across the entire dataset. - Feature Vectorization: Multiply TF by IDF to obtain the TF-IDF score for each term in each document, yielding a feature vector.
Word Embeddings (e.g., Word2Vec, GloVe) [15]	Word embeddings are dense vector representations of words in a continuous vector space, where similar words are closer together in the space.	- Pre-trained Models: Utilize pre-trained word embeddings like Word2Vec [14], GloVe [15], or FastText [16] trained on large corpora to obtain vector representations for words. - Feature Representation: Represent each drug name as a vector by averaging or concatenating the embeddings of the words (tokens) within the drug name.

Table 2. Common techniques in Feature Extraction

Here's an example R code for TF-IDF vectorization:

```
# Load necessary libraries
library(text)
library(tm)

# Example drug names
drug_names <- c("Acetaminophen 500mg", "Ibuprofen 200mg", "Amoxicillin
250mg")

# Create a corpus
corpus <- Corpus(VectorSource(drug_names))

# Preprocess the corpus
# Convert to lowercase
corpus <- tm_map(corpus, content_transformer(tolower))
```

```

# Remove punctuation
corpus <- tm_map(corpus, removePunctuation)
# Remove numbers
corpus <- tm_map(corpus, removeNumbers)

# Create a Document-Term Matrix (DTM)
dtm <- DocumentTermMatrix(corpus)

# Convert DTM to a matrix
dtm_matrix <- as.matrix(dtm)

# Calculate TF-IDF
tfidf_transformer <- weightTfIdf(dtm)

# Convert TF-IDF matrix to a data frame
tfidf_matrix <- as.data.frame(as.matrix(tfidf_transformer))

rownames(tfidf_matrix) <- drug_names
colnames(tfidf_matrix) <- Terms(tfidf_transformer)

# View the TF-IDF matrix
print(tfidf_matrix)

```

	acetaminophen	ibuprofen	amoxicillin
Acetaminophen 500mg	1.584963	0.000000	0.000000
Ibuprofen 200mg	0.000000	1.584963	0.000000
Amoxicillin 250mg	0.000000	0.000000	1.584963

Output 2. Output from an example of TF-IDF vectorization

MODEL SELECTION AND TRAINING

1. Model Selection: Select appropriate classification algorithms for your task, considering:
 - Logistic Regression: Simple and interpretable for baseline performance.
 - Support Vector Machines (SVM): Effective for high-dimensional spaces.
 - Decision Trees and Random Forests: Handle non-linearity and interactions between features.
 - Gradient Boosting Machines (GBM): Sequentially improve predictive performance.
 - Neural Networks: Capture complex relationships but require more data and computational resources.
2. Hyperparameter Tuning: Optimize the chosen algorithms by tuning parameters such as regularization strength, learning rate, number of trees (for ensemble methods), and neural network architecture (number of layers, neurons per layer).
3. Model Training:
 - Divide the dataset into training, validation, and test datasets (e.g., 60% training, 20% validation, 20% test dataset) to evaluate model performance.
 - Use techniques like cross-validation to optimize model parameters and prevent overfitting. Cross-Validation by performing k-fold cross-validation to ensure the model's robustness by splitting the data into k subsets and training/validating the model k times.
 - Fit the selected algorithms on the training data, adjusting model parameters to minimize prediction errors.

Here's an example R code for Model Selection and Training:

```

# Split data into training and testing sets (80-20 split)

```

```

set.seed(123) # for reproducibility
trainIndex <- createDataPartition(dtm_df$category, p = 0.8, list = FALSE)
train_data <- dtm_df[trainIndex, ]
test_data <- dtm_df[-trainIndex, ]

# Convert 'category' to factor
train_data$category <- as.factor(train_data$category)
# Convert 'category' to factor
test_data$category <- as.factor(test_data$category)

# Train a Support Vector Machine (SVM) model
model <- svm(category ~ ., data = train_data, kernel = "linear")

# Predict on test data
predictions <- predict(model, newdata = test_data)

# Check levels of predictions and test_data$category
levels(predictions)
levels(test_data$category)

# Ensure levels match and align them if necessary
predictions <- factor(predictions, levels = levels(test_data$category))

```

MODEL EVALUATION

1. Define evaluation metrics such as accuracy, precision, recall, and F1-score: to measure the model's performance in correctly classifying drug names into study-specific classifications. Or create confusion matrix: to visualize the performance of a classification model [5].
2. Assess model performance on the validation set to measure its ability to categorize drug names accurately.
3. Comparison: Compare the performance of different models and feature sets to identify the most effective approach.
 - Validation Performance: Evaluate each model's performance on the validation set and select the best-performing model based on predefined evaluation metrics.
 - Overfitting Prevention: Ensure the selected model generalizes well to new data by monitoring performance on both training and validation sets.

