

CAX-Gate: a Context-Aware Explainability Fusion Architecture for Dynamical Adjustment of LIME-to-SHAP attributions ratio with application in Colonoscopy

Pınar Karadayı Atas^{1,*†}, Luca Longo^{2†}

¹Department of Software Engineering, Istanbul Arel University, Istanbul, Türkiye

²University College Cork, Ireland

Abstract

Early diagnosis of colorectal polyps, the main antecedents of colorectal cancer, is crucial to stopping the disease's growth. Although deep learning-based polyp recognition systems have performed well on benchmark datasets, existing explainability techniques employ uniform attribution strategies, which prevent interpretive depth from being tailored to the diagnostic ambiguity or clinical relevance of specific cases. While borderline or visually ambiguous lesions require more rigorous reasoning, regular detections in practice require quick interpretability. This research contributes to the body of knowledge by introducing CAX-Gate, a context-aware explainability fusion architecture that dynamically adjusts the LIME-to-SHAP ratio based on sample-specific diagnostic conditions to overcome this constraint. This is applied in gastrointestinal imaging, incorporating dynamic, context-level attribution fusion, and supporting the deployment viability and interpretive accuracy of AI-assisted colonoscopy. Experimental work on the Hyper-Kvasir dataset suggests that CAX-Gate avoids unnecessary computational overhead in routine findings while maintaining diagnostic clarity in challenging cases, such as flat and low-contrast polyps. When diagnostic uncertainty increases, the performance difference between CAX-Gate and baseline approaches widens. Here, CAX-Gate mitigates instability observed in local perturbation-based explanations but preserves interpretative fidelity by prioritising globally consistent SHAP attributions in ambiguous contexts.

Keywords

Explainable Artificial Intelligence (XAI), Medical Image Analysis, Uncertainty-Aware Deep Learning, Colorectal Polyp Detection, Interpretability, Colonoscopy, LIME-to-SHAP attributions ratio, XAI fusion architecture.

1. Introduction

Most malignant cases of colorectal cancer (CRC) originate from initially benign colorectal polyps [1, 2]. CRC remains one of the leading causes of cancer-related mortality worldwide, making the early and accurate detection of polyps during colonoscopy the most critical intervention point for preventing malignant progression. Despite recent advances in computer-assisted polyp detection and classification, deep learning models still function largely as static and opaque decision systems [3]. Although their predictions are often accurate, they rarely provide meaningful justifications for classifying lesions as benign, suspicious, or malignant, nor do they adapt the depth of explanation to diagnostic uncertainty. This limitation has important clinical implications. Subtle lesions such as flat, faint, or low-contrast polyps require deeper interpretive support, while routine high-confidence detections benefit from faster and computationally lighter explanations. The mismatch between clinical variability and fixed interpretability strategies highlights the need for adaptive explainability mechanisms that can dynamically adjust explanatory granularity according to case-specific diagnostic conditions.

Most explainability approaches rely on a fixed attribution strategy regardless of lesion complexity, visual ambiguity, or clinical risk. To address this limitation, this research study introduces CAX-Gate,

Late-breaking work, Demos and Doctoral Consortium, colocated with the 4th World Conference on eXplainable Artificial Intelligence: July 01–03, 2026, Fortaleza, Brazil

*Corresponding author

† These authors contributed equally to this work.

✉ pinaratas@arel.edu.tr (P. K. Atas); luca.longo@tudublin.ie (L. Longo)

ORCID 0000-0002-9429-8463 (P. K. Atas); 0000-0002-2718-5426 (L. Longo)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

a context-aware explainability framework that dynamically regulates attribution between LIME and SHAP [4] based on prediction confidence, predictive uncertainty, image-derived diagnostic risk proxies, and computational constraints. By adapting explanatory depth to diagnostic uncertainty, CAX-Gate transforms explainability from a static post-hoc visualisation into a flexible clinical support mechanism for colonoscopic decision-making. Such a framework supports practical deployment in endoscopic environments by prioritising lightweight explanations for high-confidence predictions, thereby reducing latency and enabling real-time interpretation. The approach seeks to improve interpretability in challenging visual scenarios, such as flat or low-contrast polyps, by triggering deeper SHAP-based explanations and aligning explainable outputs with clinical diagnostic reasoning.

