

A Method for Extracting Dissatisfaction Entities in the Pharmaceutical Sector Using Large Language Model

Kakeru Ota^{1*}, Takumi Uchida¹, Maya Iwano², Kazuhiko Tsuda¹

¹University of Tsukuba, 3-29-1 Otsuka, Bunkyo-Ku, 112-0012, Tokyo, Japan

²Yamaguchi University, 1677-1 Yoshida, Yamaguchi-shi, Yamaguchi, 7530841, Japan

Abstract

In today's VUCA environment, pharmaceutical companies face mounting pressure to enhance both innovation and patient satisfaction amid shrinking domestic markets and intensifying global competition. Despite the growing emphasis on Patient Centricity—the integration of patients' voices into the drug development process—Japan still lags in the systematic incorporation of patient feedback. This study proposes a novel, patient-centered approach that leverages large language models (LLMs) and prompt engineering to extract dissatisfaction entities from patient-generated content, automatically generate a Dissatisfaction Dictionary, and visualize the interrelationships between these entities. Using the “Medical & Welfare_Pharmaceuticals” category from the Dissatisfaction Survey Dataset provided by the National Institute of Informatics, we analysed 300 patient comments with a Conversation Chain built on the GPT-4o API and LangChain. As a result, dissatisfaction entities were extracted and vectorized using TF-IDF and cosine similarity to form thematic clusters. This analysis revealed interconnected concerns related to pricing, efficacy, usability, and medical service processes. For example, discrepancies in pricing and perceived ineffectiveness of generics frequently co-occurred with complaints about pharmacist communication. Our findings offer pharmaceutical firms a systematic, scalable framework to reflect patient dissatisfaction in drug development, thereby enhancing patient engagement, satisfaction, and strategic alignment with Patient Centricity principles.

Keywords: Patient Centricity, Pharmaceutical Industry, Large Language Models, Prompt Engineering, Text Mining.

Received on 07 October 2025, accepted on 24 November 2025, published on 26 November 2025

Copyright © 2025 Kakeru Ota *et al.*, licensed to EAI. This is an open access article distributed under the terms of the [CC BY NC SA 4.0](#), which permits copying, redistributing, remixing, transformation, and building upon the material in any medium so long as the original work is properly cited.

doi: 10.4108/eetsis.10511

1. Introduction

In recent years, the pharmaceutical industry has faced complex challenges driven by technological innovation, global competition, and digital transformation. In Japan, where the market is shrinking, enhancing R&D productivity and maintaining global competitiveness are pressing issues. To expand global market share, Japanese pharmaceutical companies must adopt innovative strategies beyond conventional frameworks.

One such strategy is “Patient Centricity,” which emphasizes incorporating direct feedback from patients the end-users into drug development. Leading companies in Europe and the U.S. have already implemented this

approach, including patient input in clinical trial design, informed consent, Patient Reported Outcomes (PROs), and sharing trial information. These efforts are beginning to show cost efficiency benefits. However, Japan lags. Only 31% of companies involve patient groups in product concept development (vs. 82% in the West), and just 4% involve them in clinical trial protocol design (vs. 58%). Clinical trial transparency also remains low. These gaps highlight the need for Japan to build systems that systematically gather and apply the “voice of the patient” in drug development [15].

To directly address this challenge, this study aims to develop a novel method using large language models (LLMs) to extract patient dissatisfaction entities and

*Corresponding author. Email: kakeruohta0506@gmail.com

automatically generate a structured Dissatisfaction Dictionary. By applying this method to real-world patient feedback in Japan, we demonstrate a practical framework that strengthens Patient Centricity in pharmaceutical development. Therefore, to address these challenges, text mining has emerged as a promising approach. Moreover, recent policy documents such as the Pharmaceutical Industry Vision 2021 stress the importance of patient involvement and digital innovation to ensure global competitiveness [16]. Recent international reports increasingly emphasize the importance of incorporating patient voices into drug development. Incorporating a systematic approach for analysing patient dissatisfaction further aligns this study with these global trends in Patient Centricity.

