

Explainable Artificial Intelligence for Social Media Sentiment Mining: Integrating Transformer-Based Attention and Human-Centric Interpretation

Giorgio Marsh

Department of Computer Science and Engineering, University of Nevada, Reno, Reno, NV, USA.

giorgiomarsh@unr.edu

Mukesh C. Basu

Department of Electrical Engineering and Computer Science, University of Missouri, Columbia, MO, USA.

mukesh.c.basu@missouri.edu

Karan Seaingh

School of Computing, Clemson University, Clemson, SC, USA.

karans@clemson.edu

Abstract

The proliferation of user-generated content on social media platforms has created enormous demand for automated sentiment mining tools that can accurately capture nuanced opinion, emotion, and stance at scale. While transformer-based deep learning architectures have achieved state-of-the-art performance in text sentiment classification, their inherent opacity presents a critical barrier to adoption in high-stakes domains such as public health surveillance, financial market analysis, and political discourse monitoring. This paper presents a comprehensive interdisciplinary examination of explainable artificial intelligence (XAI) approaches tailored to social media sentiment mining, focusing on the integration of transformer-based attention mechanisms with human-centric interpretation frameworks. We survey contemporary architectural paradigms including bidirectional encoder representations, generative pre-trained transformers, and their fine-tuning strategies, then systematically analyze how intrinsic attention weights can be repurposed to generate post hoc explanations. The discussion extends to the structural tension between explanation fidelity and model complexity, the infrastructure requirements for deploying interpretable models in real-time social media pipelines, and the governance frameworks necessary to ensure fairness and accountability. By examining cross-domain case illustrations from healthcare misinformation detection, crisis response coordination, and consumer sentiment tracking, we identify key trade-offs in latency, transparency, and robustness. The paper further addresses sustainability concerns related to the computational carbon footprint of large-scale transformer fine-tuning and inference, proposing lightweight distillation and modular explanation proxies as viable paths forward. Policy implications concerning algorithmic transparency mandates and auditability standards are synthesized with technical design choices, culminating in a set of architectural principles for building socially responsible sentiment mining systems. The analysis reveals that effective explainability cannot be reduced to a single technical method but must be constructed as a socio-technical assemblage where attention visualization,

contrastive example generation, and structured natural language rationales are combined under participatory design with domain stakeholders.

Keywords

explainable artificial intelligence, sentiment mining, social media, transformer attention, human-centric interpretation, algorithmic fairness, socio-technical systems.

1. Introduction

The digital public sphere has undergone a profound transformation with the ascendancy of social media platforms as primary arenas for opinion formation, emotional expression, and collective sense-making. Every day, billions of textual artifacts, ranging from microblog posts to threaded discussions and multimedia captions, are generated across geographically and culturally diverse user populations. The automatic extraction of sentiment from this data has evolved from simple lexicon-based polarity detection into a sophisticated computational enterprise that aims to capture fine-grained affective states, sarcasm, contextual irony, and culturally coded emotional expressions. Transformer-based deep learning architectures, particularly those leveraging multi-head self-attention mechanisms, have dramatically improved the accuracy and generalizability of sentiment classifiers across languages and domains. However, the very architectural depth and parametric complexity that endow these models with superior performance simultaneously render them profoundly opaque to human understanding, creating an acute tension between predictive capability and interpretability.

This opacity crisis manifests acutely in application scenarios where sentiment mining outputs inform consequential decisions. Public health agencies that monitor vaccine sentiment on social media to design communication interventions require not only accurate trend detection but also a clear understanding of the linguistic and contextual features that drive model predictions, lest they misallocate resources based on spurious correlations. Financial regulators scrutinizing market manipulation through coordinated sentiment campaigns demand audit trails that can be presented to legal authorities and that satisfy due process requirements. Electoral oversight bodies tracking hate speech and inflammatory polarization need to differentiate between genuine grassroots sentiment shifts and coordinated inauthentic behavior, a task that demands transparent reasoning chains. In each of these domains, the inability to explain why a particular piece of social media text was classified as expressing anger rather than fear, or irony rather than sincerity, undermines trust, limits accountability, and erodes the practical utility of even the most accurate automated systems.

The emerging field of explainable artificial intelligence (XAI) has responded to this challenge by developing a diverse repertoire of techniques for rendering black-box models interpretable. These range from intrinsically interpretable proxy models like linear classifiers with sparse feature weights to post hoc explanation methods that perturb inputs and measure output sensitivity, to attention weight visualization that maps the internal focus of transformer layers back onto input tokens. The integration of XAI into social media sentiment mining is particularly demanding because the linguistic phenomena involved are deeply contextual, culturally situated, and often deliberately ambiguous. Sarcasm detection, for instance, frequently hinges on subtle incongruities between a statement and its surrounding conversational context, a feature that may be distributed across multiple attention heads and layers of a deep network. Explaining such predictions in human-understandable terms therefore requires not merely surfacing which words contributed most to the output but

constructing coherent narratives that align with human reasoning patterns about intention, tone, and social pragmatics.

