

Highlights

Polarized Message-Passing in Graph Neural Networks

Tiantian He, Yang Liu, Yew-Soon Ong, Xiaohu Wu, Xin Luo

- A novel paradigm named Polarized message-passing (PMP) is proposed to enable graph neural networks (GNNs) to learn more expressive representations.
- Novel GNN architectures, namely PMP graph convolutional network, PMP graph attention network, and PMP graph PageRank network are constructed for various downstream tasks.
- The expressiveness of the PMP-based GNNs is theoretically guaranteed, and the promising experimental results demonstrate the effectiveness of the proposed PMP paradigm.

Polarized Message-Passing in Graph Neural Networks

Tiantian He^{a,b}, Yang Liu^{c,*}, Yew-Soon Ong^{a,d}, Xiaohu Wu^e, Xin Luo^f

^a*Center for Frontier AI Research, Institute of High Performance Computing, Agency for Science, Technology and Research, Singapore*

^b*Singapore Institute of Manufacturing Technology, Agency for Science, Technology and Research, Singapore*

^c*Department of Computer Science, Hong Kong Baptist University, Hong Kong SAR*

^d*School of Computer Science and Engineering, Nanyang Technological University, Singapore*

^e*National Engineering Research Center of Mobile Network Technologies, Beijing University of Posts and Telecommunications, China PR*

^f*College of Computer and Information Science, Southwest University, China PR*

Abstract

In this paper, we present Polarized message-passing (PMP), a novel paradigm to revolutionize the design of message-passing graph neural networks (GNNs). In contrast to existing methods, PMP captures the power of node-node similarity and dissimilarity to acquire dual sources of messages from neighbors. The messages are then coalesced to enable GNNs to learn expressive representations from sparse but strongly correlated neighbors. Three novel GNNs based on the PMP paradigm, namely PMP graph convolutional network (PMP-GCN), PMP graph attention network (PMP-GAT), and PMP graph PageRank network (PMP-GPN) are proposed to perform various downstream tasks. Theoretical analysis is also conducted to verify the high expressiveness of the proposed PMP-based GNNs. In addition, an empirical study of five learning tasks based on 12 real-world datasets

*Corresponding author.

is conducted to validate the performances of PMP-GCN, PMP-GAT, and PMP-GPN. The proposed PMP-GCN, PMP-GAT, and PMP-GPN outperform numerous strong message-passing GNNs across all five learning tasks, demonstrating the effectiveness of the proposed PMP paradigm.

Keywords: Graph neural networks, Message-passing graph neural networks, Representation learning, Graph analysis

1. Introduction

Message-passing graph neural networks (MPGNNs) [1, 2] are prominent tools for analyzing graph-structured data. MPGNNs heavily rely on a two-stage message-passing paradigm to learn representations for diverse downstream tasks. In the first stage, messages (i.e., node features) are conveyed to each central node from all of its neighbors. Simultaneously, the correlation (similarity) of each neighbor to the central node is computed, typically as weights, such as attention coefficients [3, 4] or transition probabilities [5], and transmitted to the central node. In the second stage, the representation of each central node is generated through the weighted aggregation of received messages.

Diverse properties of real-world graphs, such as the friendship paradox [6] in social graphs and impact imbalance [7] in citation graphs, have demonstrated that many nodes are created differently. However, the message-passing paradigms followed by most MPGNNs focus on conveying similarity-based messages, that is, the representation is learned through the aggregation of the features of neighbors, in which the relevant neighbors are determined by some similarity measures. Properties regarding dissimilarity in the graphs

19 are consequently under-exploited by existing MPGNNs.

20 In this paper, we hypothesize that fully exploiting the properties of dis-
21 similarity present in graph data can significantly enhance the learning per-
22 formance of MPGNNs. To this end, we propose Polarized message-passing
23 (PMP) for graph neural networks (GNNs). Unlike the learning paradigm
24 adopted by conventional MPGNNs, PMP allows for the concurrent prop-
25 agation of similarity- and dissimilarity-based messages. By leveraging the
26 composite effect of both similarity and dissimilarity learned from graph data,
27 PMP is capable of preserving, reducing, or negating the messages conveyed
28 from neighboring nodes. PMP-based GNNs can learn more expressive rep-
29 resentations by largely concentrating on a significantly reduced number of
30 strongly correlated neighbors.

