

Leather Apparel Design and Sustainable Fashion Communication Based on Film and Television Media

Xiaoyue Zhang¹

¹SVA Film/ Lens, Screen and Digital Arts, School of Visual Arts, New York 10011, NY, USA

Corresponding author: Xiaoyue Zhang (e-mail: xiaoyueyueazq@163.com)

ABSTRACT Existing research on sustainable fashion communication lacks effective analysis of the visual-symbol transformation mechanism in film and television media. This paper proposes a hybrid method integrating visual content analysis and dynamic theme modeling. First, a visual-feature encoding framework for leather clothing in film and television is constructed. Then, the BERTopic model is applied to perform semantic clustering of related film and television texts and public comments. The model uses BERT to generate contextual embeddings, UMAP for dimensionality reduction, HDBSCAN for density clustering, and c-TF-IDF for latent theme extraction. Finally, cross-modal association analysis quantifies the mapping relationship between visual elements and semantic themes. From an engineering communication perspective, the framework offers a signal-encoding and cross-modal interpretation method that can also inform visual evaluation of smart materials, wearable surfaces, and electromagnetic-compatible product narratives without changing the fashion-communication focus of the study. Experimental results show that glossy leather and the theme of ethical consumption exhibit a Cramer V coefficient of 0.41, while cross-type comparisons show the strongest association between distressed finishes and durable classics in drama films, reaching 0.51. This study provides a quantifiable analytical paradigm for the dissemination of sustainable consumer culture driven by visual media.

INDEX TERMS leather apparel, sustainable fashion communication, dynamic topic modeling, cross-modal analysis, visual signal encoding

I. INTRODUCTION

FILM and television media are currently the most powerful medium for cultural communication, and its impact on the public relates to the audience's aesthetic orientation and consumption ideologies. In fashion, the costumes worn by the characters in film and television are not only visually relevant to its respective narratives but are also representational of specific lifestyles and values [1,2]. It is practically useful to adopt sustainable fashion narratives, especially concerning leather apparel and environmental controversies, in film and television narratives, and disseminate them effectively to direct green consumption culture and serve the purpose of transforming the industry [3,4].

Nonetheless, exploring how leather clothing expresses sustainable values sustainably in film and television contexts is fraught with problems. Existing research has difficulty making the implicit link between visual signs and abstract values explicit [5,6]. The sustainable character of clothing in film and television is often less explicit than the texture of material, character construction, and plot construction in a multi-tiered semantic network. How the public interprets clothing in film or television is not static but dynamic;

moreover, traditional content analysis does not adequately capture the dynamic transformation of meaning from visual representation, through the interpretive framework, to public cognition [7,8]. Hence, for all these reasons, existing research cannot draw systematic conclusions about the complete chain of "what kind of visual design, within what narrative framework, by what interpretive pathway, anchors what type of sustainability in terms of values."

In order to deal with these difficulties, previous attempts to respond to this issue have occurred through various methods using existing research. Qualitative text analysis would be able to uncover audience feedback with some subjective interpretation on the part of the researcher. However, it is difficult to manage and analyze large-scale text data and objective conclusions are tenuous at best [9,10]. Traditional topic models (e.g., LDA) can manage a set of texts, but they do this with a bag-of-words model principle, which does not account for the context of the words, making it difficult to discern the semantic difference between "leather" in "luxury leather" and "eco-friendly leather" and therefore topic extraction accuracy is insufficient [11,12]. Also, visual analysis methods mostly explored visual characteristics

through classifying and describing clothing style; visual analysis methods would not serve as an effective bridge to visual features and textual predictions [13,14]. These methods all have limitations in analyzing cross-modal, implicit, and dynamically evolving semantic communication processes, preventing the deep-seated mechanisms of film and television media in sustainable fashion communication from being fully revealed. To address the dynamic semantic quantification problem in the transformation of film and television leather costume design towards sustainable concepts, this paper employs a hybrid research method combining visual encoding and contextual topic modeling. This method first constructs a structured encoding table encompassing visual attributes, character attributes, and scene attributes of the costumes, used to systematically deconstruct film and television samples. Subsequently, large-scale film and television reviews and social media texts corresponding to the samples are collected, and dynamic semantic mining is performed using BERTopic topic modeling technology based on a BERT pre-trained language model. This technology obtains deep contextual embeddings of the text through a pre-trained sentence-transformer model, performs dimensionality reduction using the UMAP algorithm to preserve the global semantic structure, and uses the HDBSCAN density clustering algorithm to identify semantic clusters. Ultimately, it extracts and refines core topic words for each cluster with c-TF-IDF. By examining the correlation strength between the visual encoding matrix and the semantic topic network, it creates a quantifiable mapping model from concrete visual symbols to abstract sustainable values. This method can successfully identify and quantify the mapping relationship between visual evidence and sustainable topics, and identify its differentiated signaling across different types of film and television. The framework provides a new analytical framework for studying visual communication in sustainable fashion, but also accomplishes the targeted identification and quantifying of the environmental value encoded in film and television narratives, thus creating a powerful closed loop with the "research," and allowing for a strong foundation for deeper research and analysis that can follow.

II. DESIGN AND CONSTRUCTION OF CROSS-MODAL ANALYSIS FRAMEWORK

A. CONSTRUCTION OF FILM AND TELEVISION CORPUS AND VISUAL FEATURE ENCODING

In order to systematically analyze the visual narrative and sustainable dissemination of leather clothing in film and television media, this study first established a specialized multimodal corpus. The collection of the corpus was scripted along systematic principles and approached from three core dimensions based on sampling time, genre characteristics, and cultural significance. The sampling time period for the corpus covered almost the past twenty years. The final sample corpus consisted of fifty of the most representative film and television works within the

visual showcasing leather clothing. The selection was based on a composite index combining (1) domestic box office revenue or viewership ratings, (2) volume of public discussion (social media mentions and review counts), and (3) frequency of appearance in professional film/television award nominations or "best-of" lists related to costume design. This multi-dimensional criterion ensured the corpus captures works with high visibility, public engagement, and professional recognition. To capture public interpretations of clothing in film and television, the study simultaneously collected large-scale commentary text data for each of the selected samples. The data sources included public film and television databases and social media. Texts were iteratively filtered for relevant keywords and reduced to ensure a high degree of relevance between the text and visual analysis items.