This paper undertakes a systematic, interdisciplinary examination of the design space for XAI systems in social media sentiment mining, with a particular emphasis on the intersection of transformer-based attention mechanisms and human-centric interpretation. We adopt a systems-level perspective, recognizing that explainability is not an isolated property of a model but an emergent characteristic of the broader socio-technical assemblage encompassing data preprocessing pipelines, model architecture, deployment infrastructure, user interface design, and governance oversight. Our analysis traverses multiple layers of abstraction, from the internal dynamics of attention heads within fine-tuned language models to the institutional policy frameworks that increasingly mandate algorithmic transparency. By drawing on cross-domain case illustrations and synthesizing insights from natural language processing, human-computer interaction, critical algorithm studies, and infrastructure engineering, we articulate a set of architectural principles and structural trade-offs that can guide the design of responsible sentiment mining systems capable of serving both technical performance goals and societal legitimacy requirements.

2. Background and Related Work

The trajectory of sentiment analysis research mirrors the broader evolution of natural language processing from rule-based expert systems through statistical feature engineering to contemporary deep learning paradigms. Early approaches relied on curated sentiment lexicons and hand-crafted syntactic patterns, yielding systems that were inherently interpretable because their decision logic could be directly inspected. Lexicon-based classifiers produced explanations by simply listing the words that matched entries in positive or negative dictionaries, providing a form of transparency that, while often inaccurate, was intuitively satisfying. The era of supervised machine learning introduced support vector machines and naive Bayes classifiers trained on bag-of-words representations, where feature weights could be examined to identify discriminating terms, still maintaining a degree of interpretability through linear model coefficients. The decisive rupture in the interpretability-performance trade-off occurred with the introduction of deep recurrent networks and, subsequently, transformer architectures equipped with self-attention mechanisms [1, 2].

The transformer architecture, originally proposed for machine translation, rapidly emerged as the foundation for pre-trained language models such as BERT, RoBERTa, and GPT variants that could be fine-tuned on downstream sentiment classification tasks with remarkable accuracy. These models process text through multiple layers of multi-head self-attention, where each token attends to every other token in the input sequence through learned query, key, and value projections. The resulting contextualized representations capture rich semantic and syntactic information, enabling the models to disambiguate polysemous words, resolve long-range dependencies, and detect subtle pragmatic cues that are essential for sentiment understanding in social media text, which is often characterized by non-standard orthography, code-switching, and emoji usage. However, the very density and interactivity of these internal representations make it extremely challenging to trace how a particular prediction emerges from the input, as each token's final embedding is a complex non-linear aggregation of information from all other tokens across multiple representational subspaces [3].

Researchers in XAI have developed several families of methods to address this opacity, each with distinct trade-offs in terms of fidelity, computational cost, and cognitive compatibility with human reasoning. Feature attribution methods, exemplified by LIME and SHAP,

approximate the local decision boundary of a black-box model using simpler surrogate models, generating importance scores for input features that can be visualized as heatmaps or word importance lists. These post hoc techniques are model-agnostic, meaning they can be applied to any classifier without requiring access to internal architecture, a property that facilitates their integration into modular sentiment analysis pipelines where the underlying model may be updated or replaced. Integrated gradients and DeepLIFT propagate prediction differences back to the input layer through path integrals over gradients, providing more theoretically grounded attributions that satisfy axioms such as sensitivity and implementation invariance [4]. However, the computational overhead of these methods, particularly the need for multiple forward passes or gradient computations per explanation, can become prohibitive in real-time social media monitoring systems that must process high-velocity data streams with low latency requirements.

Attention-based explanation has attracted particular interest because it leverages the very mechanism that powers transformer performance, potentially providing explanations with zero additional forward pass cost beyond what is already computed during inference. Early studies demonstrated that attention weights from recurrent sequence-to-sequence models correlate with human judgments of word importance in tasks such as machine translation and image captioning. In the context of fine-tuned transformer classifiers, attention weights from the final layers can be aggregated across heads and layers to produce token-level heatmaps that highlight the parts of the input text most influential for a given sentiment prediction [5]. The intuitive appeal of these visualizations is substantial: they can be directly overlaid on the original social media post, allowing users to see which phrases the model focused on when reaching its conclusion. However, a growing body of critical scholarship has raised important caveats, demonstrating that attention weights do not always reliably indicate feature importance in the sense of causal contribution and can be manipulated without changing model predictions, and that alternative attention distributions can yield the same output, undermining the uniqueness of attention-based explanations [6, 7].