2. Related Work

Recent studies have explored various machine learning approaches for prognostic and risk prediction tasks in colorectal cancer. Talebi et al. [5] applied several conventional machine learning models to predict metastasis in rectal cancer patients under active monitoring, demonstrating robust early-detection performance despite limited clinical data. Similarly, Yang et al. [6] proposed a predictive model capable of identifying abnormal metastatic patterns following surgical resection, although the study relied on relatively simple model architectures and a small sample size. Sardari et al. [7] extended predictive modelling to genetic pathways associated with inflammatory bowel disease and colorectal cancer, achieving high accuracy but highlighting computational complexity challenges. Microbiome-based CRC classification has also attracted increasing attention. Freitas et al. [8] utilised microbial signatures for CRC detection, while Lu et al. [9] used gut microbiota-derived features to differentiate disease stages with high accuracy. However, both studies were constrained by dataset size and biological confounders. Jiang et al. [10] employed multi-model ensembles to predict mismatch repair deficiency (dMMR), achieving fast and accurate discrimination despite limited data. Other studies focused on disease progression and survival prediction. Galadima et al. [11] investigated late-stage CRC prediction, while Cardoso et al. [12] and Rodriguez et al. [13] analysed large registry and EHR cohorts to model survival trajectories, recurrence risk, and biomarker integration. Additional work by Lam et al. [14] and Mohammadi et al. [15] demonstrated strong predictive performance but limited generalizability due to institutional data constraints. Research on molecular characterisation has also emerged; for example, Badisy et al. [16] proposed ML-based survival modelling, while Zhong et al. [17] identified oxidative stress-related molecular subtypes using TCGA and GEO datasets. Ahmadih-Yazdi et al. [18] further investigated metastatic biomarkers, though external validation was limited. Tahir et al. [19] proposed a fixed LIME–SHAP integration framework for real-time interpretability, highlighting the potential of combining local and global explanations in time-sensitive settings. Most colorectal cancer prediction studies focus on improving predictive accuracy using machine learning and deep learning, while model interpretability has received less attention. Some research has applied XAI methods such as saliency maps, feature importance analysis, and attribution techniques like LIME and SHAP to highlight influential variables or image regions. However, these approaches typically treat explainability as a static post-processing step and apply the same attribution strategy to all cases. This ignores differences in diagnostic uncertainty and visual ambiguity in real clinical scenarios, leading to explanations that may be too simple for complex cases or unnecessarily costly for routine detections. These limitations motivate the need for adaptive explainability mechanisms that adjust explanation depth according to case-specific diagnostic conditions.

3. Material and Methods

Materials – HyperKvasir, a large-scale gastrointestinal (GI) endoscopy dataset published in 2020 [20], contains four primary data types: segmented images, labelled images, unlabeled images, and labelled videos. For this study, only the tagged image subset associated with the *Lower GI Tract* was utilised. This collection includes 7220 manually annotated endoscopic images from routine gastroenterological

examinations conducted at Norway’s Baerum Hospital, with several annotations verified by qualified endoscopists. Fig. 1 displays representative visual samples from the Hyper-Kvasir lower GI dataset, highlighting characteristic mucosal structures commonly observed during colonoscopic evaluation. Accordingly, the lower GI portion of HyperKvasir was selected for this analysis. Twelve subclasses and four clinical groupings were derived from the sixteen diagnostic classes originally defined in the dataset. However, several subclasses suffer from severe data sparsity, leading to class imbalance during deep neural network training. To mitigate this issue, five subclasses—ileum, haemorrhoids, and ulcerative colitis grades 0–1, 1–2, and 2–3—were removed, resulting in a refined corpus of 7131 images. The raw images were converted into NumPy tensors and resized to $224 \times 224 \times 3$. Pixel intensities were normalised by dividing by 255.0. Class labels were encoded using *LabelEncoder* followed by *One-Hot Encoding* to enable multiclass inference. The dataset was split into training (70%), validation (10%), and testing (20%) sets while preserving class distributions. All backbone networks were pre-trained on ImageNet using the *Adam* optimiser with *categorical cross-entropy* loss.