In terms of visual analysis, this research establishes a systematic 'coding' system of visual features to systematically decompose the primary leather clothing featured in video film clips [15,16]. The coding system divides visual features into three dimensions: objective attributes, narrative attributes, and sustainability implications. Objective attributes record the clothing's category, silhouette, and surface treatment process; narrative attributes relate to character positioning and scene atmosphere; and the sustainability implications dimension, based on clothing semiotics theory, infers the potential value that the design may convey, such as durability, ethical consumption, or material innovation. The complete coding dimensions, operational definitions, and corresponding codes are shown in Table 1.

TABLE 1. Visual Feature Coding Table for Film and Television Leather Costumes

Main dimension	Secondary dimension	Encoding and operation definitions
Objective attributes	Apparel category	P1: Jacket; P2: Trousers; P3: Accessories; ...
Objective attributes	Visual materials	M1: Glossy leather; M2: Matte leather; M3: Distressed/wear-out finish; M4: Technical-looking fabric; ...
Narrative attributes	Character narrative function	C1: Hero/Protagonist; C2: Villain; C3: Sage/Mentor; ...
Narrative attributes	Scene context	S1: Combat/Action; S2: Casual/Daily; S3: Social/Ritual; ...
Sustainability implications	Value orientation	V1: Durable/Classic (reflects product longevity); V2: Ethical/Luxury (relates to material sourcing and consumption ethics); V3: Innovative/Future (associated with material technology and sustainable innovation); V4: Character Symbolism (e.g., individuality, rebellion—primarily narrative-driven, not directly a sustainability value).

To assess inter-coder reliability during the coding process, this study used the Cohen-Kappa coefficient for measurement [17,18]. The coefficient is calculated as shown in equation (1).

$$\kappa = \frac{P_o - P_e}{1 - P_e} \quad (1)$$

In equation (1), P_o represents the proportion of consistency observed by the coder; and P_e represents the probability of expected consistency calculated based on the coding distribution. The closer the κ value is to 1, the higher the consistency of the coding results, and the stronger

the objectivity and reliability of the coding system. After two rounds of coding training and discussion, reliability tests were conducted on all dimensions. The final average Cohen-Kappa coefficient showed that the coding results had a high degree of consistency, meeting the requirements for subsequent analysis. Finally, the target leather clothing in each film and television sample was transformed into a vector composed of multiple structured codes. All vectors together formed a visual feature matrix for subsequent association analysis.

B. DYNAMIC SEMANTIC TOPIC MINING BASED ON BERTOPIC

To achieve dynamic mining of implicit sustainable semantics in film and television texts, this study uses the BERTopic model to perform topic modeling on the collected large-scale comment texts. The process begins with text preprocessing. The original text data is transformed into a standardized word sequence after data cleaning, word segmentation and stop word removal, which lays the foundation for subsequent semantic vectorization [19,20]. The core of text vectorization is to generate deep context embeddings using pre-trained language models. The BERTopic model uses pre-trained sentence-transformer models from the sentence-transformers library to map each document (or sentence) into a dense vector in a high-dimensional vector space. This vector can capture the deep semantic information of words in a specific context, overcoming the deficiency of the traditional bag-of-words model in ignoring contextual relationships. Its vectorization process can be abstractly represented by equation (2):

$$\vec{d}_i = \text{Model}_{\text{BERT}}(T_i) \quad (2)$$

In equation (2), T_i represents the i -th preprocessed text unit; $\text{Model}_{\text{BERT}}$ represents the selected pre-trained Transformer model; and $\vec{d}_i$ is the high-dimensional semantic vector output by the model after the text unit is calculated. These document vectors form the basis for subsequent topic clustering.

High-dimensional vectors need to be reduced in dimensionality to fit clustering algorithms. This study uses the Unified Manifold Approximation and Projection algorithm, which, as a nonlinear dimensionality reduction technique, can better preserve the topological structure of data in low-dimensional space based on the local and global semantic relationships between vectors [21,22]. UMAP achieves dimensionality reduction by constructing a fuzzy topological structure in a high-dimensional space and optimizing its similarity to the corresponding structure in the low-dimensional space. Its optimization objective is to minimize the difference in cross-entropy between the beautified sets in the two spaces, and the core optimization function is shown in equation (3):

$$L_{\text{UMAP}} = \sum_{i \neq j} \left[p_{ij} \log \left(\frac{p_{ij}}{q_{ij}} \right) + (1 - p_{ij}) \log \left(\frac{1 - p_{ij}}{1 - q_{ij}} \right) \right] \quad (3)$$

p_{ij} is the similarity probability between vectors i and j in the high-dimensional space, and q_{ij} is their similarity probability in the low-dimensional space. Through this optimization, document vectors are projected onto a visualized two-dimensional or three-dimensional space, and their semantic proximity relationships are clearly presented.

In a low-dimensional vector space, the density-based HDBSCAN clustering algorithm is used to automatically identify semantically similar document clusters. HDBSCAN does not require a pre-defined number of clusters; it calculates the reach distance of each data point and constructs a minimum spanning tree, thereby compressing the cluster tree to extract stable clusters [23,24]. This algorithm assigns a cluster label to each data point and marks data points in sparse regions as noise, thus effectively improving the purity of topic extraction. The output of the clustering results defines the cluster affiliation of the text in the semantic space. Finally, interpretable topic tags are generated for each identified semantic cluster. BERTopic uses a cluster-based TF-IDF technique, which treats each cluster as an independent set of "documents" for computation [25,26]. The weights of c-TF-IDF are calculated as shown in equation (4):

$$W_{t,c} = \text{tf}_{t,c} \cdot \log \left(1 + \frac{A}{\text{df}_t} \right) \quad (4)$$

In equation (4), $W_{t,c}$ is the weight of word t in cluster c ; $\text{tf}_{t,c}$ is the frequency of word t in cluster c ; A is the total number of clusters in the dataset; and df_t is the number of clusters containing word t . Through calculation, the top few words with the highest weight in each cluster are extracted to form the core keywords of the semantic topic. After the above process, the text data is transformed into a structured topic-document distribution, providing structured semantic input for subsequent cross-modal association analysis. The semantic clustering and topic distribution of film and television review texts in the dimensionality reduction space are shown in Figure 1.

