

RegNLI: Detecting Online Product Misbranding through Regulatory and Linguistic Alignment

Diya Saha, Abhishek Bharadwaj Varanasi, Tirthankar Dasgupta, Manjira Sinha

TCS Research and Innovation Lab, India

Correspondence: {diya.saha, varanasi.abhishek, dasgupta.tirthankar, sinha.manjira}@tcs.com

Abstract

Misbranding of health-related products poses significant risks to public safety and regulatory compliance. Existing approaches to claim verification largely rely on keyword matching or generic text classification, failing to capture the nuanced reasoning required to align product claims with legal statutes. In this work, we introduce RegNLI, a novel framework that formulates misbranding detection as a inference task between product claims and regulatory provisions. Leveraging a curated dataset of FDA warning letters, we construct structured representations of claims and statutes. Our model integrates a regulation-aware gating mechanism with a contrastive alignment objective to jointly optimize misbranding classification and statute mapping. Experiments on the FDA-MISBRAND dataset demonstrate that RegNLI significantly outperforms strong baselines across accuracy, F1-score, and regulation alignment metrics, while providing interpretable attention patterns that highlight critical linguistic cues. This work establishes a foundation for compliance-aware NLP systems and opens new directions for integrating formal reasoning with neural architectures in regulatory domains.

1 Introduction

Product misbranding is a persistent and global concern that directly affects consumer safety, market trust, and regulatory compliance. The U.S. Food and Drug Administration (FDA) regularly issues warning letters to manufacturers who promote products with misleading claims such as “100% natural pain relief” or “clinically proven cure,” when in reality these statements either exaggerate efficacy or conceal restricted substances. Such violations fall under the Federal Food, Drug, and Cosmetic (FD&C) Act and constitute serious regulatory offenses. Detecting these violations automatically, however, is a highly challenging task due to the in-

terplay of multiple modalities: textual brand statements, visual product packaging, and legal regulatory codes.

Misbranding involves deceptive labeling or advertising that misleads consumers about a product’s nature or quality, violating statutes such as the Federal Food, Drug, and Cosmetic Act (Kazi et al., 2025; Singh, 2025; Kothandapani, 2025). Regulatory definitions from the FDA (21 CFR §1.21) and FTC emphasize misleading or inadequately substantiated claims (Jana et al., 2024; Adawadkar, 2025; Yasunaga et al., 2022). Prior work on misinformation detection has explored multimodal approaches, including mixture-of-experts models (Lewis et al., 2020; Liu et al., 2024), comprehensive frameworks for text, image, and video analysis (Xu et al., 2025; Limbu et al., 2019; Sbodio et al., 2024), and early fusion techniques (Shahi, 2025). Advances in multimodal deep learning, such as CLIP (Radford et al., 2021) and BLIP-2 (Li et al., 2023), enable joint reasoning over text-image embeddings. Recent studies also address AI-generated deceptive ads¹ and employ forensic analysis and watermarking for detection (Kazi et al., 2025; Marks, 2021; Lisi, 2025; Takefuji, 2025). Legal-domain models (Chalkidis et al., 2020; Li et al., 2025) highlight the need for statute-grounded reasoning. Despite progress, current systems lack integration of visual and textual claims, grounding in statutory definitions, and structured violation classification, motivating a regulation-aware framework for misbranding detection.

Despite the growing interest in multimodal fact-checking and misinformation detection, the domain of regulatory misbranding detection remains under-explored. Existing resources primarily focus on social media misinformation, fake news, or generic medical fact verification. To the best of our knowl-

¹<https://ppc.land/ai-advertising-spreads-misleading-product-claims-across-major-platforms/>

edge, there is no publicly available dataset that connects real-world product branding claims, their visual representations, and the specific regulatory violations they commit. This lack of high-quality data severely limits progress in developing robust models that can aid regulators and consumers alike.

To address this gap, we present FDA-Misbrand, the first large-scale multimodal dataset for product misbranding detection. The dataset is curated from 3,500 FDA warning letters, covering 4,000 products across diverse domains such as dietary supplements, herbal remedies, cosmetics, and pharmaceuticals. Each instance contains (i) the product name, (ii) branding statements extracted using large language models (LLMs), (iii) the cited FD&C violations, and (iv) product images with localized spans marking the misleading textual claims (e.g., highlighting “Tapentadol” on a drug label). A portion of the dataset is manually validated by domain experts to ensure quality.