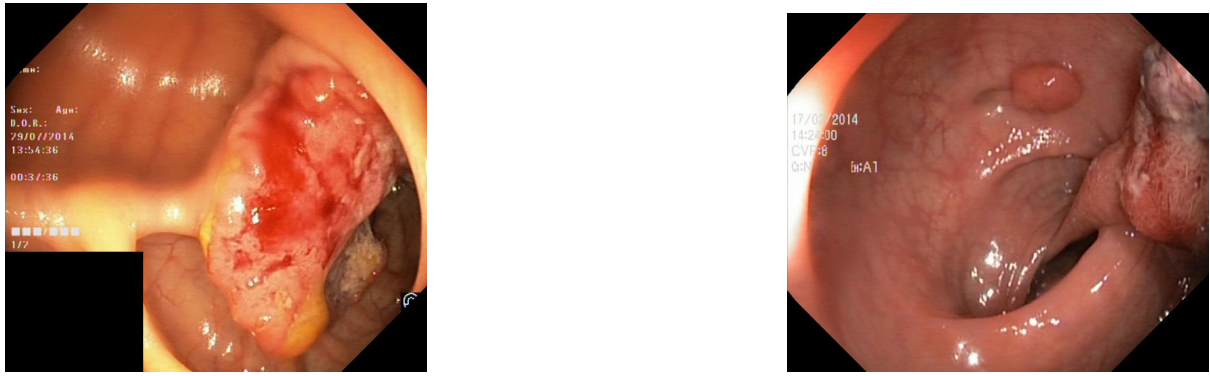


Figure 1: Two representative lower gastrointestinal images obtained from the Hyper-Kvasir dataset [20].

To account for morphological variance and prevent overfitting, controlled augmentation techniques were applied to the training set. These included random cropping, contrast-limited adaptive histogram equalisation (CLAHE), horizontal flipping ($p = 0.5$), moderate rotations ($\pm 10^\circ$), and Gaussian noise injection. No artificial lesion distortion was introduced in order to maintain clinically realistic visual patterns. HyperKvasir is publicly available and contains no personally identifiable information; therefore, ethics approval was not required.

Background Methods – LIME is a local surrogate-based interpretability technique that approximates the behaviour of complex classifiers around a specific instance [21]. Instead of explaining the entire decision surface, LIME generates perturbed input samples and fits a simplified interpretable model in the neighbourhood of the original instance. It assumes that even highly nonlinear decision boundaries may be approximated locally by a simpler model. SHAP is an explainability framework derived from cooperative game theory [22]. It assigns a Shapley value to each feature, representing its contribution to the final prediction across all possible feature combinations. This approach provides theoretically consistent feature attributions but can be computationally expensive for high-dimensional data such as medical images. Predictive uncertainty plays a critical role in medical decision support systems. When prediction confidence is high, simpler explanations may suffice, whereas uncertain predictions may require more comprehensive interpretability. In this research, predictive uncertainty is estimated using Shannon entropy to measure the dispersion of class probabilities.

$$H(x) = - \sum_{c=1}^C p(c|x) \log p(c|x) \quad (1)$$

where $p(c|x)$ represents the predicted probability of class c for input sample x .

For each sample x , a compact context vector $z(x)$ is constructed from maximum softmax confidence $c(x)$, predictive entropy $H(x)$, a diagnostic-risk proxy $r(x)$, and a normalized computational-budget term $b(x)$. Here, $c(x) = \max_c p(c | x)$ measures confidence, while $H(x)$ captures uncertainty. As no patient-level metadata were used, clinical risk was represented through image- and prediction-derived proxies, with higher $r(x)$ values assigned to more ambiguous cases.