3. Transformer-Based Architectures for Sentiment Analysis

The current generation of transformer-based sentiment classifiers represents a family of architectures that share the core self-attention mechanism but diverge in pre-training objectives, model scale, and fine-tuning strategies. Bidirectional encoder models like BERT and its domain-adapted variants pre-train on large corpora using masked language modeling and next-sentence prediction tasks, then are fine-tuned by appending a classification head and training on labeled sentiment data. The bidirectional context is particularly valuable for social media sentiment because the meaning of a given word in a post often depends on both its preceding and following context, especially in the case of negation constructions, contrastive conjunctions, and sarcastic reversals where the sentiment polarity of a statement is inverted by subsequent clauses. RoBERTa optimizes the pre-training recipe by removing the next-sentence prediction objective, training on larger batches, and dynamically changing the masking pattern, yielding improvements on several sentiment benchmarks [8]. Generative pre-trained transformer models such as GPT-3 and its successors adopt a left-to-right autoregressive architecture that excels at generating coherent continuations and can be adapted for sentiment classification through careful prompting or by extracting embeddings from the final hidden states for use in downstream classifiers. The trade-off between encoder-only and decoder-only architectures involves competing considerations of bidirectional

context integration versus generative capability, the latter being potentially valuable for generating natural language explanations alongside classification decisions.

The fine-tuning process itself introduces significant structural considerations for explainability. When a pre-trained language model is fine-tuned on a specific sentiment dataset, the internal representations undergo substantial reorganization as the model adapts its attention patterns to the particular lexical, syntactic, and stylistic features of the target domain. Social media domains such as Twitter, Reddit, or Weibo exhibit distinctive linguistic properties including heavy use of hashtags, user mentions, URL tokens, and platform-specific conventions that may not be well represented in the general-domain pre-training corpora. Fine-tuning on these domains causes certain attention heads to specialize in tracking sentiment-bearing terms and their modifiers, while others may focus on structural tokens that serve as proxy indicators of sentiment, such as the presence of certain emojis or the density of capitalization. Understanding this specialization is critical for designing explanation methods that can surface what the model has actually learned, as opposed to what its designers assume it has learned [9].

A particularly promising line of research explores dynamic adaptive attention mechanisms that can modulate the contribution of different attention heads based on input characteristics, rather than using static aggregation schemes [10]. Such approaches allow the explanation system to highlight the specific subspaces that were most actively engaged in processing a given input, providing a more fine-grained and faithful account of the model’s internal decision process. The structural challenge lies in designing the adaptation controller such that it does not itself become a source of opacity or introduce coupling that complicates model updates. Architectures that separate the explanation pathway from the prediction pathway, for instance by training an auxiliary explanation model to predict attention distributions from hidden states without affecting the forward pass of the main classifier, offer a modular design pattern that supports independent evolution of performance and interpretability components.

4. Explainability Mechanisms in Deep Learning Models

The landscape of XAI methods applicable to transformer-based sentiment classifiers can be systematically analyzed along several axes: the distinction between intrinsic and post hoc explainability, the granularity of explanation from token-level to concept-level to example-level, and the format of the explanation output ranging from visual heatmaps to natural language rationales to counterfactual instances. Intrinsically interpretable models embed explanation into their structure by design; for example, a model architecture that explicitly separates sentiment-bearing content words from function words through constrained attention masks can produce inherently transparent reasoning traces. However, the performance of such architectures on complex social media sentiment tasks often lags behind that of unconstrained deep models, reinforcing the performance-transparency trade-off that motivates much post hoc XAI research [11].

Post hoc explanation methods for transformer models can be broadly categorized into gradient-based, perturbation-based, and attention-based approaches. Gradient-based methods such as integrated gradients and SmoothGrad compute the gradient of the output with respect to each input token embedding and propagate this signal back through the network layers, producing saliency maps that highlight tokens with high sensitivity. These methods have the advantage of being grounded in the functional relationship between inputs and outputs, but they can be noisy, especially for inputs with many tokens where gradients become vanishingly small, and they require careful choice of baseline inputs for the integral path.

