

Beyond Marketplace Reviews: Transformer-Based Sentiment Analysis of YouTube Influencer Ecosystems

Laura Nucera, Davide D’Emiliano, Andrea De Mauro
LUISS University, Rome, Italy

Abstract

Online reviews and influencer-centred social media conversations jointly shape product understanding in digital markets, yet they are still often studied as separate streams. This gap is especially relevant in the beauty sector, where product evaluation depends on subjective experience, perceived efficacy, tolerability, routine fit, and trust in peer-generated accounts. This study develops a cross-platform comparison of Amazon reviews and YouTube comments related to three active skincare products from the same brand, to examine whether the emotional profiles emerging from marketplace reviews differ from those found in influencer-centred comment spaces, and whether YouTube comments reveal needs, doubts, and advice-seeking signals that remain less visible in marketplace environments. The analysis is based on a replicable pipeline that combines data collection, filtering, text preparation, and transformer-based emotion detection using RoBERTa GoEmotions. The findings show that Amazon and YouTube do not produce equivalent emotional profiles. Across the three products, Amazon tends to concentrate emotional expression around admiration and approval, whereas YouTube comments display a broader and more fragmented emotional repertoire. The paper contributes theoretically by integrating review research and influencer ecosystems within a common comparative perspective, methodologically by proposing a transparent cross-platform workflow, and managerially by showing how YouTube-centred social listening can support the detection of consumer concerns that remain underrepresented in marketplace reviews.

Keywords: business model; digital platforms; human–AI interface; intelligent knowledge; data governance.

Paper type: Academic Research paper

1 Introduction

In contemporary digital ecosystems, online reviews are central to consumer decision-making. Despite this centrality, much of the literature continues to favour sentiment analysis of text-based reviews published on brand-led environments like Amazon, while paying less attention to evaluative content produced in influencer-led communication spaces, such as YouTube. This gap is particularly relevant in the beauty sector, where brands operate within multi-platform ecosystems in which company-sanctioned spaces coexist with more spontaneous, socially mediated discursive environments. In this scenario, YouTube's relevance to the customer journey is supported by market evidence and literature, which suggests that conversations around video reviews act as an 'extended focus group,' producing rich insights into consumer needs, preferences, and evaluations (Kabel et al., 2024). From these premises, the research gap that motivates this study is twofold. First, the literature still tends to treat marketplace reviews and conversations developed in influencer spaces separately, without building a systematic comparative framework between sources that, from the consumer's perspective, jointly contribute to the formation of product perceptions. Second, much research continues to operate through aggregate categories of positive or negative polarity, which are not always sufficient to capture finer differences in emotional profiles. Emotions with the same valence can produce different effects, and emotional language is often contextual and difficult to reduce to simplified classifications (Mustak et al., 2024). For this reason, cross-platform comparison requires an analytical framework capable of preserving differences across socio-technical contexts while simultaneously making different forms of evaluative expression comparable. From this perspective, transformer models such as BERT and RoBERTa are particularly relevant for analysing sentiment and emotions, although they require explicit attention to the limitations and biases of available benchmarks (Aka Uymaz & Kumova Metin, 2022; Demszyk et al., 2020). In light of this, the study develops a cross-platform comparative analysis of online opinions on beauty products. It compares texts from company-sanctioned platforms and influencer spaces. The goal is to understand how each environment contributes to the development of emotional sentiment profiles and which emotional dimensions differ across channels. The expected contribution is threefold. Theoretically, the study links eWOM, online reviews, and research on influencer communication. Methodologically, it offers a comparative logic for analysing diverse texts. Empirically and managerially, a multi-source sentiment reading helps identify emerging themes, perceived issues, and weak signals.