Proposed Method – The proposed framework introduces a dynamically adaptive mechanism for combining local (LIME-based) and global (SHAP-based) explanations. Conventional post-hoc explainability techniques often rely on fixed weighting strategies or treat each method independently. Such static approaches may be suboptimal in clinical environments where interpretability requirements and computational constraints vary across cases. To address this limitation, the proposed approach learns a context-dependent fusion strategy that dynamically balances LIME and SHAP explanations. Rather than manually defining fusion weights, the framework uses a learnable gating mechanism that takes the context vector as input and generates a mixing coefficient controlling the contribution of LIME and SHAP explanations. The gating module is implemented as a lightweight multilayer perceptron that maps $z(x)$ to a scalar fusion coefficient $\alpha(x) \in [0, 1]$. It consists of an input layer, fully connected hidden layer(s) with ReLU activation, and a sigmoid output unit. Higher α values favour LIME, whereas lower values increase the contribution of SHAP, while keeping the added computational overhead minimal.

$$\phi_{hyb}(i) = \alpha\phi_{LIME}(i) + (1 - \alpha)\phi_{SHAP}(i) \quad (2)$$

where $\phi_{LIME}(i)$ and $\phi_{SHAP}(i)$ denote the attribution of feature i obtained from LIME and SHAP respectively. In Eq. (2), α is predicted once per sample and then applied across all feature attributions of that sample. The gating mechanism is trained using an objective that balances explanation fidelity and computational cost. During training, high-fidelity SHAP explanations computed offline were used as reference targets. The gating network was optimized so that the fused explanation remained close to these reference attributions while discouraging unnecessary SHAP-dominant behaviour in low-uncertainty cases, thereby preserving computational efficiency for routine detections. During inference, the context vector is extracted from the prediction, the gating mechanism estimates α , and LIME and SHAP are fused into a hybrid explanation.

Evaluation strategy – CAX-Gate was evaluated on Hyper-Kvasir for deep learning-based polyp detection. Evaluation covered classification performance, fidelity, stability, and computational efficiency. Each experiment is repeated five times with different random seeds for training/validation/test splits to reduce stochastic bias, and means and standard deviations are considered. Statistical significance between explainability methods is assessed using paired two-tailed Student’s t -tests with Bonferroni correction ($p < 0.05$). No severe deviations from normality were observed in the paired differences. LIME explanations are generated using 1000 perturbations per instance, while SHAP explanations are computed using a set of 100 samples and up to 2000 Monte Carlo iterations. CAX-Gate model dynamically learned a fusion coefficient α to combine LIME and SHAP attributions. The gating module was trained so that the fused attribution map approximates an offline SHAP reference while avoiding unnecessary SHAP-dominant behaviour in low-uncertainty cases. L2 similarity, cosine similarity, and SSIM were used to quantify agreement between each method and the reference attribution maps.

4. Results and Discussion

Prior to analysing explainability behaviour, the predictive performance of the underlying classifier was verified. Table 1 and Figure 2 show that the proposed CNN backbone achieves competitive classification results compared with baselines (SVM, Random Forest, ResNet-50, and EfficientNet-B0). Overall, it achieved the best performance with an accuracy of 94.2%, balanced accuracy of 93.6%, and an AUC of 0.972.

The proposed CAX-Gate framework achieved the highest fidelity across all evaluation metrics (Table 2), outperforming LIME, approximate SHAP, and static fusion. In particular, CAX-Gate achieved an L2 fidelity of 0.91, compared with 0.63 for LIME and 0.82 for static fusion. These improvements demonstrate that context-aware fusion produces explanations that are more consistent with reference attributions.

Table 1

Comparative classification performance of different baseline models on the Hyper-Kvasir test set.

