

Combining generative AI and SenticNet for aspect-based sentiment analysis of restaurant reviews on Google Maps

Minh-Phuong Han *

Faculty of Economic Information Systems and E-Commerce, Thuong Mai University, Hanoi, Vietnam.

International Journal of Science and Research Archive, 2026, 19(01), 1026-1030

Publication history: Received on 15 March 2026; revised on 22 April 2026; accepted on 24 April 2026

Article DOI: <https://doi.org/10.30574/ijrsra.2026.19.1.0852>

Abstract

This paper proposes a pipeline for aspect-based sentiment analysis (ABSA) that combines Generative AI (Gemini API) and SenticNet to measure customer satisfaction from restaurant reviews on Google Maps. The pipeline operates in three stages: (1) automatic translation of Vietnamese reviews into English using Gemini 2.5 Flash; (2) extraction of aspect-sentiment pairs through few-shot prompting; and (3) polarity scoring using a SenticNet lexicon of 135 adjectives and 62 adverbs. Experiments on 688 reviews collected from 80 buffet and hot-pot restaurant branches across Vietnam yielded 1,432 sentiment words distributed over six aspects. Food received the most mentions (570 words) and the highest satisfaction score (3.86/5.0), whereas Service scored lowest (3.42/5.0). The overall positive-sentiment ratio was 73.0%. Pearson correlation between the SenticNet-derived scores and original Google Maps star ratings reached $r = 0.676$ ($p < 0.001$), confirming the validity of the approach.

Keywords: Aspect-based sentiment analysis; Generative AI; SenticNet; Restaurant reviews; Google Maps; Customer satisfaction; Natural language processing; Few-shot prompting

1. Introduction

Online review platforms such as Google Maps have become a primary channel through which consumers share dining experiences and influence the purchasing decisions of others. With over one billion monthly active users and more than 200 million listed businesses, Google Maps generates a vast volume of unstructured textual feedback. Manually analysing thousands of reviews to extract actionable insights is infeasible in terms of both time and cost, creating a pressing need for automated sentiment analysis solutions.

Aspect-based sentiment analysis (ABSA) addresses this need by identifying the sentiment expressed toward specific service attributes—such as food quality, staff attitude, price, or ambiance—within the same review [1]. Shin et al. [2] analysed over 5,400 Google Maps restaurant reviews using TF-IDF and Random Forest to classify sentiment across four criteria. Mathayomchan and Taecharungroj [3] applied the VADER algorithm to nearly one million Google Maps reviews in the United Kingdom and found significant correlations between aspect-level sentiments and overall star ratings.

The emergence of large language models (LLMs) has opened new avenues for ABSA. Krugmann and Hartmann [4] benchmarked GPT-3.5, GPT-4, and Llama 2 against fine-tuned transfer-learning models on over 3,900 documents and showed that LLMs can match or surpass traditional methods in zero-shot sentiment classification. Zhang et al. [5] evaluated LLMs on 13 sentiment tasks across 26 datasets, reporting that LLMs excel in few-shot settings yet struggle with complex structured sentiment phenomena. In parallel, SenticNet [6]—a neuro-symbolic AI framework—provides concept-level sentiment analysis grounded in commonsense knowledge, achieving higher accuracy than ChatGPT, RoBERTa, word2vec, and bag-of-words while remaining fully interpretable.

* Corresponding author: Minh-Phuong Han

Despite these advances, no prior study has combined the flexible extraction capabilities of Generative AI with the transparent, knowledge-grounded scoring of SenticNet for ABSA of restaurant reviews. Moreover, existing work predominantly targets English-language data, leaving multilingual contexts—especially Vietnamese—underexplored.

This paper makes three contributions: (1) it designs a GenAI–SenticNet ABSA pipeline that handles multilingual reviews through machine translation; (2) it presents an empirical evaluation on 688 Vietnamese restaurant reviews from Google Maps; and (3) it validates the pipeline by correlating its output with original star ratings ($r = 0.676$).

2. Literature Survey

2.1. Aspect-Based Sentiment Analysis in the Restaurant Domain

ABSA aims to identify fine-grained sentiments toward individual aspects mentioned in a text [1]. In the restaurant domain, Shin et al. [2] vectorised 5,427 Google Maps reviews using TF-IDF and built a sentiment dictionary across four categories: food, price, service, and atmosphere. Mathayomchan and Taecharungroj [3] used VADER to process 935,386 Google Maps reviews in London, Birmingham, and Manchester, demonstrating that food, service, atmosphere, and value significantly predict overall star ratings via logistic regression. Carrasco and Dias [7] compared BART, DeBERTa, and ChatGPT on 1,000 restaurant reviews and found that ChatGPT achieved the highest F1-score of 0.77, approaching human-level performance.