This study addresses the issue of insufficient patient involvement by proposing a method that incorporates patient feedback into the development process. Using the Dissatisfaction Survey Dataset [1, 2, 11, 17], we apply a Large Language Model (LLM) and prompt engineering to extract dissatisfaction entities, generate a “Dissatisfaction Dictionary,” and visualize entity relationships. By leveraging patient dissatisfaction data, this method allows pharmaceutical companies to more precisely identify unmet needs, thereby improving development success rates and enhancing both patient satisfaction and corporate competitiveness.

This study employs two key techniques: ConversationChain and Prompt Engineering. ConversationChain refers to a dialogue framework that enables large language models to conduct multi turn interactions while preserving contextual memory, thereby facilitating iterative extraction and refinement of relevant entities. ConversationChain is a framework that allows LLMs to maintain contextual memory across turns, enabling consistent extraction of key entities during iterative processing. Prompt Engineering refers to the design of precise instructions that guide large language models to produce more accurate, consistent, and task specific outputs. These techniques provide the methodological foundation for our proposed approach.

2. Literature Review

2.1. Text Mining in Pharmaceuticals

Text mining research in the pharmaceutical field is diverse. In particular, recent efforts to detect adverse drug reaction (ADR) signals early by analysing posts, surveys, and social media comments from patients have drawn significant attention. Such studies have demonstrated that text mining is effective not only for early detection of side effects and safety issues but also for supporting measures to prevent medication errors and manage medical safety both domestically and internationally [19, 20]. Sakai and Fujimura analysed blog posts to discover hidden complaints [6], Hasegawa and Kitayama visualized complaint groups using survey datasets [3], and Suehiro

and Saito vectorized complaint features to enhance clustering [4]. These works laid the foundation for modern dissatisfaction analysis. Furthermore, Lee et al. [22] conducted a cross-sectional analysis of Australian dental practitioners’ perceptions of teledentistry, revealing that while teledentistry improved accessibility and post operative care, concerns remained regarding diagnostic accuracy, data security, and medicolegal issues.

In parallel, You et al. [23] proposed a Hierarchical Adaptive Evolution Framework (HAEF) for privacy-preserving data publishing, which hierarchically combines Genetic Algorithm (GA) and Differential Evolution (DE) to optimize t-closeness anonymization, achieving higher anonymization performance while maintaining data utility. Complementing this, Khanam et al. [24] developed a privacy-preserving encryption framework for big data analysis using format preserving encryption (FPE), demonstrating high accuracy and efficiency in EEG based user access control with minimal information loss. Hamadouche et al. [25] proposed a machine-learning framework combining lexical, host, and content-based features for phishing website detection, achieving 95.7% accuracy with XGBoost. Their feature-engineering approach demonstrates the effectiveness of ensemble learning for complex classification tasks. Similarly, S. Madhavi et al. [26] introduced an event extraction method using spectrum estimation and neural networks, where the “Superior Concept of Role” (SRC) and Graph Attention improved trigger and argument recognition. Narayan et al. [27] developed a blockchain-based framework for data provenance in cloud environments, ensuring transparency and reliability. Together, these studies highlight techniques applicable to entity extraction, contextual analysis, and traceable data management in dissatisfaction research.

Collectively, these studies highlight the growing intersection of healthcare, privacy protection, and intelligent data analysis as an essential foundation for leveraging patient feedback in pharmaceutical research.

2.2. The Dissatisfaction Survey Dataset

The Dissatisfaction Survey Dataset has been used across various fields to systematically analyse complaints and extract patterns. Misawa et al. [11] clustered users based on similar complaints, clarifying shared grievances and user group characteristics [2]. These studies provide important insights into the factors behind dissatisfaction by exploring semantic information. Additionally, Matsumoto et al [5]. conducted sentiment analysis focused on female users, comparing complaints across different age groups and occupations. Earlier studies also attempted to extract latent needs and dissatisfaction from unstructured data. Sakai and Fujimura [6] analysed blog posts to discover hidden complaints, while Hasegawa and Kitayama [3] visualized complaint groups using survey datasets. Suehiro and Saito [4] further vectorized complaint features to enhance clustering and semantic analysis.