Method	Accuracy	Balanced Acc.	Precision	F1-score	AUC
SVM (HOG features)	86.7 ± 0.9	84.2 ± 1.1	85.6 ± 1.0	85.1 ± 0.8	0.901 ± 0.007
Random Forest	88.3 ± 0.8	86.9 ± 0.9	87.8 ± 0.7	87.4 ± 0.6	0.923 ± 0.006
ResNet-50	92.1 ± 0.6	91.4 ± 0.6	91.8 ± 0.5	91.9 ± 0.5	0.958 ± 0.004
EfficientNet-B0	93.4 ± 0.5	92.7 ± 0.6	93.1 ± 0.4	93.0 ± 0.4	0.966 ± 0.003
Proposed CNN Backbone	94.2 ± 0.4	93.6 ± 0.5	93.8 ± 0.5	94.1 ± 0.4	0.972 ± 0.003

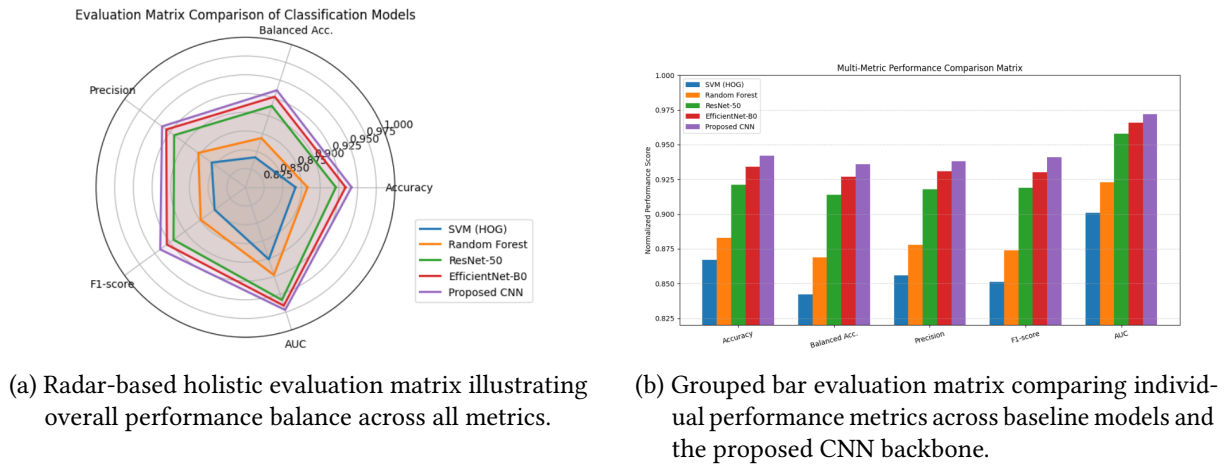


Figure 2: Complementary evaluation visualizations of classification performance on the Hyper-Kvasir dataset. The radar plot emphasizes holistic performance consistency, while the grouped bar chart highlights metric-wise comparative superiority of the proposed CNN backbone.

Table 2

Global explainability fidelity comparison across attribution strategies using multiple fidelity metrics.

Method	L2 Fidelity $\uparrow$	Cosine Sim. $\uparrow$	SSIM $\uparrow$	Std. Dev.	Effect Size (d)
LIME	0.63	0.66	0.59	0.05	—
SHAP (Approx.)	0.89	0.91	0.87	0.03	1.21
Static Fusion ($\alpha = 0.5$)	0.82	0.84	0.79	0.04	0.87
CAX-Gate (Proposed)	0.91	0.93	0.89	0.02	1.34

To analyse the adaptivity of the gating mechanism, test samples were divided into low-, medium-, and high-uncertainty groups based on prediction entropy. Table 3 shows that the learned fusion coefficient α varies systematically with uncertainty. Higher L2 fidelity, cosine similarity, and SSIM indicate closer agreement with the reference SHAP explanations, whereas lower standard deviation indicates more consistent attribution quality across runs. In low-uncertainty cases, higher α values indicate stronger reliance on LIME explanations, enabling faster interpretation. As uncertainty increases, the gating mechanism shifts toward SHAP-based explanations, which provide more globally consistent attribution. The Spearman correlation analysis confirmed a strong negative relationship between entropy and α ($\rho = -0.81$, $p < 0.001$).

Table 4 further examines explainability fidelity under different uncertainty regimes. While all

Table 3Extended statistical characterization of the learned fusion coefficient α across predictive uncertainty regimes.