2.2. Large Language Models for Sentiment Analysis

Krugmann and Hartmann [4] conducted a comprehensive benchmark of GPT-3.5, GPT-4, and Llama 2 against established transfer-learning models (ULMFiT, BERT, XLNet, RoBERTa, SiEBERT) using over 3,900 text documents from 20 datasets. In zero-shot settings, LLMs matched or exceeded fine-tuned baselines. Zhang et al. [5] extended this evaluation to 13 tasks on 26 datasets and reported that LLMs perform well on simple tasks (binary, three-class) but lag behind on complex structured sentiment tasks; however, they significantly outperform small models in few-shot learning. Sun et al. [8] proposed an LLM negotiation strategy between a generator and a discriminator to improve sentiment accuracy. Niimi [9] applied majority voting across multiple local LLM runs for restaurant review analysis, achieving over 80% accuracy. Ding et al. [10] demonstrated that GPT-3 can serve as a reliable data annotator for NLP tasks. Joseph [13] highlighted the growing role of NLP-based sentiment analysis for social media platforms. Magdaleno et al. [14] applied a GPT-based approach to sentiment analysis and rating prediction in the food service domain.

2.3. SenticNet

SenticNet [6] is a neuro-symbolic framework that fuses commonsense knowledge graphs with neural attention networks for concept-level sentiment analysis. Each concept is assigned a polarity score in $[-1, +1]$ derived from the Hourglass of Emotions model [11], which decomposes affect into four dimensions: Pleasantness, Attention, Sensitivity, and Aptitude. The polarity formula is: $p = (\text{Pleasantness} + |\text{Attention}| - |\text{Sensitivity}| + \text{Aptitude}) / 9$. SenticNet 8 [6] outperformed ChatGPT, RoBERTa, word2vec, and bag-of-words across sentiment analysis, personality prediction, and suicidal ideation detection while maintaining full interpretability. Unlike deep-learning black boxes, SenticNet provides transparent, explainable polarity assignments traceable to commonsense knowledge links [12].

3. Materials and Methods

3.1. Data Collection

Reviews were scraped from Google Maps using the google-search-scraper tool in three batches during October 2024–January 2025. After merging and deduplication, the final dataset comprised 688 reviews from 80 branches of popular buffet and hot-pot chains in Vietnam, including Kichi-Kichi, Dookki, Lầu Wang, and GoGi House. Each record contains the restaurant name, review ID, original star rating (1–5), and review text (predominantly Vietnamese). The star-rating distribution is skewed: 5★ = 385 (56.0%), 4★ = 107 (15.6%), 3★ = 59 (8.6%), 2★ = 23 (3.3%), 1★ = 114 (16.6%).

3.2. Pipeline Architecture

The pipeline follows a Workflow Automation pattern (Trigger → Process → Action) and consists of three stages:

Stage 1 — Translation. Because the SenticNet lexicon and the prompting framework operate on English text, all Vietnamese reviews are translated to English using Gemini 2.5 Flash API. The system rotates among six API keys to circumvent rate limits and implements checkpoint recovery to resume upon interruption.

Stage 2 — Aspect–Sentiment Extraction (GenAI). Each translated review is submitted to Gemini 2.5 Flash with a carefully designed few-shot prompt containing five worked examples. The prompt instructs the model to: (a) extract only adjectives and adverbs expressing sentiment; (b) assign each word to one of six aspects—Staff, Location, Service, Food, Space, Price; (c) label each word as positive or negative; (d) handle negation (e.g., “not good” → negative); (e) return a structured JSON object. The output is saved in a JSON file with 688 entries.

Stage 3 — Polarity Scoring (SenticNet). A Python pipeline loads a curated SenticNet lexicon of 135 adjectives (polarity $\in [-0.964, 0.964]$) and 62 adverbs (modifier weight $\in [-1.50, 1.50]$). For each sentiment word extracted in Stage 2, the system: (i) looks up the base polarity of the adjective; (ii) multiplies by the adverb’s modifier weight if present; (iii) inverts polarity upon negation; (iv) clamps the result to $[-1, +1]$.