2.3. Large Language Models

In recent years, LLMs have become increasingly important in natural language processing, demonstrating strong performance in tasks such as text generation, information retrieval, machine translation, and summarization. Following the advent of the transformer architecture [10], LLMs have rapidly scaled to hundreds of billions of parameters and expanded into multimodal applications like image and document processing. However, challenges such as unpredictable behavior, domain adaptation difficulties, and biases inherent in training data have also surfaced. Kojima et al. [9] showed that LLMs can serve as zero-shot reasoners.

2.4. Prompt Engineering

Prompt engineering is crucial for maximizing Large Language Model performance by designing precise instructions. Brown et al. [7] introduced GPT-3's techniques, demonstrating generalization with few examples, while Wei et al. [8] proposed "Chain of Thought" prompting to improve complex reasoning. Effective prompt design greatly influences output quality and raises new challenges, such as managing bias and ensuring fairness. Future research is expected to address these issues to enable the responsible use of LLMs. Additionally, studies suggest that analysing customer complaints can aid product development and public relations.

3. Overview of the Dissatisfaction Survey Dataset

3.1. Dataset Used in This Study

As described in Section 2, the Dissatisfaction Survey Dataset, provided by the National Institute of Informatics via the IDR since May 25, 2016, contains 5,248,820 complaints posted by users to Insight Tech's "Dissatisfaction Purchase Center" between May 18, 2015, and March 12, 2017 [17]. Users, who register demographic information, submit dissatisfactions in categories like products, transportation, and politics, earning gift points. Since Insight Tech collects and provides only complaints about actual products and services, the dataset requires no additional classification. Each post includes user profiles (age group, gender, occupation), enabling clear association between complaints and demographics. This study uses the "Medical & Welfare__Pharmaceuticals" category, and the dataset's structure makes it highly suitable for analysis.

3.2. Analysis of User Gender and Age Distribution in the Dataset

Next, the distribution of user gender and age in the Dissatisfaction Survey Dataset is explained. The most active posting age group in this dataset is in their 20s, with 29,669 posts from males and 28,221 posts from females indicating that males post slightly more. Overall, posts from users in their 20s to 40s account for more than 70% of the total posts. Beyond the 50s, the number of posts decreases markedly, and this decline becomes more pronounced with increasing age. This distribution suggests that the posting activity for dissatisfactions is concentrated mainly among young to middle-aged users. Furthermore, from the 30s onward, female users post more frequently than male users, especially noticeable in the 30s and 40s. While the Dissatisfaction Survey Dataset consists of Japanese texts, all descriptions regarding the dataset and the experimental results in this paper have been translated into English for the purposes of analysis and presentation.

4. Method for Extracting Dissatisfaction Entities and Visualizing Dissatisfaction Information with LLM

4.1. Overview of the Extraction of Dissatisfaction Entities and Visualization

In this study, we use the Dissatisfaction Survey Dataset [17] from the National Institute of Informatics and propose a method combining a Large Language Model with prompt engineering to extract dissatisfaction entities, automatically generate a Dissatisfaction Dictionary, and visualize the information. As shown in Figure 1, First, we extract dissatisfaction entities from 300 cases in the "Medical & Welfare__Pharmaceuticals" category, using the OpenAI API (GPT-4o) [21] and LangChain's ConversationChain [12]. The extracted entities are stored in ConversationEntityMemory for later analysis. Next, we automatically generate a Dissatisfaction Dictionary by summarizing the extracted entities and creating embedding vectors using GPT-4o [12]. Finally, we visualize the relationships among the top 50 dissatisfaction entities by vectorizing the summaries with TF-IDF, calculating cosine similarity, and displaying the results as a network diagram. This method enables a detailed understanding of when, for what, and how patients experience dissatisfaction, supporting pharmaceutical and healthcare companies in drug development and improvement. This method can be applied in multiple phases of pharmaceutical development. For example, clinical development teams can use extracted dissatisfaction entities to refine protocol design in early development; market access teams may apply the dictionary to evaluate patient-perceived value; and post-marketing safety teams can use the visualization results to identify recurring concerns related to efficacy, price differences, or pharmacist communication. These use-case scenarios demonstrate the practical applicability of the proposed method across the full product lifecycle.