The research questions follow from this framework and guide comparison between official brand communications and influencer-led social conversations:

- *RQ1: How do emotional profiles differ between marketplace review environments and influencer-led social conversations in beauty product discourse?*
- *RQ2: What consumer concerns, doubts, and interpretive signals emerge in influencer-led social conversations that remain less visible in marketplace reviews?*

2 Literature Review

The literature on user-generated content highlights the value of online reviews in digital marketplaces, demonstrating how they reflect more than simple satisfaction or dissatisfaction. In the beauty sector, this centrality is even more pronounced, as reviews of cosmetics and skincare products help interpret practical aspects such as skin compatibility, tolerability, side effects, suitability for individual needs, and perceived risk reduction (Oh, 2020). Kim et al. (2025) emphasise the role of readability, perceived reviewer trustworthiness, and identity cues, while Jia et al. (2023) show that reviews published on marketplaces have a more immediate impact on sales than content on third-party platforms because they are closer to the moment of purchase. Overall, online reviews guide the interpretation of products characterised by high uncertainty and strong personal involvement. However, their effectiveness depends not only on the text, but on the interaction between information quality and the platform's selective mechanisms: visibility and position influence review adoption (Alzate et al., 2024), authenticity and trust signals strengthen their reliability (Liu and Dawod, 2025), while review environments impact brand loyalty and eWOM (Tran et al., 2025). Platforms like Amazon are therefore not simple repositories of opinions, but socio-technical environments that select, amplify, and make certain voices more persuasive. This is precisely where a limitation of the literature emerges: studies on marketplaces clarify helpfulness, trust, and risk reduction, but often treat the marketplace as an almost self-sufficient context, isolating it from other discursive environments that contribute to product perception (Matassi & Boczkowski, 2021; Mustak et al., 2024). From this perspective, YouTube represents a profoundly different context. Unlike marketplaces, which are dominated by review-oriented texts, YouTube creates a richer, more relational, and narrative environment. Beauty channels are perceived not only as entertainment spaces, but also as places for pre-purchase research, learning, and indirect product observation. Lee and Lee (2022) show that these channels foster both information acquisition and vicarious experiences that partially compensate for the lack of direct contact with the product. Kabel et al. (2024) suggest that spontaneous online conversations can be valuable sources for identifying needs expressed in natural language, as they offer rich, contextualised data on the current use of products and services. This literature explains influencers' persuasive mechanisms but focuses primarily on outcomes such as purchase intention, trust, and engagement. The gap in this research, therefore, lies in the fact that we know fairly well why creators are influential, but less so how evaluative and emotional signals differ between marketplace reviews and conversations generated around creator-driven content. A second line of inquiry concerns the most suitable analytical level for studying these texts. The distinction between sentiment analysis and emotion detection is crucial because it refers to two different levels of interpretation of the evaluative discourse. Traditional sentiment analysis reduces the text to a general polarity - positive, negative, or neutral - offering a synthetic measure of the overall orientation. However, this reduction is often insufficient in the beauty sector, where texts rarely express univocal judgments and combine appreciation, doubt, disappointment, curiosity, or relief with respect to different dimensions of the consumer experience. For this reason, recent literature is moving towards more granular models, capable of distinguishing not only the general tone, but also specific aspects and differentiated emotions (Aka Uymaz & Kumova Metin, 2022; Mustak et al., 2024; Mouyassir et al., 2024; Ng et al., 2024). Mouyassir et al. (2024) and Ng et al. (2024) show that transformer models improve the ability to capture this internal plurality. Emotion detection thus adds a further layer, as texts with the same valence can convey different affective states and different interpretive implications. In this context, the GoEmotions benchmark by Demszky et al. (2020), with 27 emotional categories and one neutral one, represents an important step towards finer taxonomies. Finally, the literature on cross-platform comparison shows that cross-platform comparison cannot be treated as a simple descriptive extension of analysis. Digital platforms offer different affordances, usage practices, and relational logics, which influence the production, circulation, and interpretation of content (Matassi & Boczkowski, 2021). This gives rise to the concept of analytical comparability: reviews and comments are not comparable because they are structurally identical, but because they represent different and complementary expressions of evaluation, interpretation, and reaction within the same consumption domain. The value of comparison lies in recognizing these differences as part of the object of study, not as methodological noise. The interpretative robustness of research therefore

depends not on the total elimination of biases, but on their explicitness and treatment as an integral part of the comparison itself.