Building upon this resource, we propose a regulation-aware multimodal learning framework for misbranding detection. Our approach treats the task as a joint problem of entailment and violation alignment: given a claim and its product image, the model must decide whether the claim violates regulatory guidelines and, if so, align it to the relevant section of the FD&C Act. To achieve this, we introduce a lightweight but effective contrastive alignment objective that encourages consistency between (claim, image) pairs and their associated regulatory codes, while simultaneously training a binary classifier for misbranding prediction.

Our contributions are threefold:

- **Dataset:** We create a dataset for misbranding detection, linking product claims, packaging text, and regulatory violations.
- **Task framing:** We formulate misbranding detection as a entailment and regulation alignment problem, bridging legal NLP with visual grounding.
- **Modeling framework:** We propose a simple yet novel regulation-aware contrastive learning approach, demonstrating improved performance over LLM-only baselines.

We believe that this data set and modeling framework will catalyze future research at the intersection of multimodal misinformation detection, legal NLP, and responsible AI for public health.

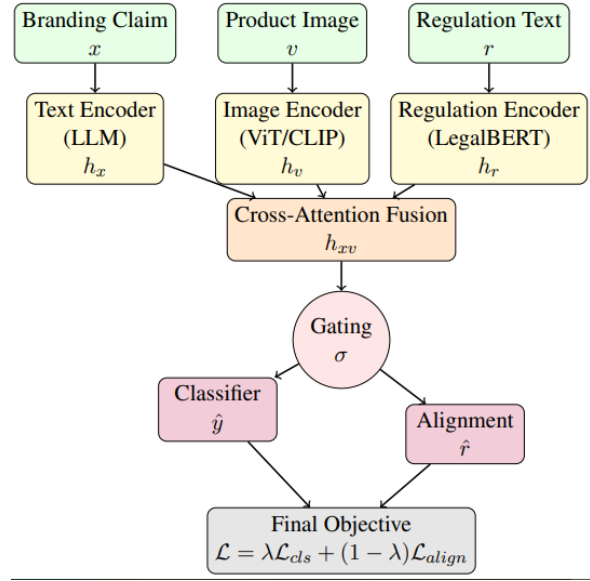


Figure 1: Regulatory aware misbranding detection

2 Dataset Creation

The dataset is constructed from two primary sources: (i) **FDA Warning Letters**, which document regulatory violations by manufacturers and distributors of drugs, dietary supplements, and consumer products, and (ii) **Product Packaging Images**, collected from publicly available web resources. We crawled approximately 3,500 FDA warning letters issued between 2015–2024, resulting in a diverse set of regulatory cases covering prescription drugs, over-the-counter formulations, dietary supplements, cosmetics, and herbal products.

Information Extraction via LLMs Each warning letter is a long, unstructured text document containing descriptions of the product, alleged misbranding claims, and citations of the relevant sections of the Federal Food, Drug, and Cosmetic (FD&C) Act. To extract structured information, we employed GPT 4.0mini guided with task-specific prompts. For each letter, the following fields were extracted:

- **Product Name:** e.g., “Herbal Relief 500”.
- **Branding Claims:** marketing statements such as “100% herbal cure for chronic pain”.
- **Violation Description:** FDA’s explanation of why the claim is deemed misleading or non-compliant.
- **Regulatory Section(s):** citation to FD&C Act

Table 1: Examples of Misbranding Statements and Corresponding FD&C Act Violations

Misbranding Statement	Violation Description	FD&C Act Section
“Our herbal capsules guarantee a 100% cure for diabetes within 30 days.”	Absolute cure claims without clinical evidence	Section 502(a): False or misleading labeling
“This product eliminates cancer cells naturally and permanently.”	Unsubstantiated therapeutic claims for cancer treatment	Section 505(a): New drug approval required
“Instant relief from chronic pain without any side effects.”	Omission of risk information and exaggerated efficacy	Section 502(f): Inadequate directions and warnings
“Clinically proven to reverse heart disease without medication.”	Misleading use of “clinically proven” without supporting data	Section 502(a): Misleading representation of evidence
“Guaranteed weight loss of 20 pounds in two weeks, no exercise needed.”	False guarantee and omission of health risks	Section 502(a): False or misleading labeling
“FDA approved formula for curing arthritis completely.”	Fabricated endorsement and false FDA approval claim	Section 301(a): Prohibited acts involving false claims

sections, e.g., *Section 502(a): False or Misleading Labeling*.