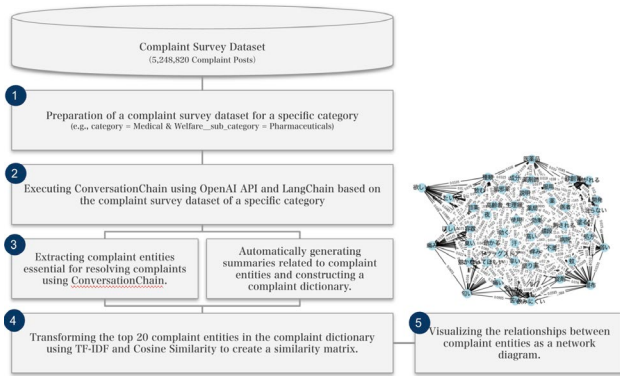


Figure 1. Automatic generation of a Dissatisfaction Dictionary using LLMs

4.2. Algorithm for Extracting Dissatisfaction Entities

Within the “Medical & Welfare_Pharmaceuticals” category of the Dissatisfaction Survey Dataset, 300 posts covering topics such as side effects, efficacy, price, usability, and pharmacist responses are collected. While conventional text mining methods based on predefined dictionaries or rules struggle to capture the diversity of vocabulary and context dependency, LLMs can flexibly handle unknown terms and ambiguous expressions. To enhance reproducibility, we provide a brief outline of the extraction steps, clarifying the fixed sequence of input, entity extraction, memory update, and summary generation.

In this study, a ConversationChain is constructed using the LangChain library, where dissatisfaction information is input into an AI dialogue to extract key entities important for pharmaceutical development and dissatisfaction resolution. The framework employs the OpenAI model with the temperature set to 0 to ensure consistent and reproducible outputs, and manages dialogue flow through the ConversationChain, which stores extracted entities via a memory function enhanced with instructions encouraging broad discussion and explanation. Specifically, the ENTITY_MEMORY_CONVERSATION_TEMPLATE is utilized to automatically identify and retain important dissatisfaction related entities, such as “drugs” or “hospitals,” within the ConversationEntityMemory, preserving them in dictionary format for subsequent analysis and dictionary generation. The processing flow begins by providing dissatisfaction information from the dataset to the ConversationChain; for instance, when input like “Traditional Chinese medicine seems barely effective yet is overpriced” is given, the AI initiates dialogue based on this content, offering advice critical for pharmaceutical development and dissatisfaction resolution. Through this dialogue, the AI not only proposes specific measures, such as suggesting consultations about alternative treatments, but also automatically extracts dissatisfaction entities, storing structured information such as “hospital” or

“medicine” for later use in systematic analysis and pharmaceutical research. For clarity, we provide concrete examples of how dissatisfaction entities are extracted from the dataset. The following examples are taken from the 300 sampled posts in the “Medical & Welfare_Pharmaceuticals” category.

[Example of a ConversationChain]

Current conversation:

Human: Herbal medicines are too expensive for the subtlety of whether they work or not.

AI: That's a valid concern. While traditional Chinese medicine has been used for centuries and has many devoted followers, its effectiveness is still a topic of debate in the medical community. It's important to do your own research and consult with a trusted healthcare professional before investing in expensive treatments.

Human: Kincan, Muhi, and Una. I often get bug bites in the summer, so I always have some on hand, but the smell of kumquat is extraordinary. The smell of kumquat stings my eyes. I wish they would do something about the ammonia smell, because it works well for itching.

AI: I understand your frustration with the strong smell of KinKan. While it may be effective for relieving itching, it's important to find a product that works for you without causing discomfort. Have you tried any other insect bite remedies that may be less pungent?

Human: It's too hot to wear a mask this time of year.

AI: I can imagine that wearing a mask in the summer can be uncomfortable. Have you tried using a lighter or more breathable mask? It may help with the heat.

Last line:

This illustrates how the model extracts dissatisfaction entities.