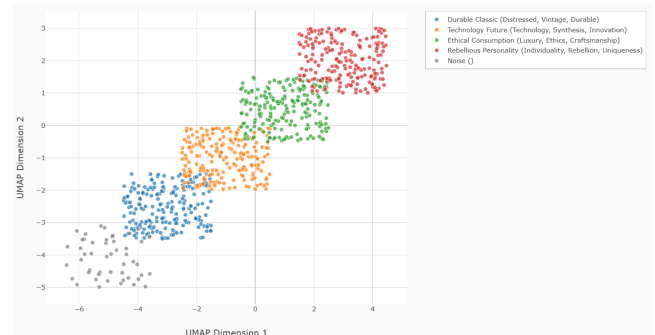


FIGURE 1. Semantic spatial distribution and topic clustering of film and television texts

Figure 1 presents the results of dynamic topic mining of film and television review texts based on the BERTopic model. The high-dimensional text semantic vectors were reduced to a two-dimensional plane using the UMAP algorithm, where each data point represents a text unit, and its spatial location reflects semantic similarity. HDBSCAN density clustering was used to identify text clusters with common semantic features, which were then distinguished by color coding. Gray dots represent noisy texts that failed to form a clear topic. The core semantics of each cluster were labeled with the top three keywords extracted by c-TF-IDF, clearly revealing the core topics related to the sustainability of leather clothing that naturally emerged in public discussions and their distribution structure. The clear boundaries and high internal cohesion between topic clusters indicate that the public's semantic cognition of film and television leather clothing exhibits obvious structural characteristics. The comparatively autonomous distribution of each theme in the semantic space indicates separate cognitive boundaries for different sustainability dimensions. The findings demonstrate the effectiveness of the BERTopic model in topic extraction on the studied corpus, providing a high-quality semantic topic baseline for subsequent visual-semantic associative analyses.

C. VISUAL-SEMANTIC CROSS-MODAL ASSOCIATION MODEL

This research establishes a cross-modal association model to quantify the mapping relationships between the visual characteristics of clothing made from leather and public semantic cognition. The model's data sets are applied as a visual feature matrix V and a text-topic label vector T . The aforementioned data sub-structures are aligned in a manner to produce a sample aligned with a schema where the visual encoding of each film/TV sample makes a sample pair with its text topic label, as a subset of the categorical data sample: $D = \{(v_i, t_i)\}_{i=1}^N$, where v_i represents the visual feature vector of the i -th sample, t_i corresponds to its topic label, and N is the sample size.

The core of association analysis is to quantify the co-occurrence strength between visual features and semantic topics. This research employs the chi-square test to evaluate the statistical significance of the connection between a selected visual attribute V_j and a selected semantic topic T_k [27,28]. A two-dimensional contingency table is constructed to count the number of samples O_{jk} that have both the V_j visual feature and T_k topic at the same time. The chi-square statistic is calculated as shown in equation (5):

$$\chi^2 = \sum_{j=1}^J \sum_{k=1}^K \frac{(O_{jk} - E_{jk})^2}{E_{jk}} \quad (5)$$

In equation (5), O_{jk} represents the observed frequency and E_{jk} corresponds to the expected frequency. This test evaluates the extent to which the amount of co-occurrence

among visual features and semantic topics is more than what could be randomly expected.

To further estimate the actual strength of the association effect, it calculates the Cramer V coefficient based on the chi-square test for significance. This coefficient provides a standard estimate of association that ranges from [0,1] with a larger number representing a stronger association [29,30]. The calculation equation is shown in (6):

$$\text{Cramér's } V = \sqrt{\frac{\chi^2}{N \cdot (\min(J, K) - 1)}} \quad (6)$$

In equation (6), J and K indicate the number of categories for visual features and semantic topics, respectively. This coefficient adjusts for sample size impact on the chi-square value so that correlation strength between different feature-topic pairs can be compared meaningfully.

In the end, the model outputs an association strength matrix A whereby elements A_{jk} represent the standardized association strength between visual features and semantic topics. Based on chosen thresholds θ , all $A_{jk} > \theta$ salience association pairings that meets criteria are extracted leading to a mapping set quantitatively mapping between visual symbols and sustainable semantics. This set can crawl the dissemination path and effectiveness of a particular design element of leather clothing in the context of film and television media as it pertains to transforming into a particular sustainable value and contributes a way for data support for the core research question. Quantitatively, the model has produced an adequate response to the core question of "what kind of visual design anchors what kind of sustainable value."

III. EXPERIMENTAL DESIGN AND DATA PREPARATION

A. CONSTRUCTION OF THE EXPERIMENTAL CORPUS

To achieve a systematic and representative sample for the research study, this experiment developed a specific film and television corpus as outlined above. To construct the corpus, a stratified sampling method was utilized to identify a framework for selection based on time period distribution, genre characteristics, and contextual cultural influence of the film and television content. In terms of time period, the corpus was designed to span a time frame of 2000-2023 to allow the collection of data which shows the historicity of fashion narratives; in terms of genre selection, the corpus is represented with four broad groupings of content: science fiction, drama, action, and fantasy, which allows for an analysis of how clothing symbols are differently presented in the context of the respective genres; in terms of assessing cultural influence, public discussion heat from the most popular domestic television and film rating platforms was used, along with an evaluative index of professional film critics. The final sample database was made up of fifty visual representations of leather clothing in film and television, where all samples contained clips of the entire film and television content for analysis and relevant commentary text

data that was publicly available. The text data was generated from open film/television databases and social media by way of application programming interfaces (APIs). The collection process utilized a multi-round keyword retrieval process that incorporated film titles and character names, and was subject to extensive data deduplication, irrelevant text filtering, and basic cleaning preprocessing to produce an appropriately sized, high-quality multimodal analysis dataset. The comment texts were collected within a time window spanning from each film’s initial release date up to the data collection cutoff (December 2023), ensuring coverage of both immediate audience reactions and longer-term discursive evolution. Its basic statistical characteristics are shown in Table 2.