3 Methodology

This study adopts a cross-platform comparative design to examine how consumer discourse differs between marketplace review environments and influencer-driven social conversations. Amazon reviews and YouTube comments serve as empirical contexts to compare two distinct forms of user-generated content. The empirical context is the beauty industry, where product evaluation is highly experiential. The analysis focuses on three active skincare products from the same brand, a choice that reduces cross-brand variability and keeps the comparison focused on platform effects and discursive environments, rather than differences in brand positioning. Amazon reviews were collected using the Apify scraper. Because the scraper imposed a limit on the number of reviews that could be retrieved, a proportional sampling strategy was adopted based on the distribution of star ratings observed on Amazon UK for each product. Amazon UK was selected to maintain linguistic consistency with the English-speaking scope of the study. For each product, a final set of 100 Amazon reviews was created, and review titles and bodies were merged into a single text field to preserve the evaluative content of each review as a single unit of analysis. The YouTube dataset was constructed through a multi-step workflow implemented in KNIME and integrated with Python. Three query formulations were used for each product: "review," "first impressions," and "wear test," as they are common in beauty discourse on YouTube and capable of capturing different product-centered video contexts. Candidate videos were retrieved via the YouTube Data API and then cleaned by removing duplicates, incomplete records, and entries with missing identifiers. Video-level metadata was added via a second API request, and only videos that met minimum quality criteria, including a minimum threshold of 50 comments, were retained. Comments were then extracted using a Python script that queried the commentThreads endpoint, managed pagination, and saved the comment text along with contextual metadata, including comment ID, author, publication date, like count, video title, video ID, and query type. Duplicate comments and records with missing text were removed before export. The study focuses on English-language content. Amazon reviews were collected from Amazon UK, while YouTube comments were filtered for English-language content, despite potentially coming from a broader set of English-speaking users. This difference is recognized as a limitation of the cross-platform design, as platform composition and audiences may partially overlap with geographic and cultural variations. Because YouTube comment sections contain many message types that are not directly comparable with marketplace reviews, all extracted YouTube comments were subjected to a preliminary ChatGPT-based classification phase. The same prompt structure was applied to all three products, modifying only the name of the focal product. Each comment was assigned to one of five unique categories: `comment_creator`, `target_review`, `other_product`, `advice_oriented`, and `uncertain`. The goal of this phase was to isolate comments directly relevant to the focal product and functionally comparable to Amazon reviews. Comments classified as `target_review` were used in the main cross-platform comparison of emotions, as they report evaluations, experiences, results, effects, or reactions to the product. Prior to validation, a gold standard was developed, defining the five categories through operational definitions, positive examples, borderline cases, and tiebreak rules. Comments reporting a clear personal experience with the focal product were assigned to `target_review`, while comments providing advice or asking about usage instructions, product compatibility, frequency of use, or skin type suitability were assigned to `advice_oriented`. The `uncertain` class was narrowly defined as a residual category for comments that were too short, irrelevant, incomprehensible, or impossible to reliably assign to one of the other classes. To validate the ChatGPT-based classification, a stratified sample of 300 comments was re-annotated independently by the first author using the default gold standard, without access to the model's predictions. The sample was balanced by product and by ChatGPT-predicted class, with 100 comments per product and 20 comments for each predicted class within each product.

Table 1. Validation metrics by product

product_name	n_comments	accuracy	cohen_kappa
Peeling Solution	100	0,71	0,6375
Hyaluronic Acid Serum	100	0,76	0,7
Salicylic Acid Solution	100	0,92	0,9

Table 2. Human–ChatGPT classification confusion matrix (%)

	chatgpt_comment_creator	chatgpt_target_review	chatgpt_other_product	chatgpt_advice_oriented	chatgpt_uncertain
human_comment_creator	81,25%	9,38%	3,12%	1,56%	4,69%
human_target_review	4,26%	87,23%	4,26%	2,13%	2,13%
human_other_product	0%	8,33%	76,67%	10%	5%
human_advice_oriented	5%	6,67%	10%	78,33%	0%
human_uncertain	4,35%	5,8%	5,8%	7,25%	76,81%