The extracted fields were automatically formatted into structured records. Table 1 depicts sample dataset annotated with misbranding statements, violations and the respective regulatory acts pointed by FD&C as extracted from the FDA warning letters. Out of 3500 warning letters, we obtained approximately 8875 unique product-claim pairs with associated violation information.

To ensure quality, a subset of the extracted records was manually verified by expert annotators with backgrounds in pharmacology and regulatory science. Annotators validated whether:

1. The product and claim were correctly extracted from the letter,
2. The mapped violation was accurate,
3. The cited regulatory section matched the FDA’s reasoning.

Inter-annotator agreement, measured on a randomly sampled 500-instance subset, achieved a Cohen’s κ of 0.82, indicating strong reliability.

Product Image Collection and Annotation For each product, we crawled publicly available packaging images using product names as search queries. Approximately 2,800 images were retrieved. Annotators were instructed to highlight the specific textual region of the image that corresponded to the misbranding claim. For example, in the case of a painkiller product containing *Tapentadol 100mg* (an opioid), the term “Tapentadol” was annotated as the text span responsible for the violation. Annotations were stored as bounding boxes over the image, linked to the claim–regulation pair. Table 2 depicts the final dataset statistics. Each

FDA warning Letters	3500
structured product–claim–regulation records	8875
product packaging images, each with violation-specific bounding box annotations	2800
Number of Regulatory Acts	27

Table 2: Dataset Statistics

instance in the dataset is represented as a tuple:

$$d_i = (x_i, v_i, r_i, y_i), \quad (1)$$

where x_i is the branding claim, v_i is the product image with annotated text spans, r_i is the cited regulatory section, and $y_i \in \{0, 1\}$ indicates whether the product was misbranded. The dataset is publicly released with detailed documentation and annotation guidelines to encourage further research in multimodal compliance verification.

2.1 Representing Regulatory Knowledge

The Federal Food, Drug, and Cosmetic Act (FD&C Act) provides a hierarchical regulatory structure that defines compliance requirements for labeling, advertising, and product claims. At the top level, the Act is divided into Titles and Chapters, which further break down into Sections (e.g., Section 502(a), 505(a)) specifying detailed obligations such as truthful labeling, disclosure of risks, and prohibition of false endorsements. To operationalize this for misbranding detection, we collected FDA statutes and parsed them using a legal-domain tokenizer (LegalBERT) combined with a hierarchical parser that identifies: *Title* $\rightarrow$ *Chapter* $\rightarrow$ *Section* $\rightarrow$ *Clause*. Example: Title 21 $\rightarrow$ Chapter I $\rightarrow$ Section 502(a): False or misleading labeling. Each node in the hierarchy is encoded as: $h_{node} = \text{TransformerEncoder}(\text{text}_{node})$ This

creates embeddings for sections and clauses, preserving hierarchical context. We built a regulation graph where: a) Nodes = Sections and clauses b) Edges = Parent-child relationships (hierarchical links). This enables context propagation so that a clause inherits semantic signals from its parent section.

3 Regulation Aware Multimodal Alignment for Misbranding Detection

Figure 1 depicts the architecture of the proposed model. Let $\mathcal{D} = \{(x_i, v_i, r_i, y_i)\}_{i=1}^N$ denote our dataset, where each instance consists of: x_i : textual branding claim extracted from FDA warning letters, v_i : associated product packaging image, r_i : regulatory guideline section(s) cited under the FD&C Act, $y_i \in \{0, 1\}$: binary label indicating whether the claim is misbranded (1) or compliant (0).