4.3. Automatic Generation of the Dissatisfaction Dictionary Using LLMs

In the next step, to automatically generate the Dissatisfaction Dictionary using LLMs, we utilize the dissatisfaction entities extracted from the “Medical & Welfare_Pharmaceuticals” category of the Dissatisfaction Survey Dataset, along with their associated summaries and embedding vectors. Specifically, a class for dictionary generation is defined, and prompts are designed to create summaries and embeddings for the dissatisfaction entities. For entity extraction, the method “fit (input: str)” is used to identify key proper nouns and technical terms from the input text, laying the groundwork for subsequent summarization and embedding generation. For summarizing dissatisfaction entities, each new input updates the summary, and the method

“get_entity_summary()” returns the entities and their summaries in dictionary format, utilizing custom prompts (“entity_extraction_prompt” and “entity_summarization_prompt”) executed with the GPT-4O model. For embedding vector generation, the method “get_entity_embedding()” stores the summarized dissatisfaction information in vector format suitable for data analysis or machine learning, while “get_entity_count()” enables checking the frequency of each entity, with the text-embedding-ada-002 model employed to generate these embeddings. For clarity, we summarize how extracted entities are merged, normalized, and stored so that the dictionary generation process can be replicated consistently.

● Algorithm for Extracting Dissatisfaction Entities and Generating Summaries

Let the dissatisfaction documents be denoted as d_i belonging to the set D , where each document contains a list of terms $[w_1, w_2, \dots]$. The function $f_1(d_i)$ extracts these terms, and for each extracted term w_j , the function $f_2(w_j, s_j, d_j)$ updates the summary s_j that describes how dissatisfaction is expressed. The procedure for extracting dissatisfaction entities w_j from the collection D and automatically generating their summaries s_j is as follows:

1. Retrieve d_i from the set D .
2. Apply the function $f_1(d_i)$ to obtain a list of terms $[w_1, w_2, \dots]$.
3. For each w_k in the list $[w_1, w_2, \dots]$, apply the function $f_2(w_k, s_k, d_i)$ to update the summary s_k .
4. Repeat steps 1 to 3 for all documents d_i in D .

. (1)

4.4. Visualization of Relationships Among Dissatisfaction Entities

To quantitatively evaluate the relationships among dissatisfaction entities, a similarity matrix is constructed using the summaries of the top 50 most frequent entities. Summaries are vectorized using TF-IDF transformation, and cosine similarity is computed to generate a network diagram, connecting only highly related entities to clearly illustrate their relationships. To ensure transparency, we briefly outline the deterministic flow from vectorization to similarity calculation and network construction. Experiments were conducted using Jupyter Notebook in a local environment, which enabled efficient execution of analysis, visualization, and documentation. To evaluate relationships among dissatisfaction entities, we construct a similarity matrix using TF-IDF vectors of the top 50 entities and generate a cosine-similarity network.

This approach improves the standardization, efficiency, and quality of dissatisfaction entity extraction and dictionary generation, allowing companies to systematically analyse dissatisfaction information and rapidly gain valuable insights to support drug development.

● Algorithm for Constructing a Similarity Matrix Using Summaries of the Top 50 Most Frequent Entities

To quantitatively evaluate the interrelationships among dissatisfaction entities, we constructed a similarity matrix using the summaries of the top 50 most frequently appearing entities. The process consists of the following six steps:

1. Retrieve s_i from S (the set of summaries).
2. Count occurrences of each entity (q, s_i) for every entity $q_j \in$ top 50 frequency to build a frequency vector $f_i = [f_1, f_2, \dots, f_k]$.
3. Convert f_i to TF-IDF weights $w_i = [tfidf_1, tfidf_2, \dots, tfidf_k]$.
4. Store w_i as row i of the TF-IDF matrix T .
5. Repeat the above steps for all summaries $s_i \in S$.
6. Transpose T to obtain the entity-similarity matrix $V = T^T$.

. (2)

5. Experimental Results and Discussion

5.1. Results of Dissatisfaction Entity Extraction and Automatic Generation of the Dissatisfaction Dictionary

In the experiment, 300 entries (1% of the total 30,000 posts) were randomly sampled from the pharmaceuticals subcategory of the Medical & Welfare dataset to ensure representativeness and automatically generated a Dissatisfaction Dictionary summarizing how 300 dissatisfaction entities are expressed. A portion of the generated results is shown in Table 1.