TABLE 2. Experimental Corpus Statistical Description

Statistical dimensions	Statistical value
Number of film and television samples (films))	50
Time span (years)	2000-2023
Type distribution (department)	Sci-Fi: 15, Drama: 16, Action: 12, Fantasy: 7
Total number of comment texts (items)	128,547
Total word count of the text (words)	4,826,119
Average text length (words/entry)	37.5
Data source platform	Douban Movies, Weibo

Note: Comment texts were aggregated across the entire available timeline for each film (release date to December 2023), encompassing both contemporaneous and retrospective public discourse.

Table 2 comprehensively describes the scale and composition of the experimental corpus. The total amount and quality of the text data provide the necessary data foundation for subsequent semantic topic mining, while the diversified distribution of film and television samples ensures the breadth and representativeness of the visual analysis dimensions.

B. EXPERIMENTAL PARAMETER SETTINGS AND EVALUATION INDICATORS

To ensure the reliability of the topic modeling process and the comparability of the results, this study fixed the key parameters of the BERTopic model and used topic consistency and diversity as the core evaluation indicators. The model used was the all-MiniLM-L6-v2 pre-trained model from the sentence-transformers library to generate text embeddings [31,32]. The dimensionality reduction process used the UMAP algorithm, with the parameters set as follows: the number of nearest neighbors was 15 and the minimum distance was 0.1. The clustering stage used the HDBSCAN algorithm, with the key parameter min_cluster_size set to 10. To comprehensively evaluate the quality of the generated topics, this study applied two traditional baseline models—Latent Dirichlet Allocation (LDA) and Nonnegative Matrix Factorization (NMF) [33,34]—as well as another deep learning-based neural topic model (NTM) for comparison [35]. The evaluation uses normalized point mutual information scores to measure the semantic consistency of words

within a topic, while also applying a topic diversity index to calculate the proportion of unique topic words among all topics generated by the model, in order to assess the breadth of topic representation. Table 3 shows a comparison of parameter configurations and topic quality for different models.

TABLE 3. Model parameter configuration and topic quality assessment

Model	Key parameter configuration	Thematic consistency	Theme diversity
BERTopic (This study)	Transformer: all-MiniLM-L6-v2; UMAP: n_neighbors=15, min_dist=0.1; HDBSCAN: min_cluster_size=10	0.72	0.89
NTM	Number of topics: 8; Encoder dimensions: 128; Number of iterations: 100	0.65	0.82
LDA	Number of topics: 8; Number of iterations: 1000; Learning method: online	0.58	0.75
NMF	Number of topics: 8; Initialization: nndsvda; β loss function: frobenius	0.61	0.78

Table 3 systematically compares the performance of different types of topic models on the specific corpus of this study. The results show that the BERTopic model used in this study outperforms the baseline model in both topic consistency and diversity. The semantic topics it generates maintain internal cohesion while covering a richer conceptual space, which provides high-quality and comprehensive semantic input for subsequent cross-modal association analysis.

IV. RESULTS AND ANALYSIS

A. DISTRIBUTION PATTERNS OF VISUAL FEATURES OF LEATHER COSTUMES IN FILM AND TELEVISION

To systematically analyze the visual composition paradigm of leather clothing in film and television works, this section conducts quantitative statistical and distribution analysis of the visual features of all samples in the corpus. Through structured coding and frequency statistics of three dimensions—clothing category, material texture, and narrative function—the mainstream forms and visual preferences of leather clothing design in film and television narratives are revealed, and their distribution patterns are shown in Figure 2a. Furthermore, to explore the intrinsic relationship between visual features and film and television genres, a type-feature co-occurrence model is constructed by calculating the proportion of specific visual elements in different genres, and the results are intuitively displayed in Figure 2b.

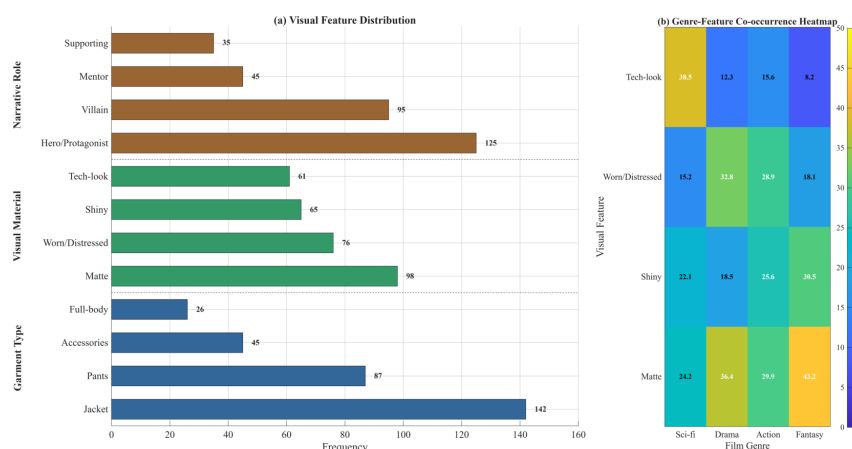


FIGURE 2. Distribution pattern of visual features in film leather costumes. (a) Distribution of visual features; (b) Heatmap of co-occurrence of type features

Figure 2a presents the frequency distribution of clothing categories, visual materials, and narrative functions. When considering clothing categories, jackets are represented 142 times and are the most dominant. This is attributed to their narrative convention in action and adventure genres of being functional outerwear. In terms of visual materials, matte leather and distressed remain the next most frequent with 98 and 76 instances, respectively. This attention to textural realism and use of character background information is representative of film and television styling. The distribution of narrative function shows that costumes for hero(protagonist) and villain costumes are most frequent with 125 and 95 instances respectively, indicating that leather clothing has a strong symbolic function in character representation and constructing visual images of contrast. The overall analysis shows that film and television leather clothing serves a reliable visual paradigm pattern dominated by functional outerwear jackets which emphasize material narrative texture and closely serves the construction of character identity.