The validation of the ChatGPT-based classification showed an overall accuracy of 79.67% and a Cohen’s κ of 0.746. Following the pre-set decision rule ($\kappa \geq 0.61$ = substantial, Landis & Koch, 1977), the LLM-assigned labels were retained for the full corpus. The confusion matrix shows strong overall agreement between human and ChatGPT annotations, with the highest consistency for target_review, followed by comment_creator, advice_oriented, uncertain, and other_product. However, the matrix also reveals systematic ambiguities, especially between other_product and advice_oriented, where product comparisons are often interpreted as requests for guidance. Additional confusion appears between comment_creator and target_review, suggesting that comments directed at the creator may sometimes be read as evaluations of the target. The uncertain category present some misclassifications spread across advice_oriented, target_review, and other_product, indicating that vague or short comments are often forced into more specific categories by the model. The final product-level datasets underwent light text cleaning in Python. User mentions, hashtags, URLs, and redundant whitespace were removed, but no aggressive normalisation was applied. This choice was made to reduce technical noise while preserving the informal and context-dependent texture of user-generated discourse.

Emotion detection was carried out using SamLowe/roberta-base-go_emotions, a RoBERTa-based model fine-tuned on the GoEmotions benchmark. Instead of retaining only the first predicted emotion, the model was run with top_k=None in order to extract the full score distribution across the 28 GoEmotions labels. For each comment, the pipeline stored the scores for all 28 GoEmotions labels, the top three emotions, and their associated model scores. The full emotion-score distribution was then aggregated into five macro-emotional categories: positive, negative, uncertainty, cognitive, and neutral. The positive category includes emotions such as admiration, approval, gratitude, joy, love, optimism, and relief. The negative category includes anger, annoyance, disappointment, disapproval, disgust, fear, nervousness, remorse, and sadness. The uncertainty category includes confusion and curiosity. The cognitive category includes realization and surprise. The neutral category includes the neutral label. For each comment, macro-emotion scores were computed by summing the scores of the emotions belonging to each macro-category. The dominant macro-emotion was also identified as the macro-category with the highest aggregated score. However, the main comparative analysis relies on the weighted macro-emotion distribution rather than only on the dominant label, since weighted scores preserve the multidimensional emotional profile of each comment. Because GoEmotions was fine-tuned on Reddit comments, model performance on Amazon and YouTube text is not guaranteed. To validate the classifier in-domain, we manually re-annotated a stratified random sample of 150 valid texts. We report Cohen's κ between the model and the human gold standard at the five macro-category level (Positive / Negative / Cognitive / Uncertainty / Neutral), $\kappa = 0.75$ (accuracy = 0.83). Since the validation supports more reliable agreement at the macro-category level, the inferential analysis in §4 is conducted at the macro-category level, while fine-grained results are presented as descriptive.

4 Results

4.1 Differences in weighted macro-emotions across platforms

The aggregate distribution of weighted macro-emotions shows a clear difference between marketplace reviews and influencer-led conversations. Amazon reviews are strongly concentrated around positive emotions, which represent approximately two-thirds of the overall weighted emotional profile. YouTube comments, on the other hand, exhibit a more distributed emotional structure: positive emotions remain present, but have a much lower relative weight, while negative, uncertainty, and neutral components increase.

