

A. Cross-Modal Attention over local representations

The pairwise local interaction for WSIs and RNA-Seq is defined as:

$$X_{WSI \leftrightarrow RNA}^L = \frac{1}{X_{WSI}^{L*}} \sum_a (z_a^{WSI} + \sum_b (\gamma_{WSI \leftrightarrow RNA}^{a,b} \cdot z_b^{RNA})) \quad (1)$$

where X_{WSI}^{L*} and X_{RNA}^{L*} are the sets of K key local features of RNA and WSI (i.e., $\{z_j^{WSI} | j = 1, \dots, K\}$ and $\{z_j^{RNA} | j = 1, \dots, K\}$).

$\gamma_{WSI \leftrightarrow RNA}^{a,b}$ is the normalized correlation coefficient computed as $\gamma_{WSI \leftrightarrow RNA}^{a,b} = \frac{\exp((z_a^{WSI})^\top z_b^{RNA})}{\sum_b \exp((z_a^{WSI})^\top z_b^{RNA})}$.

This process is similarly applied for interactions between other modalities at the local level. After computing the pairwise interactions, the final aggregated representations at the local level are obtained using the same methodology as for the global level using Eqs. 5-8 in the main paper.

Dataset	BRCA	NSCLC	RCC
Two Modalities			
Total Patients	IDC: 786 ILC: 197	LUAD: 472 LUSC: 476	KIRC: 510 KIRP: 272 KICH: 66
Total Slides	IDC: 838 ILC: 210	LUAD: 534 LUSC: 510	KIRC: 516 KIRP: 296 KICH: 66
Three Modalities			
Total Patients	IDC: 737 ILC: 190	LUAD: 438 LUSC: 440	KIRC: 498 KIRP: 261 KICH: 65
Total Slides	IDC: 788 ILC: 203	LUAD: 499 LUSC: 474	KIRC: 501 KIRP: 283 KICH: 65

Table 1. Data Statistics for BRCA, NSCLC, and RCC Datasets

B. Feature Selection via Mixture of Experts (MoE) and Top-K Activation

For an input WSI, the local representation is defined as $X_{WSI}^L = \{x_j | j = 1, \dots, C\}$. We define a set of expert networks $\{EXP_k(\cdot)\}_{k=1}^K$, where each expert EXP_k independently processes the local representations. Each expert consists of two connected layers of size 1024 with ReLU activation. A gating network $Gate(\cdot)$ determines the contribution of each expert by assigning selection weights:

$$\alpha_k = Gate_k(x_j), \quad \sum_{k=1}^K \alpha_k = 1, \quad \alpha_k \geq 0. \quad (2)$$

where α_k represents the gating score for expert k , computed using a softmax function over expert logits. To enforce sparsity and encourage selective activation, we apply a Top-K selection mechanism, retaining only the highest-scoring experts per local representation:

$$\mathcal{S}_j = \text{Top-K}(\alpha_1, \alpha_2, \dots, \alpha_K), \quad |\mathcal{S}_j| = K' \quad (3)$$

where $\mathcal{S}_j$ denotes the selected subset of experts, and K' is the number of active experts, which is set to 5. The final expert representation is computed as a weighted sum over the selected experts:

$$z_j = \sum_{k \in \mathcal{S}_j} \alpha_k EXP_k(x_j) \quad (4)$$

To further refine local representations, we apply an activation gating mechanism $Act(\cdot)$ that assigns selection scores ($score_j = Act(z_j)$). We retain only the k most discriminative local features:

$$\mathcal{A} = \text{Top-k}(score_1, score_2, \dots, score_C) \quad (5)$$

where $\mathcal{A}$ denotes the k selected active local representations out of C (i.e. number of local representations per WSI). We set k to 3. The final feature set is obtained via:

$$X_{WSI}^{L*} = Z \odot \mathbf{1}_{\mathcal{A}} = \{z_1^{WSI}, z_2^{WSI}, \dots, z_k^{WSI}\} \quad (6)$$

where $\odot$ represents element-wise multiplication of the set of MoE outputs for local representations, denoted as $Z = \{z_j | j = 1, \dots, C\}$, with an activation mask $\mathbf{1}_{\mathcal{A}}$. This Mixture of Experts with Top-K activation adaptively selects the most informative features while suppressing redundancy, enhancing model efficiency and robustness in multi-modal learning.