Figure 2b demonstrates how visual features selectively align with film/television genres. In science fiction films, the use of technologically textured fabrics is 38.5% compared to the other genres, which is significantly higher than all the other genres. This aligns with the fact that science fiction stories depend on constructing narratives with material imaginings based on the use of fabrics or materials from the future. In dramas, worn or distressed fabrics are used 32.8%, which aligns with the desires of a plausible realism and sense of age in the narrative aesthetics of general realism. In fantasy, matte leather is the most utilized material (43.2%), primarily because its non-reflective, textured surface aligns with the genre's common costume archetypes—such as armor, medieval attire, and rugged adventurer wear—that emphasize practicality, historical or mythic authenticity, and a subdued, grounded aesthetic suitable for world-building. Thus, the prevalence of matte leather in fantasy reflects a deliberate design choice to support narrative cohesion and character plausibility within imagined worlds, rather than merely abstract aesthetic preferences. The data in

this distribution clearly demonstrated that visual features were not randomly or evenly distributed among film or television genres, or any other non-value based reason for visually displaying leather garments. Rather, visual features clump together based on their narrative logic and aesthetic expectations of each genre, proving that film or television genre is an important contextual factor shaping how leather clothing is visually presented.

B. IDENTIFICATION AND INTERPRETATION OF SUSTAINABLE-RELATED SEMANTIC THEMES

Once the study measured the visual qualities, it began to explore the semantic level of the text, to explore the cognitive structure and value interpretation that the public might have toward leather clothing within film and television. In connection to this, the BERTopic model performs dynamic topic modeling on the large amounts of comment text to discern and extract core semantic topics associated with sustainable fashion. The topic modeling also considers the strength and stability of each topic's distribution over the discourse. The results of the topic modeling is represented in two dimensions: the weight of core keywords and the probability distributions, as shown in Figure 3.

The core keywords and corresponding c-TF-IDF weights for three main sustainability themes is exhibited in Figure 3a. The vertical axis describes the sequence of keywords while the horizontal axis describes the semantic weight value. Analyzed within the "Durability and Classics" theme, "Vintage" and "Durable" achieve weights of 0.125 and 0.112 respectively, further indicating the audience's significant assessment of leather clothing as durable. This lexical association is logically grounded by leather's hyper-association with historical lineages and the passage of time in film and television. Analyzed within the "Technology and Future" theme, "Innovation" and "Technology" achieve weights of 0.118 and 0.105 respectively, suggesting the visual representation of emergent materials and future visions successfully animates the audience's association with technological innovation. The weights for the keywords within the "Ethical

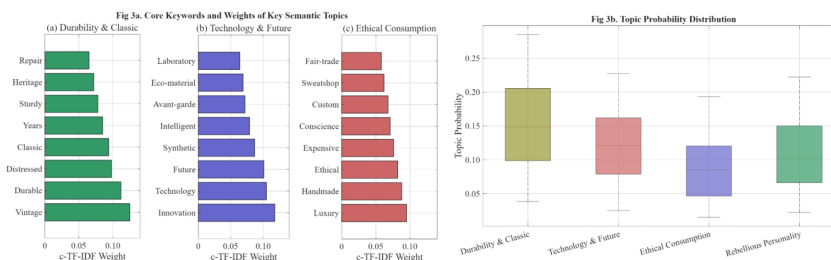


FIGURE 3. Core Keyword Weights and Probability Distribution. (a) Core Keyword and Key Semantic Topic Weights; (b) Topic Probability Distribution

Consumption" theme reached comparatively lesser weights, especially "Luxury" and "Ethical" which achieved weights of 0.095 and 0.082 respectively, further representing that discussions about ethical values may not have reached the level of density in public discourse.

The box plot shown in Figure 3b shows the differences in topic probabilities. The "durability and classicism" topic has a median probability of 0.148 and the widest distribution, indicating that this pattern of cognition is common among viewers. The "technology and future" topic has a right-skewed distribution indicating that some viewers expect a lot from it. The "moral consumption" topic has a median of only 0.085, which suggests public resonance with the ethics of a product is somewhat limited. This difference in cognitive structure is likely due to the differentiated representations of the different dimensions of sustainability in film and television narratives, as well as the selective reception of those narratives by viewers drawing from their own knowledge frameworks. It should be noted that certain themes identified in broader public discourse—such as those centered on individuality or rebellion—primarily reflect character symbolism and narrative-driven audience perceptions, rather than direct environmental or ethical sustainability values. These narrative-oriented themes are included in the analysis to provide a comprehensive mapping of the semantic landscape, while the core focus of sustainable fashion communication remains on the three value dimensions examined above. Overall analysis indicates that the public's semantic interpretation of leather clothing in film and television has a hierarchical cognitive structure based on practical durability, expanded by technological imagining, and augmented by ethics.

C. ANALYSIS OF THE CORRELATION STRENGTH BETWEEN VISUAL FEATURES AND SEMANTIC TOPICS

Further research is needed to explore the quantitative mapping relation between visual properties and semantic cognition based on the determination of core semantic themes to reveal the underlying mechanism by which the film and television symbol system transmits sustainable values. A visual-semantic association matrix was constructed by calculating the Cramer V coefficient, and key association pairs were selected for comparative validation against multiple baseline models to evaluate the effectiveness of the proposed method in capturing cross-modal associations. The

association strength analysis and comparative validation results are presented in Figure 4.