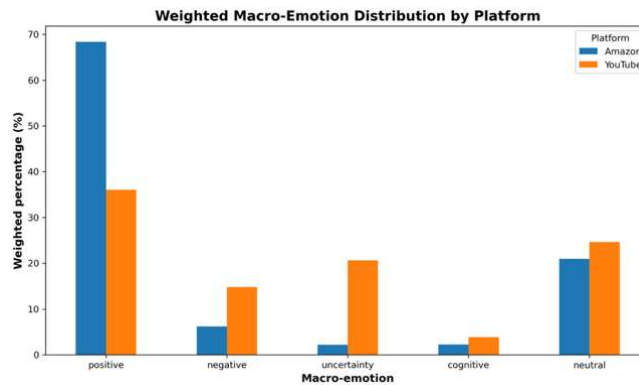


Figure 1 – Weighted Macro-Emotion Distribution by Platform

4.2 Product-level robustness check

The product-level heatmap confirms that a single product does not drive the aggregate pattern. In all three cases analysed, YouTube comments show a lower positive component than Amazon reviews. The difference is particularly marked for Peeling Solution, followed by Hyaluronic Acid and Salicylic Acid. At the same time, YouTube shows higher values for the negative components for all products and higher values for uncertainty, especially for Peeling Solution, potentially associated with concerns about frequency of application, skin sensitivity, irritation, and compatibility with other products. These results suggest that influencer-led conversations become particularly informative when the product requires interpretation, caution, or practical guidance. The product-by-product comparison reinforces the aggregate result, but also shows that the intensity of the difference between platforms varies based on product characteristics and perceived risk level.

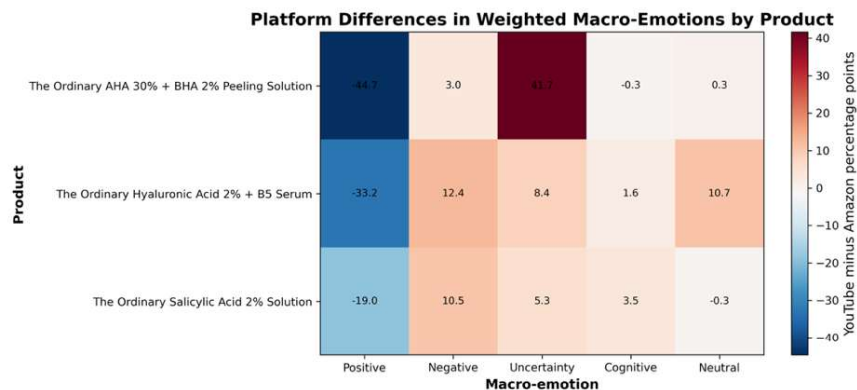


Figure 2 – Platform Differences in weighted Macro-Emotions by Product

4.3 Specific emotions and emotional ambiguity

The distribution of the top ten emotions helps better interpret the results observed at the macro-category level. Amazon reviews are strongly driven by positive emotions, particularly admiration and, to a lesser extent, approval. This confirms that marketplace reviews tend to express direct appreciation and more concentrated evaluative judgments. YouTube comments, instead, are less concentrated around a single positive emotion. The neutral category is particularly visible, as is curiosity, along with other emotions such as confusion, sadness, gratitude, disapproval, and love.

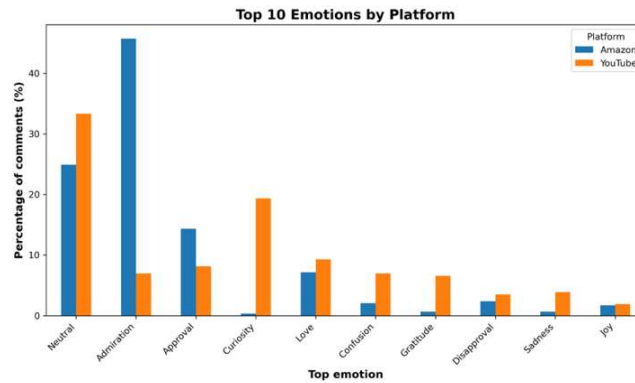


Figure 3 – Top 10 Emotions by Platform

Further evidence emerges from the average margin between the first and second emotions detected by the model. Amazon reviews show a higher average margin, indicating that RoBERTa more clearly identifies a dominant emotion. YouTube comments, however, show a lower average margin, suggesting greater emotional ambiguity or the coexistence of different emotional signals.