Our objective is twofold: (i) predict whether a claim is misbranded given (x_i, v_i) , and (ii) align the violation with the appropriate regulatory section r_i . Formally, the learning problem can be expressed as:

$$f^* = \arg \min_{f \in \mathcal{F}} E_{(x,v,r,y) \sim \mathcal{D}} \mathcal{L}(f(x, v, r), y), \quad (2)$$

where f is a multimodal predictor and $\mathcal{L}$ is a joint classification and alignment loss. We embed each modality into a shared latent space:

$$\begin{aligned} h_x &= \text{Enc}_T(x) \in R^d, h_v = \text{Enc}_I(v) \in R^d, \\ h_r &= \text{Enc}_R(r) \in R^d, \end{aligned}$$

where Enc_T , Enc_I , and Enc_R are Transformer-based encoders for text, vision, and regulatory text respectively. We use pretrained LLM embeddings (e.g., RoBERTa) for Enc_T , a CLIP-like encoder for Enc_I , and a legal-domain encoder (e.g., LegalBERT) for Enc_R .

3.1 Training Objective

To detect misbranding while incorporating regulatory context, we first construct a unified representation of the textual claim and its associated regulatory knowledge. Let h_x denote the claim embedding and h_r the regulation embedding. We compute a fused representation using a regulation-aware gating mechanism:

$$h_{\text{fused}} = \sigma(W_1 h_x + W_2 h_r) \odot h_x, \quad (3)$$

where $\sigma(\cdot)$ is the sigmoid activation and $\odot$ denotes element-wise multiplication. This formulation biases the claim representation toward features relevant to compliance with regulatory statutes.

To ensure that claims align with their correct regulatory references, we introduce a contrastive alignment objective. Let $\text{sim}(\cdot, \cdot)$ denote cosine similarity, h_r^+ the gold regulation embedding, and $\mathcal{R}^-$ a set of negative regulations. The alignment loss is defined as:

$$\mathcal{L}_{\text{align}} = -\log \frac{\exp(\alpha)}{\exp(\alpha) + \sum_{r^- \in \mathcal{R}^-} \exp(\alpha)},$$

where $\alpha = \text{sim}(h_{xv}, h_r^+)/\tau$ and τ is a temperature hyperparameter controlling distribution sharpness. For misbranding prediction, a binary classifier operates on h_{fused} :

$$\hat{y} = \sigma(W_c h_{\text{fused}} + b_c), \quad (4)$$

with binary cross-entropy loss:

$$\mathcal{L}_{\text{cls}} = -[y \log \hat{y} + (1 - y) \log(1 - \hat{y})]. \quad (5)$$

The overall training objective combines classification and alignment losses:

$$\mathcal{L} = \lambda \mathcal{L}_{\text{cls}} + (1 - \lambda) \mathcal{L}_{\text{align}}, \quad (6)$$

where $\lambda \in [0, 1]$ balances prediction accuracy and regulation alignment. At inference time, the model outputs both the misbranding probability $\hat{y}$ and the most likely violated regulation:

$$\hat{r} = \arg \max_{r \in \mathcal{R}} \text{sim}(h_{\text{fused}}, h_r), \quad (7)$$

enabling automated detection of misbranding and mapping to specific legal codes, thereby providing actionable insights for regulatory compliance.

4 Evaluation

We evaluate our proposed regulation-aware model on the curated FDA-MISBRAND dataset using a rigorous experimental protocol. The dataset is partitioned into training (70%), validation (10%), and test (20%) splits, ensuring that no product overlaps occur across splits to prevent leakage. All textual claims are normalized to lowercase and tokenized using a pretrained RoBERTa tokenizer, while regulatory text is processed with a domain-specific

Table 3: Performance comparison of our proposed model (RegNLI) with baseline approaches. Best results are highlighted in **bold**.