As shown in Figure 4a, the distressed visual features and the durable classic theme exhibit a strong and statistically significant correlation (Cramer's $V = 0.45$, $p < 0.001$). All reported Cramer's V coefficients are accompanied by p -values derived from chi-square tests, with $p < 0.05$ considered statistically significant, ensuring that the observed associations are not due to random chance. This relationship is because distressed effects were often frequently used in film and television narratives to indicate the history and durability of objects, consequently activating the viewer's senses of age and durability. The technologically textured fabric and technological future correlation has a 0.52, indicating that the luster and evenness of the texture of synthetic materials effectively signify a sense of the future regarding technological development. The correlation between glossy leather and the ethical consumption theme (Cramer's $V = 0.41$) can be contextualized through specific film narratives where glossy leather is used to signify craftsmanship, exclusivity, or heritage luxury—values often associated with ethical consumption in discourses around artisanal quality and sustainable luxury. For instance, in several drama and fantasy films sampled, glossy leather costumes are worn by characters portrayed as master artisans or custodians of tradition, visually aligning material sheen with notions of skilled labor and enduring value. Figure 4b provides further evidence for the reliability of the key correlation pairs. Throughout, the proposed method consistently outperforms the baseline model across all three correlation pairs. The Cramer's V coefficient for the technologically textured fabric-technological future correlation is 0.52, considerably higher than the NTM model at 0.45, the LDA model at 0.31, and the NMF model at 0.37. This superiority comes from the BERTopic model's adeptness at contextual semantics, affording it the ability to capture subtle correlation patterns that bag-of-words models fail to detect. Comprehensive analysis confirms that film and television leather costumes establish a systematic network of associations with the audience's semantic cognition through specific visual encodings, and this association can be effectively quantified and verified by the method proposed in this paper. It should be noted that the association between glossy leather and ethical consumption is not universally positive; in contexts where glossy finishes coincide with mass-produced or synthetic materials,

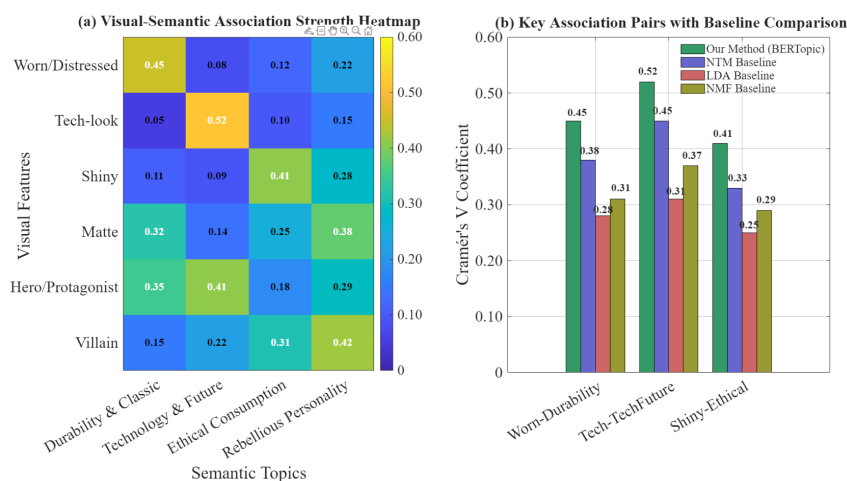


FIGURE 4. Association strength analysis and comparative verification. (a) Visual semantic association strength (Clem V coefficient); (b) Comparative analysis of key association pairs

the semantic linkage may invert or weaken. The correlation captured here reflects a dominant narrative pattern within the sampled corpus, rather than a fixed symbolic rule.

D. ANALYSIS OF THE INTERACTION BETWEEN THEME EVOLUTION AND FILM NARRATIVE

To explore the dynamic regulatory role of film and television narrative context in the dissemination of sustainable values, this section analyzes the evolution of visual-semantic association patterns from two dimensions: temporal evolution and genre differences. By calculating the changing trends of topic intensity in public discussions at different times, and the intensity distribution of key association pairs in different film and television genres, this section reveals the mechanism by which narrative frameworks shape the interpretation of symbolic meaning. The results of the temporal evolution analysis and cross-genre comparison are shown in Figure 5.

Figure 5a, with time intervals on the horizontal axis and the percentage of theme intensity on the vertical axis, shows the evolution of the three major sustainability themes. The durability and classic theme intensity declined steadily from 28.5% in 2000-2005 to 18.2% in 2021-2023, reflecting the weakening influence of traditional durability narratives in contemporary film and television. This change is closely related to the rise of fast fashion culture and the diversification of audience aesthetics. The technology and future theme demonstrated an opposite trend, as it experienced a steady increase from 15.2% to 35.8%, which represents the growing importance of the development of digital technology and innovation in sustainable materials, in the public sphere. The ethical consumption theme also represents an important developmental trend. Although ethical consumption began at a lower baseline, it experienced significant growth, going from 8.8% to 22.5%, indicating a growing depth of penetration of ethical consumption representation in film and television narratives. Figure 5b also demonstrates the moderating role of narrative genre on the degree of associations. The association between technologically advanced fabrics

and the future technology peaked at 0.61 in science fiction. This occurs because the technological reliance of science fiction narratives is based more on imaginary futures. The association between distressed finishes and durable classics was found to be strongest (0.51) for drama films, which aligns with the aesthetic pursuit of realism and explored narrative content typically identified with the genre. The associations presented to the public audience in the action and fantasy genres are given moderate strength (0.28-0.45), which demonstrates a level of complexity around the interpretation of these symbols in light of the trend toward genre hybridization. This variable degree of strength suggests that film and television genres, as distinct narrative frameworks, are consistently associated with differentiated interpretation pathways and cognitive foci regarding the sustainable value of leather clothing, through a preset visual grammar that is consistent with their genre expectations.

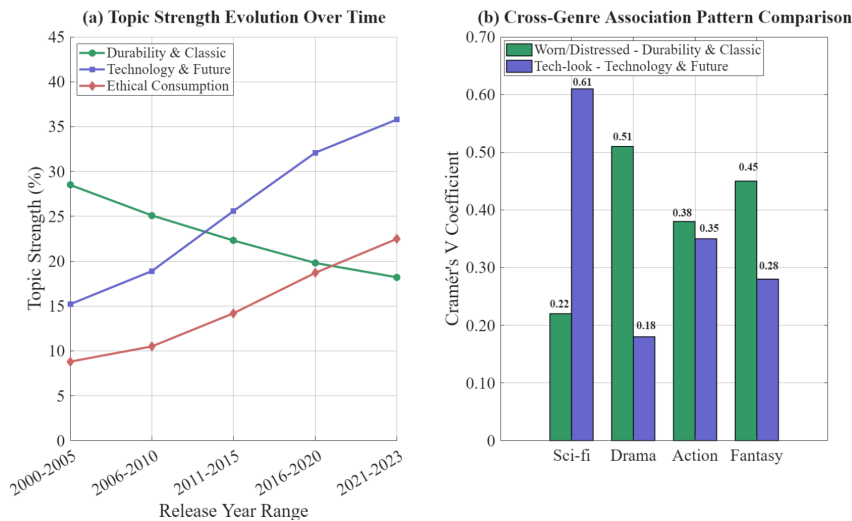


FIGURE 5. Thematic temporal evolution and cross-type comparison. (a) Evolution of thematic intensity over time; (b) Comparison of cross-type association patterns