Perturbation-based methods such as LIME and SHAP systematically remove or replace input tokens and measure the resulting change in prediction probability, constructing a local linear approximation of the decision boundary. While these methods are model-agnostic and flexible, their sampling-based nature introduces variance into the explanations, and the choice of perturbation strategy (e.g., masking with zero vectors versus replacing with out-of-distribution tokens) can significantly influence the resulting attribution scores. For social media text containing informal language, emojis, and platform-specific tokens, perturbation with out-of-vocabulary substitutes risks generating unrealistic input variants that lie outside the model’s training distribution, leading to unreliable explanations [12].

Attention-based explanation inherits the architectural affinity with the underlying model, as the attention weights are already computed as part of the standard inference pass. The typical workflow for producing an attention-based explanation from a fine-tuned transformer involves extracting the self-attention matrices from one or more layers, often focusing on the final layer or computing a weighted average across layers and heads, and then mapping the attention scores from the classification token or the pooled output back to the input tokens. A common approach is to extract the attention vector for the special classification token in the final layer, which aggregates information from all input tokens, and use this vector as token importance scores. More sophisticated methods consider attention flow across layers, accounting for the fact that attention in deeper layers operates on already contextualized representations that themselves incorporate attention from earlier layers, a phenomenon known as attention rollout or attention flow [13]. The visualization of these attention maps as color-coded overlays on the original social media post provides an immediate and intuitive interface for end users, aligning well with the requirements of human-centric interpretation.

Despite their intuitive appeal, attention-based explanations face fundamental challenges that have sparked vigorous debate in the XAI community. Recent work has demonstrated that attention weights can be adversarial manipulated without changing model predictions, that different attention distributions can produce identical outputs, and that attention often correlates poorly with gradient-based importance measures that have stronger causal grounding [6]. These findings do not necessarily invalidate attention-based explanation but rather underscore the need for careful validation against human judgments and task-specific benchmarks. In the context of social media sentiment, where the phenomena of interest are inherently subjective and culturally variable, the evaluation of explanation quality cannot rely solely on computational metrics but must engage with the interpretive frameworks and domain expertise of the human stakeholders who will ultimately act on the explanations.

5. Human-Centric Interpretation Frameworks

The shift from model-centric to human-centric XAI represents a critical evolution in the design of interpretable sentiment mining systems. Human-centric interpretation acknowledges that the ultimate measure of an explanation’s quality is not its fidelity to the model’s internal computations but its effectiveness in supporting human decision-making, cultivating appropriate trust, enabling error detection, and satisfying institutional accountability requirements. This perspective draws on insights from cognitive psychology, human-computer interaction, and science and technology studies to characterize the cognitive processes by which human users construct mental models of automated systems and the social contexts in which explanations are requested, provided, and contested [14].

A foundational concept in human-centric XAI is the distinction between explanation as a product and explanation as a process. Product-oriented approaches deliver a static attribution

map or a generated rationale to the user, treating the explanation as a fixed artifact. Process-oriented approaches, in contrast, support interactive exploration where users can probe the model's behavior by modifying inputs, querying counterfactual scenarios, and receiving iterative feedback that helps them refine their understanding. For social media sentiment mining applications, process-oriented explanation is particularly valuable because the interpretation of a single post is often ambiguous without considering the conversational thread, the user's historical posting behavior, and the broader discourse context. An interactive explanation interface might allow a public health analyst to click on a tweet classified as expressing vaccine hesitancy and see not only which tokens contributed to that classification but also how the classification would change if certain phrases were removed or replaced, and how the tweet's sentiment relates to the sentiment of temporally or topically adjacent tweets [15].

The design of human-centric explanation systems for sentiment mining must also contend with the diversity of stakeholders, each possessing distinct expertise, goals, and cognitive styles. A data scientist responsible for model debugging may require detailed attention head visualizations and statistical summaries of feature importance across a validation corpus, while a public health communication officer may need plain-language summaries of prevailing sentiment themes and specific exemplar posts that illustrate each theme, and a legal auditor may need a formal, reproducible explanation that can be documented in compliance reports. Accommodating this diversity necessitates modular explanation architectures where the core model can generate rich internal representations that are then rendered into different output modalities and abstraction levels by separate interpretation layers. This separation of concerns also facilitates independent evolution of the model and the explanation interface, allowing domain-specific customization without retraining the underlying classifier [16].

Natural language explanation generation represents a frontier of human-centric XAI for sentiment mining. Rather than presenting numerical importance scores that require cognitive translation by the user, generative explanation models produce coherent textual rationales that articulate the reasoning behind a classification in human language. For instance, instead of outputting a heatmap highlighting certain words, the system might generate the statement: "This post was classified as expressing anger because the author uses accusatory language directed at a specific public figure and employs intensifying adverbs associated with high arousal negative affect." The generation of such rationales can be achieved by fine-tuning large language models on datasets that pair posts with human-written explanations, or by using chain-of-thought prompting to elicit step-by-step reasoning from generative models [17]. The challenge lies in ensuring that the generated rationales are faithful to the actual decision process of the classifier, as it is well documented that generative models can produce plausible-sounding but confabulated explanations that do not accurately reflect the basis of their predictions. Addressing this challenge requires careful architectural coupling between the classifier and the explanation generator, possibly through shared latent representations or through consistency training that penalizes discrepancies between attention-based importance and generated rationale content.