Uncertainty Level	Mean α	Median	IQR	Std. Dev.	Min–Max	Mean SHAP Weight ($1 - \alpha$)
Low	0.76	0.78	0.11	0.09	0.58 – 0.91	0.24
Medium	0.48	0.46	0.14	0.11	0.31 – 0.67	0.52
High	0.23	0.21	0.10	0.08	0.09 – 0.41	0.77

methods perform reasonably well in low-uncertainty cases, the advantage of CAX-Gate becomes more pronounced as uncertainty increases.

Table 4

Explainability fidelity stratified by diagnostic uncertainty using reference SHAP explanations.

Uncertainty Level	LIME	SHAP (Approx.)	Static Fusion	CAX-Gate (Proposed)	Relative Gain vs LIME (%)
Low	0.71 ± 0.06	0.88 ± 0.03	0.86 ± 0.04	0.89 ± 0.02	+25.4
Medium	0.61 ± 0.07	0.89 ± 0.03	0.84 ± 0.05	0.92 ± 0.02	+50.8
High	0.54 ± 0.08	0.90 ± 0.04	0.80 ± 0.06	0.94 ± 0.02	+74.1

In the high-uncertainty regime, CAX-Gate achieved a fidelity of 0.94, corresponding to a 74% improvement over LIME and a clear improvement over static fusion. This result indicates that adaptive attribution fusion is particularly beneficial for diagnostically ambiguous cases. Explanation stability was analysed by measuring fidelity variance under repeated perturbations. As shown in Table 5, LIME exhibits the highest variance due to its sensitivity to local perturbations. SHAP provides more stable explanations but at higher computational cost.

Table 5

Explanation stability comparison across attribution methods (lower is better).

Method	Mean Fidelity Variance	Std. Dev.	Coefficient of Variation	Relative Reduction (%)
LIME	0.014	0.003	21.4%	–
SHAP	0.006	0.001	16.7%	57.1
Static Fusion	0.009	0.002	22.2%	35.7
CAX-Gate (Proposed)	0.005	0.001	20.0%	64.3

CAX-Gate achieved the lowest fidelity variance (0.005), corresponding to a 64% reduction compared with LIME. This demonstrates that uncertainty-aware attribution fusion improves explanation consistency. Computational cost was evaluated by measuring the per-sample runtime of each explanation method (Table 6). SHAP exhibited the highest cost (310ms per sample), while LIME was the fastest (42ms).

Table 6

Detailed computational cost comparison per sample.

Method	Mean Time (ms)	Std. Dev.	Min–Max (ms)	Throughput (fps)	Speed-up vs SHAP
LIME	42	6	31–55	23.8	$\times 7.4$
SHAP	310	28	265–352	3.2	1.0
Static Fusion	182	19	151–216	5.5	$\times 1.7$
CAX-Gate (Proposed)	97	11	81–122	10.3	$\times 3.2$

CAX-Gate achieved a balanced trade-off between speed and fidelity, with an average runtime of 97 ms per sample. This represents a $3.2\times$ speed improvement over SHAP while maintaining significantly higher explanation fidelity than LIME. Overall, the experimental results demonstrate that CAX-Gate consistently improves explainability fidelity, stability, and efficiency compared with single-method attribution techniques and static fusion baselines. By dynamically adapting the fusion of local and global explanations based on prediction uncertainty, the proposed framework provides more reliable and clinically meaningful interpretations. These findings suggest that context-aware explainability can play a critical role in enabling trustworthy deployment of deep learning models in colonoscopy.