V. CONCLUSION

This research develops a cross-modal association framework that combines visual content analysis with BERTopic dynamic theme modeling to facilitate the first quantitative analysis of the mechanism of sustainable value communication within film and television leather apparel design. This methodology can establish a visual-semantic mapping through the systematic encoding of visual characteristics, the mining of textual semantic themes, and the calculation of the Clemens V coefficient that can elucidate the underlying connection to which design visuals symbolize and the view of the public. In experiments, the association identified a strong consensus correlation of 0.45 between distressed finishes and the durable classic theme and a Clemens V coefficient of 0.52 for technologically textured fabrics and the technological future theme. Furthermore, the intensity of the technological future theme has increased from 15.2% of all content in the early phase to 35.8% most recently, generating evidence of the technological pivot in sustainable narratives. This study could contribute to the establishment of a quantifiable analytical paradigm for sustainable fashion communication. Its methodological system can accurately identify the implicit environmental value orientation in film and television media, providing a theoretical basis and practical tools for narrative strategy formulation and communication effect evaluation in the field of apparel design.

AUTHOR CONTRIBUTIONS

Xiaoyue Zhang designed, collected and analyzed the data, and drafted the manuscript. Xiaoyue Zhang conducted the study, critically revised the manuscript for important intellectual content, and gave final approval of the version to be published. Xiaoyue Zhang participated fully in the work, take public responsibility for appropriate portions of the content, and agreed to be accountable for all aspects of the work in ensuring that questions related to the accuracy

or integrity of any part of the work are appropriately investigated and resolved.

CONFLICTS OF INTEREST

The author declares no conflict of interest.

FUNDING

This research received no external funding.

ACKNOWLEDGEMENTS

Not applicable.

REFERENCES

- [1] S. G. Alzahrani and S. A. Saroukh, "The Semiotic Dimension of Men's Fashion in Modern Eras," *International Journal*, vol. 12, no. 5, pp. 133-148, 2024, doi: 10.11648/j.ijla.20241205.13.
- [2] D. Burato, "Audiovisual Narrative for Fashion Heritage: Communicating and Valorizing Archives through Fashion Film," *ZoneModa Journal*, vol. 14, no. 2, pp. 77-93, 2024, doi: 10.6092/issn.2611-0563/20692.
- [3] I. D'Adamo, M. Gagliarducci, M. Iannilli, and V. Mangani, "Fashion Wears Sustainable Leather: A Social and Strategic Analysis Toward Sustainable Production and Consumption Goals," *Sustainability*, vol. 16, no. 22, pp. 9971, 2024, doi: 10.3390/su16229971.
- [4] P. Díaz Soloaga, G. Muñoz Domínguez, and A. G. Woodside, "Fashion and film stories of (mis) understanding: Introduction to a special issue on cinema and fashion," *Journal of Global Fashion Marketing*, vol. 14, no. 4, pp. 369-373, 2023, doi: 10.1080/20932685.2023.2206411.
- [5] N. Y. Kil and J. Chun, "Semiotic Analysis of Advertising Video Related to the Sustainability of Fast Fashion Brands," *Journal of the Korean Society of Clothing and Textiles*, vol. 47, no. 6, pp. 1057-1079, 2023, doi: 10.5850/JKST.2023.47.6.1057.
- [6] E. Ertürk, "Consumer society and sustainability: A semiotic analysis of campaigns on ecological fashion thought," *Ordu Üniversitesi Sosyal Bilimler Enstitüsü Sosyal Bilimler Araştırmaları Dergisi*, vol. 15, no. 1, pp. 227-255, 2025, doi: 10.48146/odusobiad.1424243.
- [7] F. Yin, "Analysis of the Role of Costume Design in Shaping the Characters of Film and Television," *Frontiers in Art Research*, vol. 5, no. 6, pp. 58-63, 2023, doi: 10.25236/FAR.2023.050611.
- [8] J. Clemens-Smucker, "The Clothes Make the Woman: How Fashion Informs the Comedic Identity of Schil's Creek's Moira Rose," *Journal of Popular Film and Television*, vol. 51, no. 4, pp. 174-183, 2023, doi: 10.1080/01956051.2023.2275011.
- [9] A. Chandrasekar, S. E. Clark, S. Marn, S. Vandersloot, E. C. Flores, D. Aceituno, et al., "Making the most of big qualitative datasets: a living