6. System Integration and Infrastructure Design

Deploying explainable sentiment mining in real-world settings demands careful attention to system integration and infrastructure design, as the addition of explanation generation capabilities can substantially increase computational load, storage requirements, and data flow complexity. A typical production sentiment mining pipeline for social media involves

continuous data ingestion from platform APIs or firehose streams, real-time text preprocessing and normalization, batched inference through the fine-tuned transformer model, and indexing of predictions and associated metadata for downstream analytics and visualization. Integrating XAI into this pipeline requires additional stages for explanation computation, explanation storage, and explanation serving, each of which must be designed for scalability, fault tolerance, and cost efficiency [18].

The latency implications of different XAI methods are a primary infrastructure consideration. Attention-based explanations, because they reuse internal states already computed during the forward pass, add minimal incremental latency, typically only the cost of extracting and serializing attention weights from the model output tensors. This makes them well suited for real-time monitoring applications where users expect near-instantaneous feedback, such as live dashboards tracking sentiment shifts during a televised political debate or a product launch event. Gradient-based and perturbation-based methods, by contrast, require multiple additional forward or backward passes for each explained prediction, multiplying inference cost by factors ranging from tens to hundreds. These methods may be deployed in asynchronous explanation pipelines that operate on a sampled subset of predictions or that serve on-demand explanation requests from users who click on specific posts for deeper investigation. The architectural pattern of a two-tier explanation system, with lightweight attention-based explanations computed in the hot path and deeper attribution analyses computed in a separate cold path, balances the competing demands of responsiveness and analytical depth [19].

Storage and indexing of explanations present further challenges at scale. A high-throughput sentiment monitoring system processing millions of posts per day would generate an enormous volume of explanation artifacts if full attention tensors were stored for every prediction. Dimensionality reduction techniques such as storing only the top-k attended tokens per post or computing aggregate attention statistics across time windows can substantially reduce storage footprints while preserving sufficient information for trend analysis and spot auditing. The choice of explanation persistence strategy must also account for data retention policies and privacy regulations, particularly in jurisdictions where social media data may contain personally identifiable information that is subject to deletion requests. Explanation storage systems that maintain referential integrity with the underlying posts while allowing for selective purging are necessary to comply with evolving data governance requirements [20].

The hardware and energy infrastructure underlying large-scale transformer inference is increasingly subject to sustainability scrutiny. Fine-tuning and running large language models for sentiment mining across billions of social media posts consumes significant electrical energy, contributing to the carbon footprint of AI operations. Research into model distillation, quantization, and pruning has demonstrated that substantially smaller models can retain competitive sentiment classification accuracy while dramatically reducing inference cost, making it feasible to deploy explainable sentiment mining on edge devices or within carbon-constrained data centers [21]. The integration of XAI into sustainable AI strategies involves an additional consideration: the explanation generation process itself must not reverse the efficiency gains achieved through model compression. This motivates the development of lightweight explanation proxies, where a compact student model is trained to generate explanations that closely approximate those of a larger teacher model, enabling deployment of

interpretable sentiment mining in resource-limited settings such as community health monitoring in low-connectivity regions [22].

7. Robustness, Fairness, and Ethical Governance

The integration of XAI into social media sentiment mining intersects critically with concerns about model robustness, fairness, and ethical governance. Social media text exhibits substantial linguistic variation across demographic groups, geographies, and subcultures, and models trained predominantly on data from majority populations may systematically underperform on and misrepresent the sentiment of minority communities. Explainability mechanisms can serve as diagnostic tools for detecting such disparities, enabling developers and auditors to inspect whether the features driving sentiment predictions differ across demographic subgroups in ways that reflect genuine linguistic patterns or unwanted biases. For example, if an attention-based explanation reveals that a classifier consistently assigns negative sentiment to posts containing African American Vernacular English features regardless of actual sentiment content, this signals a fairness violation that must be remediated through data augmentation, adversarial debiasing, or model retraining [23].