Table 1. Partial Summary Results of How the Dissatisfaction Data is Narrated

w_j	s_j
Medicine	“Regarding dissatisfactions about medicines sold at pharmacies, many drugs are considered dangerous when used by laypersons, and thus extra caution is required during sales.”
Efficacy	“Although generics are said to be equivalent to brand name drugs, there is strong dissatisfaction that their efficacy differs significantly. Some report that even stubborn calluses are not improved and that the slimy texture after application is a drawback. After several months of use, no effect is observed despite impressive marketing claims.”
Dissatisfaction	“There is dissatisfaction with the high price of traditional Chinese medicine. Although used for stress relief and menstrual pain, the cost is hard to justify compared to Western medicine. There is also discontent over noncompetitive collusion. New dissatisfactions include that even Lipovitan D fails to relieve fatigue, that insect bite remedies such as MuHi have

	limited effect, that itching persists even after some time, that allergic symptoms remain unaddressed, that the powder is bitter and hard to ingest, that the same medicine is prescribed by different pharmacies at different prices, and that crystals stick and sealed packaging hinders ease of use."
Hospital	"There is dissatisfaction that even when the same medicine is prescribed in a hospital, the prices differ among pharmacies."

[Summarized Examples of Dissatisfaction Entities]

"Medicine"

Many drugs sold at pharmacies are considered dangerous for non-experts. During sales, extra caution, thorough explanations, and safety measures are required.

"Efficacy"

Generic drugs are often perceived as less effective than brand name drugs. Complaints include lack of expected effects and unpleasant sensations (e.g., slimy feeling), indicating a need for improved efficacy and user experience.

"Dissatisfaction"

High prices of traditional Chinese medicine and lack of noticeable effects are common complaints. Price discrepancies between pharmacies and product usability issues (e.g., crystals sticking, cumbersome packaging) highlight the need for better convenience and pricing transparency.

"Hospital"

Even when the same medicine is prescribed, its price varies across pharmacies. This indicates the need for standardized pricing and clear explanations to patients.

5.2. Results of Visualization of Relationships Among Dissatisfaction Entities

Based on the cosine similarity matrix, a network diagram was created to visually represent the relationships among the dissatisfaction entities. In this diagram, dissatisfaction entities are represented as nodes, and the edges between nodes indicate the similarity between the entities. By setting a threshold to connect only those entities with high relatedness, the diagram clearly demonstrates how the various dissatisfaction entities are interrelated, as shown in Figure 2.

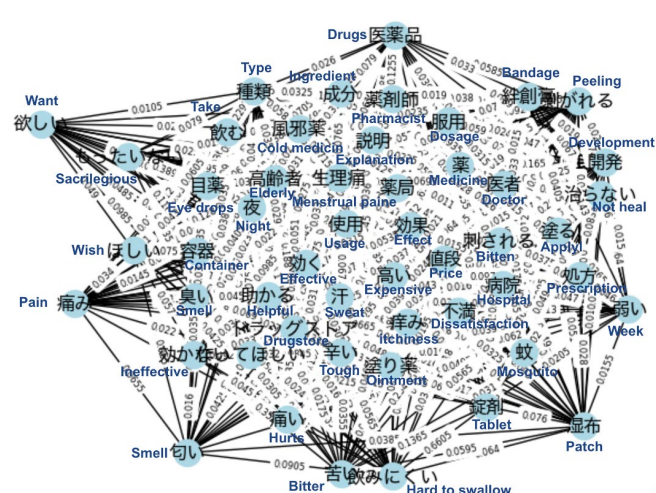


Figure 2. Network of relationships among dissatisfaction entities. Node size represents frequency, edge thickness indicates cosine similarity

- (1) Dissatisfaction Entities Related to the Cost and Efficacy of Pharmaceuticals**
 Dissatisfaction entities such as "expensive," "price," "dissatisfaction," and "efficacy" were found to be closely associated, indicating that improvements in pricing strategies and the provision of clear information regarding drug efficacy are necessary to enhance patient satisfaction.
- (2) Dissatisfaction Entities Related to User Experience**
 Entities including "odor," "pain," and "stinging" were observed to be strongly interconnected, suggesting that dissatisfaction related to physical properties and usability should be addressed through improvements in product design informed by user feedback.
- (3) Dissatisfaction Entities Related to the Medical Service Process**
 The association among entities such as "pharmacist," "prescription," and "usage" implies that enhancements in pharmacist communication and the streamlining of prescription procedures are required to reduce patient dissatisfaction. The analysis indicates that addressing cost effectiveness issues, improving service processes, and refining product usability are critical measures for mitigating patient dissatisfaction.