31 To develop the PMP paradigm, we introduce the following two technical
32 contributions. First, the weight matrix used by the layers in conventional
33 MPGNNs to aggregate the features from neighboring nodes is decomposed
34 into two separate learnable matrices. The first matrix can capture the cor-
35 relations between each central node and its neighbors (e.g., node-node sim-
36 ilarities regarding features). This matrix can reflect how the messages from
37 neighboring nodes should be preserved. The second matrix can acquire the
38 difference between nodes (i.e., the exponential of negative node-node dis-
39 tances), indicating whether the messages from neighbors should be reduced
40 or negated.

41 Second, we construct a new sparse matrix that enables the composite
42 propagation of similarity- and dissimilarity-based messages (i.e., polarized
43 messages) in GNNs. The corresponding entries of this matrix are calculated

44 as the normalized products of the matrices pertaining to similarity and dis-
45 similarity. With the newly constructed sparse matrix, PMP can learn to
46 preserve, reduce, or negate the messages. More expressive representations
47 can be learned through the sparse aggregation of the messages from neigh-
48 bors.

49 The main contributions of this paper can be summarized as follows.

- 50 • We present a novel learning paradigm, named PMP, for GNNs. In
51 contrast to conventional message-passing paradigms, PMP allows the
52 concurrent propagation of similarity- and dissimilarity-based messages,
53 enabling GNNs to learn to preserve, reduce, or negate the messages
54 conveyed by neighbors.
- 55 • We propose three novel PMP-based GNNs, namely PMP graph convo-
56 lutional network (PMP-GCN), PMP graph attention network (PMP-
57 GAT), and PMP graph PageRank network (PMP-GPN). The GNNs
58 are designed for learning more expressive representations for various
59 downstream tasks by largely concentrating on sparse but strongly cor-
60 related neighbors.
- 61 • We conduct a comprehensive theoretical analysis, showing that the
62 proposed PMP-GCN, PMP-GAT, and PMP-GPN are more expressive
63 than representative MPGNNs.
- 64 • The proposed PMP-based GNNs are tested on five learning tasks,
65 namely scientific article classification, blog and content-based image
66 clustering, coauthor analysis, and rich-text classification from 12 real-
67 world datasets. The notable results demonstrate the effectiveness of

68 the proposed PMP.

69 2. The Polarized message-passing paradigm

70 In this section, the PMP paradigm is first elaborated. The theoretical
71 analysis demonstrating that the proposed PMP paradigm is more expressive
72 than those adopted by conventional MPGNNs is then presented.

73 2.1. Notations

74 In this paper, a graph containing N nodes and $|E|$ edges is denoted as
75 $G = \{V, E, \mathbf{X}\}$, where V , E , and $\mathbf{X} \in \mathbb{R}^{N \times D}$ represent the node and edge sets
76 and the input node feature matrix, respectively. Each node in G belongs to
77 one of the C classes, and their one-hop neighbors form a set denoted as $\mathcal{N}_i$.
78 $\mathbf{A} \in \{0, 1\}^{N \times N}$ denotes the adjacency matrix of G . $\mathbf{W}^l$, $\mathbf{H}^l$, and $\mathbf{P}^l$ represent
79 the matrices of learnable weights, node features, and latent positions at the
80 l th layer. Thus, $\mathbf{H}^1 = \mathbf{X}$.

81 2.2. Formulating Polarized message-passing at a single layer

82 At each layer of conventional MPGNNs, the output representations for
83 each node are generated through the addition of the node feature and the
84 messages received from its neighbors [8, 9]. The message from each neighbor
85 is typically formulated through the multiplication of the neighboring fea-
86 ture and the weight obtained via a similarity/correlation strategy. However,
87 many phenomena frequently observed in real-world graphs, e.g., the friend-
88 ship paradox, have demonstrated that many nodes in the graph are created
89 differently. These differences in the characteristics of nodes in real-world
90 graphs, which are contrary to the notion of similarity widely considered by

existing GNNs, have not been well explored. Therefore, in this paper, we propose PMP, which can leverage the composite effect of both similarity and dissimilarity to learn to preserve, reduce, and negate the messages from neighbors. PMP-based GNNs can learn more expressive representations from a significantly reduced number of correlated nodes in the graph.

Different from existing message-passing strategies used by conventional GNNs, PMP constructs two learnable matrices, $\mathbf{C}$ and $\mathbf{S}$, to quantify the correlation and differentiation between pairs of nodes, respectively. A novel matrix $\mathbf{M}$ is then constructed through the combination of the joint effect of $\mathbf{C}$ and $\mathbf{S}$. Finally, polarized messages from neighbors are composed through the multiplication of $\mathbf{M}$ -induced weights and the corresponding features of neighbors.