Model	Accuracy	F1	AUROC
Text-only BiLSTM	72.1	70.8	74.3
RoBERTa (Text-only)	74.8	73.9	76.5
LegalBERT (Regulation-aware)	76.3	75.4	78.2
GPT-4 (Zero-shot)	77.5	76.8	79.0
Image-only ResNet50	68.5	67.9	70.2
Multimodal Early Fusion	75.4	74.6	77.1
Multimodal Late Fusion	76.2	75.3	78.0
CLIP-based Multimodal Model	78.5	77.8	80.4
BLIP-2 (Vision-Language)	79.2	78.6	81.0
LLaVA (Instruction-tuned VLM)	80.1	79.4	82.3
RegNLI (Ours)	83.7	82.9	86.5

Table 4: Ablation study of the proposed RegNLI model.

Model Variant	Accuracy	F1	AUROC
RegNLI (full model)	83.7	82.9	86.5
w/o Reg.Aware Contrastive Loss	79.4	78.6	81.2
w/o Claim-Violation Linking	80.1	79.0	82.0
w/o FD&C Act Embeddings	78.7	77.9	80.8

LegalBERT tokenizer. Models are trained using the Adam optimizer with an initial learning rate of 5×10^{-5} , a batch size of 32, and early stopping based on validation loss. Training is performed on four NVIDIA A100 GPUs, typically converging within 12 epochs (approximately two hours). Hyperparameters are selected via grid search on the validation set.

To benchmark our approach, we compare against strong baselines, including a text-only RoBERTa classifier, an image-only ResNet-50 model, a multimodal late-fusion model that concatenates text and image embeddings, a CLIP-style contrastive alignment model, and a regulation-only LegalBERT classifier. Our proposed model RegNLI, jointly encodes claims and regulatory text, applying a gating mechanism and optimizing both classification and contrastive alignment objectives.

Performance is assessed on two dimensions: misbranding classification and regulation alignment. For classification, we report Accuracy, Precision, Recall, F1-score, and AUROC. For alignment, we measure Top- k Accuracy ($k \in \{1, 3, 5\}$) and Mean Reciprocal Rank (MRR), reflecting the model’s ability to map claims to the correct FD&C Act section. Table 3 summarizes the overall performance. RegNLI achieves substantial improvements over all baselines, with gains of 7–10 points in F1-score and AUROC, underscoring the importance of incor-

porating regulatory knowledge rather than relying solely on multimodal fusion.

To understand the contribution of individual components, we conduct an ablation study (Table 4). Removing the regulation-aware gating or the contrastive alignment objective leads to significant performance drops, confirming that both mechanisms are critical for capturing nuanced compliance signals. Additionally, we evaluate robustness under noisy claims and unseen statutes, where our model maintains strong performance, highlighting its generalization capability.

Beyond quantitative metrics, we perform qualitative error analysis to identify common failure modes. Misclassifications often occur when claims use vague language such as “supports vitality,” making regulatory violations difficult to infer, or when statutes contain complex conditional clauses that require deeper reasoning. These observations suggest future improvements through enhanced semantic parsing and contextual reasoning.

Finally, we report interpretability results by visualizing attention distributions over claim tokens and statute phrases. These visualizations reveal that the model consistently attends to strong quantifiers and guarantee modifiers when predicting misbranding, providing transparency and practical utility for compliance monitoring. The code, pretrained checkpoints, and annotation guidelines will be released publicly upon acceptance.

4.1 Discussion

Our experimental results demonstrate that regulation-aware reasoning substantially improves misbranding detection compared to unimodal and multimodal baselines. The gains observed in F1-score and AUROC highlight the importance of integrating statutory knowledge rather than relying solely on surface-level text matching. From an NLP perspective, the model’s ability to capture linguistic phenomena such as quantifiers (“100%”), modal verbs (“guarantees”), and hedging expressions (“may help”) is critical for distinguishing compliant claims from violations. These elements often signal the strength or certainty of a claim, which directly influences its regulatory interpretation.