The relationship between explainability and fairness is, however, not straightforward. Explanations can themselves encode and perpetuate biases if the explanation generation process is trained on data that contains stereotypical associations. Moreover, the provision of explanations may create an illusion of neutrality that masks underlying discriminatory behaviors, a phenomenon known as explanation laundering, where the superficial transparency of an attention heatmap obscures deeper structural biases embedded in the training data, the annotation guidelines, or the problem formulation itself. Mitigating these risks requires a multi-layered governance approach in which XAI outputs are subject to regular fairness audits by interdisciplinary teams that include domain experts, affected community representatives, and ethicists, not solely by the engineers who developed the models [24].

Robustness concerns extend to the susceptibility of both sentiment classifiers and their associated explanations to adversarial manipulation. Social media platforms are rife with actors who seek to game algorithmic systems for commercial, political, or malicious purposes, and sentiment mining tools are attractive targets for such manipulation. Adversarial attacks on text classifiers, which involve introducing perturbations such as synonym substitutions, character-level misspellings, or invisible Unicode characters that change model predictions while preserving human-perceived meaning, have been extensively documented. Recent research demonstrates that adversarial perturbations can also be designed to manipulate explanations while leaving predictions unchanged, a particularly insidious attack vector that could mislead human overseers into trusting a compromised system. Defending against such attacks requires adversarial training regimes that incorporate explanation consistency into the loss function and robust explanation methods that are less sensitive to input perturbations [25].

Ethical governance frameworks for explainable sentiment mining must also address the broader societal implications of deploying these technologies at scale. The capacity to monitor and interpret public sentiment in real time confers significant power on the entities that control the monitoring infrastructure, raising concerns about surveillance, manipulation of public opinion, and chilling effects on legitimate expression. Policy interventions such as the European Union’s AI Act and the Algorithmic Accountability Act proposed in various jurisdictions seek to mandate transparency, risk assessment, and human oversight for high-risk AI systems, categories that arguably encompass large-scale sentiment mining used in

electoral, public health, and financial contexts. Compliance with these regulatory regimes requires that XAI mechanisms be integrated not as afterthoughts but as foundational design requirements, with documented audit trails, standardized explanation formats, and mechanisms for contestation by affected individuals and communities.

8. Deployment and Sustainability Considerations

The deployment lifecycle of an explainable sentiment mining system encompasses model selection, training, validation, integration, monitoring, and continuous updating, each phase presenting distinct sustainability challenges. The initial training or fine-tuning of a large transformer model on domain-specific social media sentiment data requires substantial computational resources, with corresponding energy consumption and carbon emissions that have attracted increasing concern within the research community and the broader public. The principle of sustainable AI advocates for a holistic accounting of environmental costs throughout the model lifecycle, including not only the one-time training expenditure but also the ongoing inference costs incurred over the operational lifetime of the system, which for a continuously running social media monitoring service may far exceed the initial training cost [21].

Strategies for enhancing the sustainability of explainable sentiment mining include model distillation, where a compact student model is trained to replicate the predictions and, crucially, the explanation patterns of a larger teacher model, and architectural innovations that reduce the quadratic complexity of self-attention to linear or log-linear scaling, making it feasible to process longer documents or larger batches with lower energy per token. Sparse attention mechanisms, which restrict each token to attending only to a subset of other tokens based on locality or learned sparsity patterns, have demonstrated substantial efficiency gains with minimal accuracy degradation, although their impact on the interpretability of attention weights as explanations requires further study, as sparse attention patterns may introduce discontinuities that complicate human interpretation [22].

The social sustainability of sentiment mining systems depends on their long-term maintainability and adaptability to evolving language use. Social media language is characterized by rapid lexical innovation, meme propagation, and platform-specific communicative norms that shift on timescales of weeks or months. A sentiment classifier deployed without continuous monitoring and retraining will experience performance drift as the statistical properties of the input data diverge from the training distribution. Explainability mechanisms can aid in detecting such drift by enabling operators to inspect whether the contextual features driving predictions remain plausible as language evolves. For instance, if an attention-based explanation begins assigning high importance to a newly emerged slang term that the model has not encountered during training, this may signal the need for data refreshment and fine-tuning. Designing modular architectures where the base language model, the classification head, and the explanation generator can be updated independently with different cadences supports agile response to linguistic change while maintaining overall system stability.

The human infrastructure supporting sentiment mining systems, including data annotators, quality assurance personnel, and domain experts who validate explanations, represents another dimension of sustainability that is often overlooked in technical deployments. The labor conditions of crowd workers who produce labeled sentiment data and explanation annotations have been subject to ethical criticism regarding fair compensation, psychological safety when exposed to toxic content, and the precarity of platform-mediated work.

Sustainable system design must incorporate fair labor practices, adequate content moderation support for human annotators, and pathways for the workers whose contributions enable automated sentiment mining to share in the value created by these systems.