5 Discussion and Conclusions

This study compared marketplace reviews and influencer-led social conversations to understand whether these two digital environments reveal different emotional profiles and information needs in the context of beauty products. The results confirm that the two sources are not interchangeable: even when the comparison is based on balanced samples and YouTube comments filtered for product relevance, marketplace reviews and influencer-led conversations display different emotional structures. In answer to RQ1, regarding the differences between emotional profiles in the two environments, the results show that marketplace reviews are characterized by a greater concentration of positive emotions. Emotional expression on Amazon appears more stable, evaluative, and oriented toward synthetic product judgments. In contrast, YouTube comments display a more heterogeneous emotional profile, with a lower concentration of positive emotions and a greater presence of negative components, uncertainty, and cognitive processing. In response to RQ2, regarding additional signals emerging in influencer-led conversations, the results indicate that YouTube makes doubts, requests for clarification, uncertainties about use, and weak signals related to practical experience with the product more visible. This element emerges both from the emotional analysis, in which uncertainty and cognitive components are more evident in YouTube comments, and from the preliminary classification of comments, which identified a significant presence of content classifiable as advice-seeking. From a theoretical perspective, the study contributes to the literature on user-generated content, online reviews, and influencer ecosystems by showing that platforms are not simply neutral containers of opinions, but discursive environments that shape consumer experience. Marketplace reviews tend to make the final judgment observable, often expressed in a more concise and evaluative form. Influencer-led conversations, on the other hand, allow us to observe the process through which users interpret, negotiate, and contextualize product use. The methodological contribution involves the construction of a comparative and validated pipeline for analyzing heterogeneous content from different platforms. The preliminary classification of YouTube comments using ChatGPT was validated against a human gold standard, showing substantial agreement with human annotation. The emotional classification using RoBERTa-GoEmotions was also validated at the macro-category level, explicitly acknowledging the limitations of transferring a model trained on Reddit to different domains such as Amazon and YouTube. From a managerial perspective, the findings suggest that beauty companies should not rely solely on marketplace reviews to understand consumer perceptions. The latter are useful for monitoring satisfaction, approval, and summary product reviews, but they may be less effective in identifying emerging concerns, difficulties with use, or unclear information needs. Influencer-led social conversations, however, can serve as a particularly useful social listening space for identifying recurring questions, concerns and requests for product education. These insights can support brand communication, the development of FAQs, the selection of content to be shared with creators, and the improvement of instructions for use.

The study has some limitations that should guide the interpretation of the results. The analysis used only one brand and three products. The final comparison is based on balanced samples, not the whole corpus. Only English YouTube comments and Amazon UK reviews directly related to the product were included. The emotion

detection model was applied outside its training context. Thus, these results show cross-platform differences within a specific setting, not universal emotional signatures. Future research could include a program trained for the specific communication environment, more brands, larger samples, other product categories, or the analysis of excluded categories, such as advice-seeking comments.

Despite these limitations, the study shows that comparing platforms allows for a richer understanding of consumers' digital discourse and, in particular, that influencer-led environment make visible dimensions of product interpretation that are only partially observable in brand-led marketplace environment.

6 Ethics Declaration

Ethical clearance was not required for this research, as the study relied exclusively on publicly available online user-generated content and did not involve direct interaction with human participants. The data were analysed in aggregated form, and no personal identifying information is reported in the paper.

7 AI Declaration

AI tools were used to support the preliminary classification of YouTube comments into analytical categories, as described in the methodology section. AI tools were also used to support language editing and reference-format checking. All outputs were reviewed by the authors, who remain fully responsible for the research design, analysis, interpretation of results, and final content of the paper.