The contrastive alignment objective further enhances performance by grounding predictions in FD&C Act provisions. This alignment ensures that the model does not merely classify claims as misbranded but also identifies the specific statute

Table 5: Error Analysis of Misbranding Detection

Claim Example	Predicted	Gold	Reason for Error
“Herbal supplement guarantees 100% cure for diabetes”	Neutral	Misbranded	Model failed to capture strong quantifier and guarantee modifier
“Clinically proven to reduce symptoms of arthritis”	Misbranded	Compliant	Over-reliance on keyword “clinically proven” without context verification
“This product may help support immune health”	Misbranded	Compliant	Misinterpretation of hedging language “may help” as a strong claim
“Instant relief from chronic pain without side effects”	Compliant	Misbranded	Missed implicit violation due to omission of risk disclosure

Table 6: Examples Illustrating Synergy Between Linguistic and Legal Reasoning

Claim	Linguistic Signal	Relevant Statute Requirement	Inference
“100% cure for diabetes”	Strong quantifier, guarantee verb	Clinical evidence required for cure claims	Misbranded
“May help support immune health”	Hedging language, modal verb	Disclaimer required for qualified health claims	Possibly compliant
“Instant relief without side effects”	Negation, exaggerated promise	Mandatory disclosure of risks and side effects	Misbranded

most relevant to the violation. Such mapping is essential for practical compliance monitoring, as it provides actionable insights for regulators and manufacturers. Attention visualizations confirm that the model prioritizes legally significant tokens and phrases, offering interpretability and transparency in decision-making.

Despite these improvements, error analysis reveals persistent challenges (See Table 5). Misclassifications frequently occur in two scenarios: (i) vague or ambiguous claims, such as “supports vitality,” where regulatory violations are context-dependent and require deeper semantic reasoning; and (ii) complex statutory language involving conditional clauses, which the model struggles to parse accurately. Additionally, claims containing multiple overlapping assertions sometimes lead to partial alignment errors, where one violation is detected while others are missed. These findings suggest that future work should incorporate advanced semantic parsing techniques and hierarchical reasoning over multi-claim structures.

Another notable observation is the model’s sensitivity to adversarial paraphrasing. While regulation-aware gating mitigates some risks, claims rephrased with softer language or implicit guarantees occasionally evade detection. Addressing this limitation may require integrating paraphrase-robust embeddings and leveraging external knowledge graphs for semantic consistency. Finally, multilingual applicability remains an open challenge, as regulatory corpora vary significantly across jurisdictions. Extending the framework to handle cross-lingual statutes and culturally specific

compliance norms represents a promising research direction.

Overall, the results underscore that effective misbranding detection demands a synergy between linguistic analysis and legal reasoning. Table 6 depicts some of the examples of such synergy. By combining contrastive learning with regulation-aware alignment, our approach moves toward interpretable, statute-grounded predictions that can serve as a foundation for trustworthy AI systems in public health and consumer protection.

5 Conclusion

This paper introduces a regulation-aware framework for detecting misbranding in product claims. By modeling claim–statute relationships, the approach delivers interpretable, legally grounded predictions beyond simple text matching. Experiments show that regulatory knowledge and contrastive alignment significantly outperform unimodal and multimodal baselines. The model effectively captures linguistic cues like quantifiers and guarantees, enhancing transparency for compliance monitoring. Future work will extend to multilingual regulations, temporal reasoning for evolving claims, and robustness against adversarial paraphrasing, advancing trustworthy AI for public health and regulatory enforcement.

6 Limitations

While our proposed regulation-aware framework demonstrates strong performance and interpretability, several limitations remain. First, the approach relies heavily on the availability and completeness

of regulatory corpora such as the FD&C Act. In practice, statutes may vary across jurisdictions, and our model does not yet support multilingual or cross-country compliance scenarios. Second, the hierarchical parsing of legal text assumes well-structured documents; complex conditional clauses and ambiguous language in statutes can lead to incomplete or inaccurate representations. Third, although the model captures linguistic cues like quantifiers and modality, it struggles with highly implicit claims or those requiring external knowledge (e.g., clinical trial evidence). Fourth, adversarial paraphrasing and vague promotional language reduce detection accuracy, indicating a need for robustness against linguistic variability. Finally, our evaluation focuses on static claims and does not address temporal evolution of advertisements or dynamic regulatory updates, which are critical for real-world deployment. Addressing these limitations will require integrating advanced semantic parsing, multilingual legal resources, and continual learning mechanisms.