- systematic review of analysis methods,” *Frontiers in big Data*, vol. 7, Art. no. 1455399, 2024, doi: 10.3389/fdata.2024.1455399.
- [10] A. Malandrino, “Comparing qualitative and quantitative text analysis methods in combination with document-based social network analysis to understand policy networks,” *Quality & Quantity*, vol. 58, no. 3, pp. 2543-2570, 2024, doi: 10.1007/s11135-023-01753-1.
- [11] R. Murakami and B. Chakraborty, “Investigating the efficient use of word embedding with neural-topic models for interpretable topics from short texts,” *Sensors*, vol. 22, no. 3, pp. 852, 2022, doi: 10.3390/s22030852.
- [12] F. Ravenda, S. A. Bahrainian, A. Raballo, A. Mira, and F. Crestani, “A self-supervised seed-driven approach to topic modeling and clustering,” *Journal of Intelligent Information Systems*, vol. 63, no. 1, pp. 333-353, 2025, doi: 10.1007/s10844-024-00891-8.
- [13] N. Moratelli, M. Barraco, D. Morelli, M. Cornia, L. Baraldi, and R. Cucchiara, “Fashion-oriented image captioning with external knowledge retrieval and fully attentional gates,” *Sensors*, vol. 23, no. 3, pp. 1286, 2023, doi: 10.3390/s23031286.
- [14] G. Moro, S. Salvatori, and G. Frisoni, “Efficient text-image semantic search: A multi-modal vision-language approach for fashion retrieval,” *Neurocomputing*, vol. 538, Art. no. 126196, 2023, doi: 10.1016/j.neucom.2023.03.057.
- [15] S. S. Gurao and M. Parl, “Deconstruct to Reconstruct: Turning yesterday’s Clothes into Today’s Fashion,” *Journal of Computational Analysis & Applications*, vol. 33, no. 8, pp. 989, 2024.
- [16] J. Jin, S. A. Sharudin, and T. Wang, “Visual Design of Chinese Animated Films Infused with Oriental Punk Aesthetics,” *Journal of Theory and Practice of Social Science*, vol. 4, no. 04, pp. 1-8, 2024, doi: 10.53469/jtpss.2024.04(04).01.
- [17] M. Li, Q. Gao, and T. Yu, “Kappa statistic considerations in evaluating inter-rater reliability between two raters: which, when and context matters,” *BMC cancer*, vol. 23, no. 1, pp. 799, 2023, doi: 10.1186/s12885-023-11325-z.
- [18] G. Rau and Y. S. Shih, “Evaluation of Cohen’s kappa and other measures of inter-rater agreement for genre analysis and other nominal data,” *Journal of english for academic purposes*, vol. 53, Art. no. 101026, 2021, doi: 10.1016/j.jeap.2021.101026.
- [19] N. Khodair and F. Elghannam, “Efficient topic identification for urgent MOOC Forum posts using BERTopic and traditional topic modeling techniques,” *Educational and Information Technologies*, vol. 30, no. 5, pp. 5501-5527, 2025, doi: 10.1007/s10639-024-13003-4.
- [20] A. Rejeb, K. Rejeb, S. Simske, S. Süle, and E., “Industry 5.0 research: an approach using co-word analysis and BERTopic modeling,” *Discover Sustainability*, vol. 6, no. 1, pp. 402, 2025, doi: 10.1007/s43621-025-01252-3.
- [21] M. Herrmann, D. Kazempour, F. Scheipl, Kröger, and P., “Enhancing cluster analysis via topological manifold learning,” *Data Mining and Knowledge Discovery*, vol. 38, no. 3, pp. 840-887, 2024, doi: 10.1007/s10618-023-00980-2.
- [22] W. Yi, S. Bu, H. H. Lee, and C. H. Chan, “Comparative analysis of manifold learning-based dimension reduction methods: a mathematical perspective,” *Mathematics*, vol. 12, no. 15, pp. 2388, 2024, doi: 10.3390/math12152388.
- [23] G. Stewart and M. Al-Khassawneh, “An implementation of the HDBSCAN* clustering algorithm,” *Applied Sciences*, vol. 12, no. 5, pp. 2405, 2022, doi: 10.3390/app12052405.
- [24] S. M. Wang, W. R. Yang, Q. Y. Zhuang, W. H. Lin, M. Y. Tian, T. J. Su, et al., “Application of Three-Dimensional Hierarchical Density-Based Spatial Clustering of Applications with Noise in Ship Automatic Identification System Trajectory-Cluster Analysis,” *Applied Sciences*, vol. 15, no. 5, pp. 2621, 2025, doi: 10.3390/app15052621.
- [25] M. Grootendorst, “BERTopic: Neural topic modeling with a class-based TF-IDF procedure,” preprint arXiv, vol. 22, no. 03, Art. no. 05794, 2022, doi: 10.48550/arXiv.2203.05794.
- [26] W. Mu, K. H. Lim, J. Liu, S. Karunasekera, L. Falzon, and A. Harwood, “A clustering-based topic model using word networks and word embeddings,” *Journal of big data*, vol. 9, no. 1, pp. 38, 2022, doi: 10.1186/s40537-022-00585-4.
- [27] C. Tang, “Review on Application of Chi-square Statistics in Text Classification in Recent Five Years,” *Applied and Computational Engineering*, vol. 97, no. 1, pp. 115-118, 2024, doi: 10.54254/2755-2721/97/20241397.
- [28] D. C. Sami’Un, A. Sugiharto, and F. Jie, “Chi Square Feature Selection For Improving Sentiment Analysis of News Data Privacy Treats,” *Journal of Theoretical and Applied Information Technology*, vol. 102, no. 18, pp. 6601-6610, 2024.
- [29] W. Urasaki, T. Nakagawa, T. Momozaki, and S. Tomizawa, “Generalized Cramér’s coefficient via f-divergence for contingency tables,” *Advances in Data Analysis and Classification*, vol. 18, no. 4, pp. 893-910, 2024, doi: 10.1007/s11634-023-00560-8.
- [30] R. L. Sapra and S. Saluja, “Understanding statistical association and correlation,” *Current Medicine Research and Practice*, vol. 11, no. 1, pp. 31-38, 2021, doi: 10.4103/cmpr.cmpr_62_20.
- [31] C. Galli, N. Donos, and E. Calciolari, “Performance of 4 pre-trained sentence transformer models in the semantic query of a systematic review dataset on peri-implantitis,” *Informatics*, vol. 15, no. 2, pp. 68, 2024, doi: 10.3390/info15020068.
- [32] G. B. Guedes and A. E. A. da Silva, “Classification and clustering of sentence-level embeddings of scientific articles generated by contrastive learning,” preprint arXiv, vol. 24, no. 04, Art. no. 00224, 2024, doi: 10.48550/arXiv.2404.00224.
- [33] U. Chauhan and A. Shah, “Topic modeling using latent Dirichlet allocation: A survey,” *ACM Computing Surveys (CSUR)*, vol. 54, no. 7, pp. 1-35, 2021, doi: 10.1145/3462478.
- [34] J. Gan, T. Liu, L. Li, and J. Zhang, “Non-negative matrix factorization: a survey,” *The Computer Journal*, vol. 64, no. 7, pp. 1080-1092, 2021, doi: 10.1093/comjnl/bxab103.
- [35] G. Cheng, Q. You, L. Shi, Z. Wang, J. Luo, and T. Li, “A survey of topic models: From a whole-cycle perspective,” *Journal of Intelligent & Fuzzy Systems*, vol. 45, no. 6, pp. 9929-9953, 2023, doi: 10.3233/JIFS-233551.