3.3. Scoring Formulae

The polarity score $p \in [-1, +1]$ is mapped to a star rating $s \in [1, 5]$ via a linear transformation:

$$s = (p + 1) / 2 \times 4 + 1$$

The satisfaction level for aspect a_i is defined as:

$$Level_Sas(a_i) = \sum(\theta_k \times k) / n, \quad k \in \{1, 2, 3, 4, 5\}$$

where θ_k is the count of polar words at star level k , and n is the total number of polar words for that aspect. The overall satisfaction ratio is $S = n_{\text{positive}} / n_{\text{total}} \times 100\%$.

4. Results and Discussion

4.1. Aspect Extraction Summary

The pipeline extracted a total of 1,432 sentiment words from 688 reviews. Table 1 presents the detailed statistics for each aspect.

Table 1 Aspect-level sentiment extraction results (Mean = satisfaction score on a 1–5 scale)

Aspect	Reviews	Words	Pos.	Neg.	% Pos.	Mean	Std	Rank
Food	323	570	435	135	76.3	3.86	1.05	1
Location	43	49	39	10	79.6	3.80	1.13	2
Space	126	199	152	47	76.4	3.67	1.11	3
Staff	162	243	176	67	72.4	3.65	1.24	4
Price	70	102	74	28	72.5	3.53	1.09	5
Service	174	269	170	99	63.2	3.42	1.33	6
Total	—	1,432	1,046	386	73.0	—	—	—

Food dominates with 570 sentiment words (39.8% of total), followed by Service (269, 18.8%) and Staff (243, 17.0%). Location has the fewest mentions (49 words, 3.4%), likely because diners have already decided on the location before visiting. The overall positive-sentiment ratio is 73.0% (1,046 of 1,432), consistent with the original rating distribution, where 71.5% of reviews carry 4–5 stars.

4.2. Satisfaction Ranking by Aspect

Food achieves the highest Level_Sas (3.86), indicating that buffet and hot-pot customers are generally satisfied with food quality. The most frequent positive adjectives for Food are “delicious” (116 occurrences), “fresh” (47), and “good” (45). The most common negative expression is “not fresh” (12), signalling freshness as a critical differentiator.

Service scores lowest (3.42) with the highest standard deviation (1.33), suggesting inconsistent service quality across branches. The top negative terms are “bad” (12) and “not good” (9), primarily related to slow serving speed—a known challenge in the buffet format where peak-hour demand strains capacity.

The gap between the best (Food, 3.86) and worst (Service, 3.42) aspect is 0.44 points on a 5-point scale. This 12.9% relative gap is practically significant and suggests that restaurant chains should prioritise service-process improvements while maintaining food quality.

4.3. Top Polar Words

Table 2 Top-5 sentiment words per aspect (frequency in parentheses)

Food	Service	Staff	Space	Price	Location
delicious (116)	good (32)	friendly (29)	spacious (22)	reasonable (12)	good (5)
fresh (47)	bad (12)	good (15)	clean (14)	good (10)	nice (5)
good (45)	not good (9)	clean (10)	beautiful (11)	cheap (10)	spacious (4)
okay (18)	attentive (9)	attentive (8)	comfortable (9)	expensive (4)	beautiful (4)
not fresh (12)	slow (8)	rude (7)	dirty (5)	delicious (8)	great (4)

4.4. Validity Assessment

To validate the pipeline, we computed the Pearson correlation between the mean SenticNet star score per review and the original Google Maps star rating. For the Food aspect (n = 323), $r = 0.676$ ($p < 0.001$), indicating a strong positive correlation. The imperfect correlation is expected because: (a) the star rating reflects the holistic experience, whereas the pipeline scores individual aspects; (b) the SenticNet lexicon (135 words) does not cover all sentiment expressions encountered; (c) machine translation may attenuate some emotional nuances. Nevertheless, an r above 0.6 is considered strong in social-science research and confirms that the pipeline captures the dominant sentiment trends.

4.5. Comparison with Related Work

Table 3 compares the proposed pipeline with related approaches.

Feature	Shin et al. [2]	Mathayomchan & Taecharungroj [3]	Carrasco & Dias [7]	This study
Method	TF-IDF + RF	VADER	ChatGPT	GenAI + SenticNet
Reviews	5,427	935,386	1,000	688
Source	Google Maps	Google Maps	TripAdvisor	Google Maps
Aspects	4	4	5	6
Multilingual	Korean	English	Portuguese	Vietnamese → English
Explainability	Low	Medium	Low	High (SenticNet)

The proposed pipeline differentiates itself through: (a) the use of GenAI for flexible, context-aware aspect extraction rather than fixed keyword matching; (b) transparent, knowledge-grounded polarity scoring via SenticNet; and (c) built-in multilingual support via machine translation.