4.1. Discussion

Findings demonstrate that explanation adaptivity is essential for achieving clinically meaningful interpretability. Traditional explainability methods apply the same attribution procedures to both routine and ambiguous cases, assuming homogeneous diagnostic conditions. However, visual patterns in gastrointestinal endoscopy vary significantly, and the interpretive requirements for clear polyp detection differ from those for low-contrast or borderline lesions. CAX-Gate addresses this heterogeneity by conditioning explanation fusion on uncertainty and confidence. As diagnostic uncertainty increases, the performance gap between CAX-Gate and baseline approaches becomes more pronounced, which is particularly important in clinical settings where inconsistent explanations may affect clinicians' trust and decision-making. The framework reduces instability observed in local perturbation-based explanations while preserving interpretative fidelity by prioritising globally consistent SHAP attributions in ambiguous cases. Stability analysis further shows that CAX-Gate produces more reproducible explanations under repeated perturbations, supporting clinical auditing, long-term monitoring, and regulatory assessment. These improvements are achieved without substantial computational overhead, enabling near-real-time operation in live endoscopic workflows.

5. Conclusion

We introduced CAX-Gate, a context-aware fusion framework that adaptively balances LIME and SHAP for colorectal polyp analysis. Experiments on Hyper-Kvasir show improved fidelity, stability, and efficiency over single-method and static-fusion baselines, particularly in clinically ambiguous cases.

Extensive tests on the Hyper-Kvasir dataset show that CAX-Gate consistently outperforms single-method and static fusion baselines across multiple evaluation dimensions, including computational efficiency, explainability fidelity, robustness under uncertainty, and stability under perturbations. Notably, the proposed framework yields the greatest benefits in clinically ambiguous scenarios, where reliable and faithful explanations are most critical. Future research will investigate richer uncertainty representations and expand the framework to encompass additional medical imaging modalities.

Declaration on Generative AI

The author(s) have not employed any Generative AI tools.

References

- [1] C. Steup, K. B. Kennel, M. F. Neurath, S. Fichtner-Feigl, F. R. Greten, Current and emerging concepts for systemic treatment of metastatic colorectal cancer, *Gut* 74 (2025) 2070–2095.
- [2] A. H. Chu, K. Lin, H. Croker, S. Kefyalew, G. Markozannes, K. K. Tsilidis, Y. Park, J. Krebs, M. P. Weijenberg, M. L. Baskin, et al., Dietary-lifestyle patterns and colorectal cancer risk: global cancer update programme (cup global) systematic literature review, *The American Journal of Clinical Nutrition* (2025).
- [3] A. Farhoudian, A. Heidari, R. Shahhosseini, A new era in colorectal cancer: Artificial intelligence at the forefront, *Computers in Biology and Medicine* 196 (2025) 110926.
- [4] A. M. Salih, Z. Raisi-Estabragh, I. B. Galazzo, P. Radeva, S. E. Petersen, K. Lekadir, G. Menegaz, A perspective on explainable artificial intelligence methods: Shap and lime, *Advanced Intelligent Systems* 7 (2025) 2400304.
- [5] R. Talebi, C. A. Celis-Morales, A. Akbari, A. Talebi, N. Borumandnia, M. A. Pourhoseingholi, Machine learning-based classifiers to predict metastasis in colorectal cancer patients, *Frontiers in Artificial Intelligence* 7 (2024) 1285037.
- [6] X. Yang, W. Yu, F. Yang, X. Cai, Machine learning algorithms to predict atypical metastasis of colorectal cancer patients after surgical resection, *Frontiers in Surgery* 9 (2023) 1049933.