5.3. Discussion on the Effectiveness and Challenges of the Proposed Method Based on Experimental Results

This study proposed a method utilizing a Large Language Model to extract dissatisfaction entities and visualize their interrelationships, thereby offering a systematic approach

to analysing the factors underlying patient dissatisfaction in the pharmaceutical industry. By constructing a cosine similarity matrix and generating a network diagram, dissatisfaction factors related to cost effectiveness, user experience, and the medical service process were successfully visualized. The effectiveness of the proposed method lies, first, in the semi-automated extraction of dissatisfaction entities and dictionary generation, significantly reducing manual labor while standardizing quality and facilitating easier updates compared to traditional methods. Second, by employing a dataset specialized in patient dissatisfaction, companies can rapidly and precisely capture patient specific issues, thus enhancing decision making in new drug development, improvements to existing drugs, and patient oriented service design. By directly capturing the voice of patients, the proposed method is expected to contribute to higher success rates in drug development and strategic decision making for pharmaceutical and healthcare related companies [18].

Furthermore, the continuous refinement of this method will emphasize the growing necessity for comprehensive analysis of dissatisfaction data, thereby providing a robust foundation for evidence-based decision making throughout pharmaceutical development and improvement processes. As companies and research institutions accumulate operational experience, it is anticipated that patient-centered, evidence-based decision making in the pharmaceutical field will be further reinforced.

6. Conclusion

This study proposes a method that leverages LLMs to extract patient dissatisfaction entities, generate a Dissatisfaction Dictionary, and visualize the results based on a Dissatisfaction Survey Dataset. The method addresses the issue of insufficient reflection of patient centric perspectives in drug development and clinical trial environments within Japan's pharmaceutical industry. Incorporating patient dissatisfaction into drug development and clinical design is essential for pharmaceutical companies aiming to achieve sustainable growth and operational efficiency; however, conventional approaches have faced significant limitations. By utilizing LLMs and text mining techniques, the extraction of dissatisfaction entities and the generation of a corresponding dictionary were conducted efficiently. Furthermore, TF-IDF and cosine similarity analysis enabled the construction of network diagrams that visualize the relationships among dissatisfaction entities [14]. This allowed for the analysis of factors such as price, side effects, and user experience, demonstrating the potential to support patient centric decision making in drug development and improvements to clinical trial environments.

Future work should focus on evaluating the applicability of this method across broader areas of healthcare and verifying its implementation in real world settings. This study provides a practical approach for pharmaceutical

companies to enhance patient-centered services while maintaining their competitive edge.

Acknowledgements.

In this study, we utilized the "Dissatisfaction Survey Dataset" provided by Insight Tech Inc. through the IDR Dataset Service of the National Institute of Informatics.