Example R code for model evaluation:

```

# Evaluate confusion matrix and accuracy
confusionMatrix(predictions, test_data$category)

# Accuracy
accuracy <- mean(predictions == test_data$category)
print(paste("Accuracy:", accuracy))

```

DEPLOYMENT AND MONITORING

After we create a robust and accurate classification model that meets study-specific requirements, we can deploy the selected model to classify new drug names into study-specific classifications in production environments [17]. Finally, monitor the model's performance over time and retrain as necessary with new data to maintain accuracy and relevance.

Example R code for applying the model to new data:

```

# Example of new data (drug names)
new_drugs <- c("New Drug X", "Another Drug Y")
# Create a corpus from drug names
new_corpus <- Corpus(VectorSource(new_drugs$drug_name))

```



```

new_corpus <- tm_map(new_corpus, content_transformer(tolower))
new_corpus <- tm_map(new_corpus, removePunctuation)
new_corpus <- tm_map(new_corpus, removeNumbers)
new_corpus <- tm_map(new_corpus, removeWords, stopwords("en"))
new_corpus <- tm_map(new_corpus, stemDocument)
new_corpus <- tm_map(new_corpus, stripWhitespace)

# Convert to a dataframe with preprocessed drug names
new_data <- sapply(new_corpus, as.character)

# Create a document-term matrix for new data
new_dtm <- DocumentTermMatrix(Corpus(VectorSource(new_data)), control =
list(dictionary = Terms(dtm)))

# Convert to a dataframe
new_dtm_df <- as.data.frame(as.matrix(new_dtm))

# Predict using the trained model
new_predictions <- predict(model, newdata = new_dtm_df)

# Create a dataframe with drug names and predicted categories
predicted_df <- data.frame(drugname = new_drugs, predicted_category =
new_predictions)

# Print or inspect the dataframe
print(predicted_df)

```

DISCUSSION

In this case study, we explored the application of natural language processing (NLP) techniques in R for the classification of drug names into study-specific classifications. The goal was to leverage computational linguistics to automate the categorization process, which traditionally relies heavily on manual effort and domain expertise.

PERFORMANCE AND ACCURACY

Our findings indicate that the NLP-based classification models developed in R achieved promising results in accurately categorizing drug names. The use of techniques such as bag-of-words (BoW), TF-IDF, and word embeddings enabled effective representation of drug name semantics, capturing subtle nuances that contribute to classification accuracy [18]. The models demonstrated robust performance across various evaluation metrics, including accuracy, precision, recall, and F1-score, suggesting their suitability for practical applications in pharmaceutical research and healthcare settings [19].

CHALLENGES AND LIMITATIONS

Despite the overall success, several challenges and limitations were encountered during the development and evaluation of the models. One notable challenge was the variability and ambiguity inherent in drug naming conventions, especially with brand names, generics, and variations in international drug nomenclature. This variability sometimes led to misclassifications or ambiguity in model predictions, highlighting the need for further refinement and enhancement of the training datasets and feature engineering techniques.

Another limitation was the computational resources required for training models, particularly when using deep learning approaches or large-scale datasets. While R provides powerful libraries and packages for NLP tasks, such as **tm**, **textclean**, and **keras** [20], optimizing model training and hyperparameter tuning remains computationally intensive and time-consuming in some cases.

IMPLICATIONS AND APPLICATIONS

The successful implementation of NLP for drug name classification in R holds significant implications for pharmaceutical research and clinical practice [21]. By automating the classification process, researchers and healthcare professionals can streamline workflows, enhance data-driven decision-making, and potentially reduce errors associated with manual categorization. This automation is particularly valuable in large-scale studies or clinical trials where rapid and accurate drug classification is crucial for monitoring and analysis.