References

- Uymaz, H. and Kumova Metin, S. (2022) 'Vector based sentiment and emotion analysis from text: a survey', *Engineering Applications of Artificial Intelligence*, 113, 104922. doi: 10.1016/j.engappai.2022.104922.
- Alzate, M., Arce-Urriza, M. and Cebollada, J. (2024) 'Is review visibility fostering helpful votes? The role of review rank and review characteristics in the adoption of information', *Computers in Human Behavior*, 153, 108088. doi: 10.1016/j.chb.2023.108088.
- Demszky, D., Movshovitz-Attias, D., Ko, J., Cowen, A., Nemade, G. and Ravi, S. (2020) 'GoEmotions: a dataset of fine-grained emotions', in *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*. Association for Computational Linguistics, pp. 4040–4054. doi: 10.18653/v1/2020.acl-main.372.
- Jia, Y., Feng, H., Wang, X. and Alvarado, M. (2023) "'Customer reviews or vlogger reviews?' The impact of cross-platform UGC on the sales of experiential products on e-commerce platforms', *Journal of Theoretical and Applied Electronic Commerce Research*, 18(3), pp. 1257–1282. doi: 10.3390/jtaer18030064
- Kabel, D., Martin, J., Elg, M. and Witell, L. (2024) 'Capturing the voice of the customer: focus groups versus netnography?', *Total Quality Management & Business Excellence*, 35(11–12), pp. 1359–1377. doi: 10.1080/14783363.2024.2369592.
- Kim, S., Park, S., Li, X. and Kim, J.K. (2025) 'A deep learning-based review helpfulness prediction system for online beauty products', *IEEE Access*, 13, pp. 202049–202061. doi: 10.1109/ACCESS.2025.3638451
- Landis, J. R., & Koch, G. G. (1977). The measurement of observer agreement for categorical data. *Biometrics*, 33(1), 159–174. Doi: <https://doi.org/10.2307/2529310>
- Lee, M. and Lee, H.-H. (2022) 'Do parasocial interactions and vicarious experiences in the beauty YouTube channels promote consumer purchase intention?', *International Journal of Consumer Studies*, 46(1), pp. 235–248. doi: 10.1111/ijcs.12667.
- Liu, X. and Yahya Dawod, A. (2025) 'When technology signals trust: blockchain vs. traditional cues in cross-border cosmetic e-commerce', *Information*, 16(10), 913. doi: 10.3390/info16100913.
- Mahbub, F.B., Mim, M.H. and Rahman, M.S. (2024) 'Sentiment analysis of beauty product reviews in Bangladesh: a comprehensive study using transformer-based models', in *2024 27th International Conference on Computer and Information Technology (ICCIIT)*. IEEE, pp. 2712–2717. doi: 10.1109/ICCIIT64611.2024.11021928.
- Matassi, M., & Boczkowski, P.J. (2021). An Agenda for Comparative Social Media Studies: The Value of Understanding Practices From Cross-National, Cross-Media, and Cross-Platform Perspectives.
- Mouyassir, K., Fathi, A. and Assad, N. (2024) 'Elevating aspect-based sentiment analysis in the Moroccan cosmetics industry with transformer-based models', *International Journal of Advanced Computer Science and Applications*, 15(6). doi: 10.14569/IJACSA.2024.0150654.
- Mustak, M., Hallikainen, H., Laukkanen, T., Plé, L., Hollebeek, L.D. and Aleem, M. (2024) 'Using machine learning to develop customer insights from user-generated content', *Journal of Retailing and Consumer Services*, 81, 104034. doi: 10.1016/j.jretconser.2024.104034.
- Ng, C.M., Tiun, S. and Sastria, G. (2024) 'Multi-label aspect-sentiment classification on Indonesian cosmetic product reviews with IndoBERT model', *International Journal of Advanced Computer Science and Applications*, 15(11). doi: 10.14569/IJACSA.2024.0151168.
- Oh, Y.-K. (2020) 'Determinants of online review helpfulness for Korean skincare products in online retailing', *Journal of Distribution Science*, 18(10), pp. 65–75. doi: 10.15722/jds.18.10.202010.65.

Tran, S.N., Trương Bích, P., Doan, T.L.M., Lam, V.L.C. and Tran Thi Anh, T. (2025) 'The impact of online reviews on brand loyalty and e-WOM in Vietnam's cosmetics e-commerce sector: the moderating role of e-service quality', *Social Sciences & Humanities Open*, 12, 102208. doi: 10.1016/j.ssaho.2025.102208.