- [7] A. Sardari, H. Usefi, Machine learning-based meta-analysis of colorectal cancer and inflammatory bowel disease, *Plos one* 18 (2023) e0290192.
- [8] P. Freitas, F. Silva, J. V. Sousa, R. M. Ferreira, C. Figueiredo, T. Pereira, H. P. Oliveira, Machine learning-based approaches for cancer prediction using microbiome data, *Scientific Reports* 13 (2023) 11821.
- [9] F. Lu, T. Lei, J. Zhou, H. Liang, P. Cui, T. Zuo, L. Ye, H. Chen, J. Huang, Using gut microbiota as a diagnostic tool for colorectal cancer: machine learning techniques reveal promising results, *Journal of Medical Microbiology* 72 (2023) 001699.
- [10] Z. Jiang, L. Yan, S. Deng, J. Gu, L. Qin, F. Mao, Y. Xue, W. Cai, X. Nie, H. Liu, et al., Development and interpretation of a clinicopathological-based model for the identification of microsatellite instability in colorectal cancer, *Disease Markers* 2023 (2023) 5178750.
- [11] H. Galadima, R. Anson-Dwamena, A. Johnson, G. Bello, G. Adunlin, J. Blando, Machine learning as a tool for early detection: A focus on late-stage colorectal cancer across socioeconomic spectrums, *Cancers* 16 (2024) 540.
- [12] L. Buk Cardoso, V. Cunha Parro, S. Verzinhasse Peres, M. P. Curado, G. A. Fernandes, V. Wünsch Filho, T. Natasha Toporcov, Machine learning for predicting survival of colorectal cancer patients, *Scientific reports* 13 (2023) 8874.
- [13] P. J. Rodriguez, P. J. Heagerty, S. Clark, S. Khor, Y. Chen, E. Haupt, E. E. Hahn, V. Shankaran, A. Bansal, Using machine learning to leverage biomarker change and predict colorectal cancer recurrence, *JCO Clinical Cancer Informatics* 7 (2023) e2300066.
- [14] C. S. N. Lam, A. A. Bharwani, E. H. Y. Chan, V. H. Y. Chan, H. L. H. Au, M. K. Ho, S. Rashed, B. M. H. Kwong, W. Fang, K. W. Ma, et al., A machine learning model for colorectal liver metastasis post-hepatectomy prognostications, *Hepatobiliary surgery and nutrition* 12 (2022) 495.
- [15] G. Mohammadi, M. A. Looha, M. A. Pourhoseingholi, M. R. Tavirani, S. Sohrabi, A. Z. S. Khaneh, H. Piri, M. Alaei, N. Parvani, I. Vakilzadeh, et al., Classification and diagnostic prediction of colorectal cancer mortality based on machine learning algorithms: A multicenter national study, *Asian Pacific Journal of Cancer Prevention: APJCP* 25 (2024) 333.
- [16] I. El Badisy, Z. BenBrahim, M. Khalis, S. Elansari, Y. ElHitmi, F. Abbass, N. Mellas, K. El Rhazi, Risk factors affecting patients survival with colorectal cancer in morocco: Survival analysis using an interpretable machine learning approach, *Scientific Reports* 14 (2024) 3556.
- [17] H. Zhong, L. Yang, Q. Zeng, W. Chen, H. Zhao, L. Wu, L. Qin, Q.-Q. Yu, Machine learning predicts the oxidative stress subtypes provide an innovative insight into colorectal cancer, *Oxidative Medicine and Cellular Longevity* 2023 (2023) 1737501.
- [18] A. Ahmadih-Yazdi, A. Mahdavinezhad, L. Tapak, F. Nouri, A. Taherkhani, S. Afshar, Using machine learning approach for screening metastatic biomarkers in colorectal cancer and predictive modeling with experimental validation, *Scientific reports* 13 (2023) 19426.
- [19] H. A. Tahir, W. Alayed, W. U. Hassan, A. Haider, A novel hybrid xai solution for autonomous vehicles: Real-time interpretability through lime-shap integration, *Sensors* 24 (2024) 6776.
- [20] H. Borgli, V. Thambawita, P. H. Smedsrud, S. Hicks, D. Jha, S. L. Eskeland, K. R. Randel, K. Pogorelov, M. Lux, D. T. D. Nguyen, D. Johansen, C. Griwodz, H. K. Stensland, E. Garcia-Ceja, P. T. Schmidt, H. L. Hammer, M. A. Riegler, P. Halvorsen, T. de Lange, HyperKvasir, a comprehensive multi-class image and video dataset for gastrointestinal endoscopy, *Scientific Data* 7 (2020) 283.
- [21] M. T. Ribeiro, S. Singh, C. Guestrin, " 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*, 2016, pp. 1135–1144.
- [22] S. M. Lundberg, S.-I. Lee, A unified approach to interpreting model predictions, in: *Proceedings of the 31st International Conference on Neural Information Processing Systems, NIPS'17*, Curran Associates Inc., Red Hook, NY, USA, 2017, p. 4768–4777.