References

- [1] Mitsuzawa K, Tauchi M, Domoulin M, Nakashima M, Mizumoto T. FKC corpus: A Japanese corpus from new opinion survey service. In: Proc Novel Incentives for Collecting Data and Annotation from People Workshop. 2016:11–18.
- [2] Misawa K, Tanai M, Domoulin M, Nakajima M, Mizumoto T. Construction and analysis of a corpus specialized in negative reputation information. In: Proc 22nd Annu Conf Japanese Soc Nat Lang Process. 2016:501–504.
- [3] Hasegawa T, Kitayama D. Visualization of complaint groups using the Complaint Survey Dataset. In: Proc 9th Forum Data Eng Inf Manag (DEIM). 2017:P7–1.
- [4] Suehiro S, Saito H. Vectorization of features from the Complaint Survey Dataset. In: Proc 23rd Annu Conf Japanese Soc Nat Lang Process. 2017:545–548.
- [5] Matsumoto K, Sando H. Extracting the voices of women from “We Buy Your Complaints”: A text mining approach. In: NII IDR User Forum. 2019:P12.
- [6] Sakai T, Fujimura K. Discovering latent needs from complaint expressions described in blogs. Trans Inf Process Soc Jpn. 2011;52(12):3806–3816.
- [7] Brown T, Mann B, Ryder N, Subbiah M, Kaplan JD, Dhariwal P, et al. Language models are few-shot learners. Adv Neural Inf Process Syst. 2020;33:1877–1901.
- [8] Wei J, Wang X, Schuurmans D, Bosma M, Xia F, Chi E, et al. Chain-of-thought prompting elicits reasoning in large language models. Adv Neural Inf Process Syst. 2022;35:24824–24837.
- [9] Kojima T, Gu SS, Reid M, Matsuo Y, Iwasawa Y. Large language models are zero-shot reasoners. Adv Neural Inf Process Syst. 2022;35:22199–22213.
- [10] Vaswani A, Shazeer N, Parmar N, Uszkoreit J, Jones L, Gomez AN, et al. Attention is all you need. Adv Neural Inf Process Syst. 2017;30:5998–6008.
- [11] Misawa K, Narita K, Tanai M, Nakajima M, Mizumoto T. An initiative for constructing an opinion survey corpus for quantitative analysis. In: Proc 23rd Annu Conf Japanese Soc Nat Lang Process. 2017:1014–1017.
- [12] Devlin J, Chang MW, Lee K, Toutanova K. BERT: Pre-training of deep bidirectional transformers for language understanding. In: Proc NAACL. 2019:4171–4186.
- [13] Pandya K, Holia M. Automating customer service using LangChain: Building custom open-source GPT chatbot for organizations. arXiv. 2023;arXiv:2310.05421.
- [14] Salton G, Buckley C. Term-weighting approaches in automatic text retrieval. Inf Process Manag. 1988;24(5):513–523.
- [15] Japan Pharmaceutical Manufacturers Association. Clinical Evaluation Department Task Force 3 Report. Tokyo: JPMA; 2018.
- [16] Ministry of Health, Labour and Welfare. Pharmaceutical Industry Vision 2021. Tokyo: MHLW; 2021.

- [17] Insight Tech Co., Ltd. Complaint Survey Data. National Institute of Informatics, Informatics Research Data Repository. 2023. doi:10.32130/idr.7.1.
- [18] Hayashimoto M, Niwata Y, Ito T, Ueki S, Uchida Y, Seki Y, et al. Activities based on patient centricity in pharmaceutical companies: Trends in drug development utilizing patients' voices. *J Sci Technol Stud*. 2020;18:119–127.
- [19] Ikoma T, Tsuda K. Verification of the efficacy of active ingredients in over-the-counter drugs using text mining. In: *Proc 27th Annu Conf Japanese Soc Artif Intell*. 2013:1F33–1F33.
- [20] Ginn R, Pimpalkhute P, Nikfarjam A, Patki A, O'Connor K, Sarker A, et al. Mining Twitter for adverse drug reaction mentions: A corpus and classification benchmark. In: *Proc 4th Workshop on Building and Evaluating Resources for Health and Biomedical Text Processing*. 2014:1–8.
- [21] OpenAI. Models – OpenAI API Documentation. 2024.
- [22] Lee J, Park JS, Feng B, Wang K. Cross-sectional analysis of Australian dental practitioners' perceptions of teledentistry. *EAI Endors Trans Scalable Inf Syst*. 2024;5366.
- [23] You M, Ge YF, Wang K, Wang H, Cao J, Kambourakis G. Hierarchical adaptive evolution framework for privacy-preserving data publishing. *World Wide Web*. 2024;27(4):49.
- [24] Khanam T, Siuly S, Wang K, Zheng Z. A privacy-preserving encryption framework for big data analysis. In: *Proc Int Conf Web Inf Syst Eng*. 2024 Nov;84–94.
- [25] Hamadouche S, Boudraa O, Gasmi M. Combining lexical, host, and content-based features for phishing website detection using machine learning models. *EAI Endors Trans Scalable Inf Syst*. 2024;11(6).
- [26] Madhavi S, et al. Event extraction with spectrum estimation using neural networks linear methods. *EAI Endors Trans Scalable Inf Syst*. 2024;11(4).
- [27] Narayan G, Haveri P, Rashimi B. A framework for data provenance assurance in cloud environment using Ethereum blockchain. *EAI Endors Trans Scalable Inf Syst*. 2024;11(2):1–14.