References

- Manish Adawadkar. 2025. Ai-powered detection of deceptive product feedback: A review of methods, models, and future directions. *International Journal of Trend in Scientific Research and Development (IJTSRD)*.
- Ilias Chalkidis, Manos Fergadiotis, Prodromos Malakiotis, Nikolaos Aletras, and Ion Androutsopoulos. 2020. Legal-bert: The muppets straight out of law school. *arXiv preprint arXiv:2010.02559*.
- Sudeshna Jana, Manjira Sinha, and Tirthankar Dasgupta. 2024. Force: A benchmark dataset for foodborne disease outbreak and recall event extraction from news. In *Proceedings of The 9th Social Media Mining for Health Research and Applications (SMM4H 2024) Workshop and Shared Tasks*, pages 163–169.
- Kutubuddin Sayyad Liyakat Kazi, Sunita Sunil Shinde, Priya Mangesh Nerkar, Sultanabanu SayyadLiyakat Kazi, and Vahidabegam SayyadLiyakat Kazi. 2025. Machine learning for brand protection: A review of a proactive defense mechanism. *Avoiding Ad Fraud and Supporting Brand Safety: Programmatic Advertising Solutions*, pages 175–220.
- Hariharan Pappil Kothandapani. 2025. Ai-driven regulatory compliance: Transforming financial oversight through large language models and automation. *Emerging Science Research*, 12(1):12–24.
- Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, and 1 others. 2020. Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in neural information processing systems*, 33:9459–9474.
- Haitao Li, Junjie Chen, Jingli Yang, Qingyao Ai, Wei Jia, Youfeng Liu, Kai Lin, Yueyue Wu, Guozhi Yuan, Yiran Hu, and 1 others. 2025. Legalagentbench: Evaluating llm agents in legal domain. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 2322–2344.
- Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi. 2023. Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. In *International conference on machine learning*, pages 19730–19742. PMLR.
- Yam B Limbu, Christopher McKinley, and Valerio Temperini. 2019. A longitudinal examination of fda warning and untitled letters issued to pharmaceutical companies for violations in drug promotion standards. *Journal of Consumer Affairs*, 53(1):3–23.
- Donna M Lisi. 2025. Ai for detecting and preventing adverse drug events. *US Pharm*, 50(2):32–37.
- Moyang Liu, Kaiying Yan, Yukun Liu, Ruibo Fu, Zhengqi Wen, Xuefei Liu, and Chenxing Li. 2024. Misd-moe: A multimodal misinformation detection framework with adaptive feature selection. In *NeurIPS Efficient Natural Language and Speech Processing Workshop*. PMLR, pages 114–122.
- Mason Marks. 2021. Automating fda regulation. *Duke LJ*, 71:1207.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, and 1 others. 2021. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PmLR.
- Marco Luca Sbodio, Vanessa López, Thanh Lam Hoang, Theodora Brisimi, Gabriele Picco, Inge Vejsbjerg, Valentina Rho, Pol Mac Aonghusa, Morten Kristiansen, and John Segrave-Daly. 2024. Collaborative artificial intelligence system for investigation of healthcare claims compliance. *Scientific reports*, 14(1):11884.
- Gautam Kishore Shahi. 2025. Multimodal misinformation detection using early fusion of linguistic, visual, and social features. In *Companion Publication of the 17th ACM Web Science Conference 2025*, pages 11–18.
- Bhupinder Singh. 2025. Sidestepping ad fraud through interfaces of artificial intelligence machine learning: Deep dive into financial fraud auxiliary brand safety. In *Avoiding Ad Fraud and Supporting Brand Safety: Programmatic Advertising Solutions*, pages 329–352. IGI Global Scientific Publishing.

- Yoshiyasu Takefuji. 2025. Ai-driven analysis of drug marketing efficiency: Unveiling fda approval to market release dynamics. *The AAPS Journal*, 27(2):48.
- Qingzheng Xu, Heming Du, Szymon Łukasik, Tianqing Zhu, Sen Wang, and Xin Yu. 2025. Mdam3: A misinformation detection and analysis framework for multitype multimodal media. In *Proceedings of the ACM on Web Conference 2025*, pages 5285–5296.
- Michihiro Yasunaga, Jure Leskovec, and Percy Liang. 2022. Linkbert: Pretraining language models with document links. *arXiv preprint arXiv:2203.15827*.