Computing $\mathbf{C}$ and $\mathbf{S}$. To quantify the node correlations ($\mathbf{C}$) at a layer, PMP computes the feature and structural similarities between pairwise nodes that are connected:

$$\begin{aligned}\mathbf{H} &= \mathbf{H}^l \mathbf{W}_h^l, \mathbf{P} = \mathbf{P}^l \mathbf{W}_p^l, \\ \mathbf{C}_{ij} &= \mathbf{A}_{ij} \cdot [\mathbf{a}(\mathbf{H}_{i,:} || \mathbf{H}_{j,:})^T + \mathbf{P}_{i,:} \text{diag}(\mathbf{b}) \mathbf{P}_{j,:}^T],\end{aligned}\tag{1}$$

where $\mathbf{H}_{i,:}$ and $\mathbf{P}_{i,:}$ are row vectors representing the feature and latent positions of node i , $||$ is the operator for vector concatenation, $\text{diag}(\cdot)$ is the diagonal matrix, and $\mathbf{a}$ and $\mathbf{b}$ are vectors of learnable parameters. To quantify the node dissimilarities, PMP computes the weighted distances in terms of feature and structure:

$$\mathbf{S}_{ij} = \mathbf{A}_{ij} \cdot [\alpha |\mathbf{H} \mathbf{W}_{i,:} - (\mathbf{H} \mathbf{W})_{j,:}|^2 + (1 - \alpha) |\mathbf{P}_{i,:} - \mathbf{P}_{j,:}|^2],\tag{2}$$

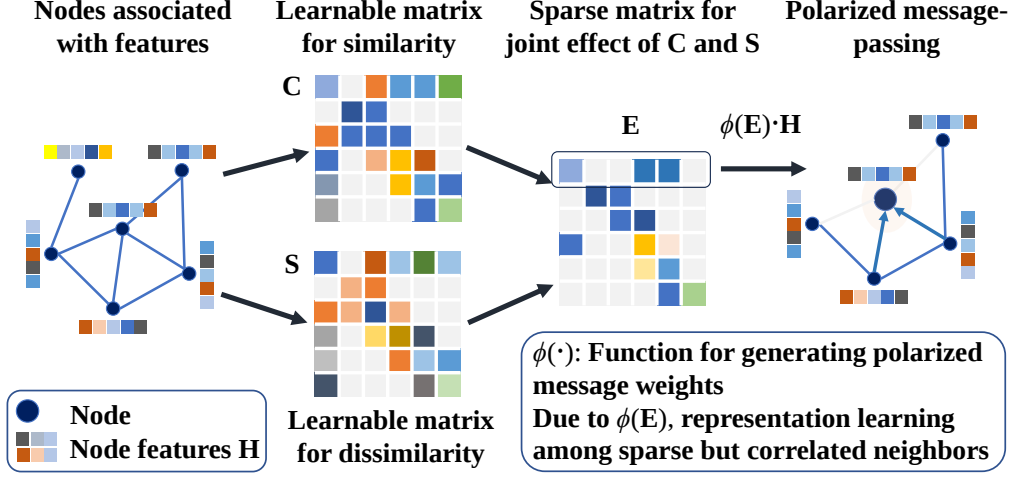


Figure 1: Graphical view of the proposed polarized message-passing (PMP) paradigm. By appropriately merging neighboring similarity- and dissimilarity-based messages, PMP allows GNNs to learn more expressive representations with sparse but strongly correlated neighbors.

where $\alpha \in (0, 1)$ is a learnable parameter balancing the relative significance of distances caused by the node feature and the graph structure.