Furthermore, the modular and scalable nature of the NLP models developed in R allows for adaptation and integration into existing information systems and databases. This flexibility supports the interoperability of healthcare IT systems, enabling seamless data exchange and integration across different platforms and applications [22].

FUTURE DIRECTIONS

Looking ahead, several avenues for future research and development are identified based on the outcomes of this study. Firstly, exploring advanced NLP techniques, such as contextual embeddings (e.g., BERT, GPT), could further enhance the semantic understanding and classification accuracy of drug names [23]. Additionally, incorporating domain-specific knowledge bases or ontologies into the classification pipeline may help mitigate challenges related to naming variations and international drug classifications.

Moreover, investigating the robustness and generalizability of the developed models across diverse datasets and languages remains a critical area of interest. Collaborative efforts with domain experts and stakeholders in pharmaceutical research and healthcare sectors can facilitate the validation and refinement of NLP-based classification frameworks for broader adoption and impact.

CONCLUSION

In conclusion, this paper underscores the potential of NLP in R as a powerful tool for automating and improving the classification of drug names into study-specific classifications. By addressing technical challenges and leveraging emerging advancements in computational linguistics, future research can further enhance the accuracy, efficiency, and applicability of NLP-based solutions in pharmaceutical research and healthcare domains.

We hope the information we've presented here provides you with some ideas or insights for building your own NLP tasks. Happy coding!

REFERENCES

- [1] Chapman, W. W., et al. (2011). Extending the NegEx Lexicon for Multiple Languages. *Studies in Health Technology and Informatics*, 169, 82-86.
- [2] Jensen, P. B., Jensen, L. J., & Brunak, S. (2012). Mining electronic health records: towards better research applications and clinical care. *Nature Reviews Genetics*, 13(6), 395-405.
- [3] Bodenreider, O. (2004). The Unified Medical Language System (UMLS): Integrating Biomedical Terminology. *Nucleic Acids Research*, 32(Suppl_1), D267-D270.
- [4] Holdsworth, Jim. "What is NLP?" Accessed June 6, 2024. IBM. <https://www.ibm.com/topics/natural-language-processing>
- [5] Silge, J., & Robinson, D. (2017). *Text Mining with R: A Tidy Approach*. O'Reilly Media.
- [6] tm <https://www.rdocumentation.org/packages/tm>
- [7] quanteda <https://quanteda.io/>
- [8] text2vec <https://github.com/dselivanov/text2vec>
- [9] NLP <https://cran.r-project.org/web/packages/NLP/index.html>
- [10] keras <https://cran.r-project.org/web/packages/keras/index.html>

- [11] RStudio TensorFlow <https://tensorflow.rstudio.com/>
- [12] tensorflow <https://www.tensorflow.org/>
- [13] spaCy <https://spacy.io/>
- [14] Mikolov, T., et al. (2013). Distributed representations of words and phrases and their compositionality. Advances in Neural Information Processing Systems.
- [15] Pennington, J., Socher, R., & Manning, C. D. (2014). GloVe: Global Vectors for Word Representation. Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP).
- [16] Bojanowski, P., et al. (2017). Enriching Word Vectors with Subword Information. Transactions of the Association for Computational Linguistics (TACL).
- [17] Kuhn, M., & Johnson, K. (2013). Applied Predictive Modeling. Springer.
- [18] "Natural Language Processing with R." Accessed January 17, 2024. GeeksforGeeks. <https://www.geeksforgeeks.org/natural-language-processing-with-r/>
- [19] Udacity Team. "Natural Language Processing With R." Accessed October 8, 2020. Udacity Blog. <https://www.udacity.com/blog/2020/10/natural-language-processing-with-r.html>
- [20] keras <https://keras.io/>
- [21] Jurafsky, D., & Martin, J. H. (2020). Speech and Language Processing (3rd ed.). Pearson Education.
- [22] Gruetzemacher, Ross. "The Power of Natural Language Processing." Accessed April 19, 2022. Harvard Business Review. <https://hbr.org/2022/04/the-power-of-natural-language-processing>
- [23] OpenNLP <https://cran.r-project.org/web/packages/openNLP/index.html>

CONTACT INFORMATION

Your comments and questions are valued and encouraged. Contact the author at:

Feifeng Zhou, Zhenchao Chen
Pfizer (China) Research and Development Co.,Ltd
Feifeng.Zhou@pfizer.com
Zhenchao.Chen@pfizer.com