9. Policy Implications and Societal Impact

The policy landscape surrounding AI transparency and accountability is evolving rapidly, with direct implications for the design and deployment of explainable sentiment mining systems. The European Union's AI Act establishes a risk-based regulatory framework that classifies AI systems according to their potential for harm, imposing requirements for transparency, human oversight, and technical documentation on high-risk applications. Social media sentiment mining deployed in contexts such as public health surveillance, financial market analysis, or electoral integrity monitoring would likely fall within the high-risk category, triggering obligations to provide meaningful explanations of automated decisions and to maintain auditable records of system operation. Similar legislative initiatives in North America, Asia, and other regions reflect a global convergence toward the principle that algorithmic systems affecting fundamental rights must be explainable and contestable [18].

The operationalization of these policy mandates into concrete technical requirements presents both challenges and opportunities for system architects. A requirement to provide "meaningful" explanations demands more than raw attention weight vectors; it requires interpretation layers that translate machine representations into formats that are accessible to the diverse stakeholders who may need to understand and challenge system outputs. Standardization efforts such as the IEEE P7001 standard for transparency of autonomous systems and the ISO/IEC JTC 1/SC 42 working group on AI trustworthiness are developing frameworks for specifying, measuring, and certifying the explainability properties of AI systems. Sentiment mining platforms that anticipate these standards by incorporating modular, auditable, and standards-compliant explanation interfaces will be better positioned for regulatory approval and public trust [24].

The societal impact of explainable sentiment mining extends beyond regulatory compliance into broader questions of democratic discourse and information ecosystem health. Automated sentiment analysis is increasingly used by news organizations, political campaigns, and advocacy groups to gauge public opinion and tailor messaging strategies. When these tools operate opaquely, they risk amplifying existing power asymmetries by enabling well-resourced actors to optimize their communication for algorithmic amplification while the broader public remains unaware of the mechanisms shaping their information environment. Explainable sentiment mining, when paired with public interest transparency requirements, could serve as a counterbalance by enabling independent researchers, journalists, and civil society organizations to audit the sentiment monitoring systems that influence public discourse. The development of open-source, explainable sentiment mining toolkits, supported by public research funding and maintained by academic consortia, would contribute to a more pluralistic and accountable ecosystem for computational social science [15].

Finally, the integration of XAI into sentiment mining raises epistemological questions about the nature of sentiment as a construct and the authority of automated systems to adjudicate human emotional expression. Sentiment is not a purely objective property of text; its interpretation is culturally mediated, context-dependent, and often contested even among human annotators. Explainable systems that make their interpretive assumptions explicit can foster more critical engagement with sentiment mining outputs, encouraging users to treat automated sentiment classifications as provisional hypotheses rather than definitive truths.

This shift from an automation mindset to an augmentation mindset, where sentiment mining tools support rather than replace human interpretation, aligns with the broader vision of human-centric AI that places human agency and dignity at the center of technological design.

10. Conclusion

The integration of explainable artificial intelligence with transformer-based sentiment mining for social media represents a frontier of socio-technical system design where technical sophistication must be matched by institutional wisdom. This paper has surveyed the architectural landscape of transformer models adapted for sentiment classification, examined the spectrum of XAI methods applicable to these architectures, articulated the principles of human-centric interpretation, and analyzed the infrastructure, governance, and policy dimensions that shape real-world deployment. Throughout this analysis, the centrality of attention mechanisms has emerged as a unifying thread, not only as the computational engine driving state-of-the-art performance but also as a potential bridge between machine representations and human understanding, provided that the limitations and potential pitfalls of attention-based explanation are acknowledged and addressed through careful validation and complementary methods.

The structural trade-offs identified in this analysis—between explanation fidelity and model complexity, between computational efficiency and interpretability depth, between standardized automation and context-sensitive human judgment—cannot be resolved through technical optimization alone. They are fundamentally design choices that embed values and priorities into sociotechnical systems. The principle of human-centric interpretation provides a guiding orientation, insisting that the ultimate purpose of explainable sentiment mining is not the generation of ever more precise attribution maps but the empowerment of human stakeholders to understand, question, and responsibly act upon the insights derived from the vast and noisy corpus of public sentiment expressed on social media. As regulatory frameworks solidify and societal expectations for algorithmic accountability intensify, the ability to design sentiment mining systems that are not only accurate but also transparent, fair, and sustainable will distinguish responsible innovators from those whose technologies risk exacerbating the very polarization and mistrust they purport to measure.