5. Conclusion

This paper presented a pipeline for aspect-based sentiment analysis that combines the extraction capabilities of Generative AI (Gemini 2.5 Flash) with the interpretable polarity scoring of SenticNet. Applied to 688 Vietnamese restaurant reviews from Google Maps, the pipeline extracted 1,432 sentiment words across six aspects and computed satisfaction scores on a 1–5 scale. Key findings include: Food is the most discussed and highest-rated aspect (3.86/5.0);

Service is the weakest aspect (3.42/5.0) with the highest variance; the overall positive-sentiment ratio is 73.0%; and the pipeline's output correlates strongly with original star ratings ($r = 0.676$).

The study contributes a practical methodology that merges GenAI's flexibility with SenticNet's transparency, offering restaurants an automated, interpretable tool for monitoring customer feedback at scale. Future work will expand the SenticNet lexicon with domain-specific culinary vocabulary, develop a native Vietnamese ABSA pipeline eliminating the translation step, benchmark multiple LLMs (GPT-4, Gemini, Llama) for extraction accuracy, and enlarge the dataset to improve generalisability.

Compliance with ethical standards

Disclosure of conflict of interest

No conflict of interest to be disclosed.

References

- [1] B. Liu, *Sentiment Analysis: Mining Opinions, Sentiments, and Emotions*. Cambridge University Press, 2015.
- [2] B. Shin, S. Ryu, Y. Kim, and D. Kim, "Analysis on review data of restaurants in Google Maps through text mining: Focusing on sentiment analysis," *Journal of Multimedia Information System*, vol. 9, no. 1, pp. 61–68, 2022.
- [3] V. Mathayomchan and V. Taecharungroj, "'How was your meal?' Examining customer experience using Google Maps reviews," *International Journal of Hospitality Management*, vol. 90, 102641, 2020.
- [4] J. O. Krugmann and J. Hartmann, "Sentiment analysis in the age of generative AI," *Customer Needs and Solutions*, vol. 11, no. 1, pp. 1–19, 2024.
- [5] W. Zhang, Y. Deng, B. Liu, S. J. Pan, and L. Bing, "Sentiment analysis in the era of large language models: A reality check," in *Findings of the Association for Computational Linguistics: NAACL 2024*, pp. 3881–3906, 2024.
- [6] E. Cambria, X. Zhang, R. Mao, M. Chen, and K. Kwok, "SenticNet 8: Fusing emotion AI and commonsense AI for interpretable, trustworthy, and explainable affective computing," in *Proceedings of HCI International (HCII)*, pp. 197–216, 2024.
- [7] P. Carrasco and S. Dias, "Enhancing restaurant management through aspect-based sentiment analysis and NLP techniques," *Procedia Computer Science*, vol. 237, pp. 129–137, 2024.
- [8] X. Sun, X. Li, S. Shi, T. Li, G. Chen, and J. Zhu, "Sentiment analysis through LLM negotiations," *arXiv preprint arXiv:2311.01876*, 2023.
- [9] J. Niimi, "Dynamic sentiment analysis with local large language models using majority voting: A study on factors affecting restaurant evaluation," *arXiv preprint arXiv:2407.13069*, 2024.
- [10] B. Ding, D. Qin, L. Liu, L. Bing, S. Joty, and B. Li, "Is GPT-3 a good data annotator?," in *Proceedings of the Annual Meeting of the Association for Computational Linguistics*, vol. 1, pp. 11173–11195, 2023.
- [11] E. Cambria, "Affective computing and sentiment analysis," *IEEE Intelligent Systems*, vol. 31, no. 2, pp. 102–107, 2016.
- [12] E. Cambria, Q. Liu, S. Decherchi, F. Xing, and K. Kwok, "SenticNet 7: A commonsense-based neurosymbolic AI framework for explainable sentiment analysis," in *Proceedings of LREC*, pp. 3829–3839, 2022.
- [13] T. Joseph, "Natural language processing (NLP) for sentiment analysis in social media," *International Journal of Computer Engineering*, vol. 6, no. 2, pp. 35–48, 2024.
- [14] D. Magdaleno, M. Montes, B. Estrada, and A. Ochoa-Zezzatti, "A GPT-based approach for sentiment analysis and bakery rating prediction," in *LNAI*, vol. 14502, pp. 61–76, 2024.