Future research directions should pursue tighter integration between generative explanation and attention-based visualization, the development of standardized benchmarks for evaluating explanation quality in subjective domains, longitudinal studies of how explanations influence user trust and decision-making in operational settings, and the creation of collaborative governance models that bring together platform companies, civil society organizations, and academic researchers to steward the public interest dimensions of sentiment monitoring infrastructure. The path forward lies not in a singular technical breakthrough but in the patient, interdisciplinary construction of explainable sentiment mining as a robust and legitimate sociotechnical practice.

References

1. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2017). Attention is all you need. In *Advances in Neural Information Processing Systems* (pp. 5998–6008).
2. Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019*

Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (pp. 4171–4186).

3. Jain, S., & Wallace, B. C. (2019). Attention is not explanation. In Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (pp. 3543–3556).
4. Sundararajan, M., Taly, A., & Yan, Q. (2017). Axiomatic attribution for deep networks. In Proceedings of the 34th International Conference on Machine Learning (pp. 3319–3328).
5. Clark, K., Khandelwal, U., Levy, O., & Manning, C. D. (2019). What does BERT look at? An analysis of BERT’s attention. In Proceedings of the 2019 ACL Workshop BlackboxNLP: Analyzing and Interpreting Neural Networks for NLP (pp. 276–286).
6. Wiegrefe, S., & Pinter, Y. (2019). Attention is not not explanation. In Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing (pp. 11–20).
7. Serrano, S., & Smith, N. A. (2019). Is attention interpretable? In Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics (pp. 2931–2951).
8. Liu, Y., Ott, M., Goyal, N., Du, J., Joshi, M., Chen, D., Levy, O., Lewis, M., Zettlemoyer, L., & Stoyanov, V. (2019). RoBERTa: A robustly optimized BERT pretraining approach. arXiv preprint arXiv:1907.11692.
9. Rogers, A., Kovaleva, O., & Rumshisky, A. (2020). A primer in BERTology: What we know about how BERT works. Transactions of the Association for Computational Linguistics, 8, 842–866.
10. Li, Q. (2026). Dynamic Adaptive Attention and Supervised Contrastive Learning: A Novel Hybrid Framework for Text Sentiment Classification. arXiv preprint arXiv:2604.10459.
11. Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). “Why should I trust you?” Explaining the predictions of any classifier. In Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (pp. 1135–1144).
12. Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. In Advances in Neural Information Processing Systems (pp. 4765–4774).
13. Abnar, S., & Zuidema, W. (2020). Quantifying attention flow in transformers. In Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics (pp. 4190–4197).
14. Miller, T. (2019). Explanation in artificial intelligence: Insights from the social sciences. Artificial Intelligence, 267, 1–38.
15. Ehsan, U., Liao, Q. V., Muller, M., Riedl, M. O., & Weisz, J. D. (2021). The human in human-centric AI: A literature review of empirical studies on XAI. arXiv preprint arXiv:2105.07191.
16. Liao, Q. V., Gruen, D., & Miller, S. (2020). Questioning the AI: Informing design practices for explainable AI user experiences. In Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems (pp. 1–15).

17. Camburu, O.-M., Rocktäschel, T., Lukasiewicz, T., & Blunsom, P. (2018). e-SNLI: Natural language inference with natural language explanations. In *Advances in Neural Information Processing Systems* (pp. 9539–9549).
18. Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., ... & Schafer, B. (2018). AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689–707.
19. Bhatt, U., Xiang, A., Sharma, S., Weller, A., Taly, A., Jia, Y., Ghosh, J., Puri, R., Moura, J. M., & Eckersley, P. (2020). Explainable machine learning in deployment. In *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency* (pp. 648–657).
20. Selbst, A. D., Boyd, D., Friedler, S. A., Venkatasubramanian, S., & Vertesi, J. (2019). Fairness and abstraction in sociotechnical systems. In *Proceedings of the Conference on Fairness, Accountability, and Transparency* (pp. 59–68).
21. Strubell, E., Ganesh, A., & McCallum, A. (2019). Energy and policy considerations for deep learning in NLP. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics* (pp. 3645–3650).
22. Sanh, V., Debut, L., Chaumond, J., & Wolf, T. (2019). DistilBERT, a distilled version of BERT: Smaller, faster, cheaper and lighter. *arXiv preprint arXiv:1910.01108*.
23. Blodgett, S. L., Barocas, S., Daumé III, H., & Wallach, H. (2020). Language (technology) is power: A critical survey of “bias” in NLP. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics* (pp. 5454–5476).
24. Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399.
25. Pruthi, D., Bansal, M., Dhingra, B., Neubig, G., & Lipton, Z. C. (2020). Evaluating explanations: How much do explanations from the teacher aid students? *Transactions of the Association for Computational Linguistics*, 8, 721–737.