Coalescing polarized messages for representation learning. To learn the output representation at each layer, the polarized messages can be directly generated according to the weights induced by C and S . However, this procedure is computationally demanding as it requires double computation, that is, the multiplication of the weights induced by C and S and the features of all neighbors. In this paper, we propose an alternative procedure to efficiently generate polarized messages. The joint effect of C and S is first computed. The output representations can then be learned using messages and weights derived from the joint effect. To achieve this, we introduce a new sparse matrix denoted as E , which is derived from C and S . The derived E is utilized

123 to generate polarized messages, which play a crucial role in the learning of
 124 output representations. Specifically, $\mathbf{E}$ is the element-wise exponential of the
 125 difference between $\mathbf{C}$ and $\mathbf{S}$. This straightforward strategy for the derivation
 126 of $\mathbf{E}$ (i.e., taking the element-wise exponential of the difference between $\mathbf{C}$
 127 and $\mathbf{S}$) has been found to perform quite well in various downstream tasks, al-
 128 though complicated alternative strategies [10] are available. With $\mathbf{E}$ at hand,
 129 PMP can be established for a graph neural layer (See Fig. 1 for a graph-
 130 ical view). The output representations for each node (say i) are acquired
 131 following the generic procedure presented below:

$$\text{FEATURE MAPPING: } \mathbf{H} = \mathbf{H}^l \mathbf{W}_h^l, \mathbf{P} = \mathbf{P}^l \mathbf{W}_p^l,$$

$$\text{POLARIZED MESSAGE: Compute } \mathbf{C} \text{ and } \mathbf{S},$$

$$\mathbf{E}_{ij} = \mathbf{A}_{ij} \cdot \exp(\mathbf{C}_{ij} - \beta \mathbf{S}_{ij}), \quad (3)$$

$$\mathbf{M}_i = \sum_{j \in \mathcal{N}_i} \phi(\mathbf{E}_{ij}) \cdot \mathbf{H}_{j,:},$$

$$\text{FEATURE UPDATE: } \mathbf{H}_{i,:}^{l+1} = \frac{\theta}{|\mathcal{N}_i|} \mathbf{H}_{i,:} + \mathbf{M}_i,$$

132 where $\phi(\cdot)$ is an appropriately function leveraging $\mathbf{E}$ to generate the weights
 133 for computing messages from neighbors, β is a learnable positive parameter,
 134 and θ is a learnable positive irrational that is independent of the learning
 135 of M_i . Given Eq. (3), the proposed PMP paradigm for GNNs substantially
 136 differs from those used in existing MPGNNs in the following three aspects.
 137 First, the messages generated by $\mathbf{E}$ are more comprehensive, as $\mathbf{E}$ considers
 138 the joint effect of similarities and dissimilarities between nodes in the graph.
 139 Second, sparse messages from neighbors can be learned with $\mathbf{E}$, as it considers
 140 the differences between correlations and dissimilarities. Lastly, the proposed
 141 PMP is highly flexible for constructing diverse graph neural layers, such as

graph convolution layers and attention layers, by properly setting $\phi(\cdot)$. This enables PMP-based GNNs to effectively tackle a wide range of downstream tasks.

2.3. Expressiveness of Polarized message-passing

As shown in Eq. (3), capturing the joint effect of similarity and dissimilarity between connected nodes for the composition of polarized messages is a core step of PMP. This step endows PMP-based GNNs with the capability of learning more expressive representations, regardless of the selection of feature aggregation strategy. Here, a qualitative analysis is conducted to demonstrate such superior capability possessed by PMP-based GNNs, compared with existing MPGNNs.

Learning tasks performed by GNNs typically concern predicting the labels of test data samples, given the representations learned from the training set. Let the ground-truth label be Y . Given a learning task, a GNN performs it by conceptually maximizing the following conditional likelihood $p(Y|\mathbf{z})$, where $\mathbf{z}$ is the representation learned by a GNN. We know that conventional MPGNNs learn representations mainly with similarity/correlation-based messages. To learn the representation $\mathbf{z}$, a conventional MPGNN has to utilize the message set generated for each central node in the graph. Thus, for node i , its representation can be represented as $\mathbf{z}_i = f_\psi(C_i)$, where C_i denotes the set of similarity-based messages for node i and $f_\psi(\cdot)$ represents the function for representation learning. And the objective function for representation learning can be conceptually formulated as follows:

$$\min_{\psi} \mathbb{E}_{i \sim I_{train}, C_i \sim p(C)} [L(Y_i, f_\psi(C_i))], \quad (4)$$

165 where L represents a function measuring the divergence between Y_i and $f_\psi(\cdot)$.

166 Given Eq. (3), we know that PMP-based GNNs leverage polarized mes-
 167 sages to learn representations for downstream tasks. For node i , the represen-
 168 tation learned by a PMP-based GNN differs from that learned by a conven-
 169 tional MPGNN: $\mathbf{z}_i = f_\psi(C_i \oplus S_i)$, where S_i denotes the set of dissimilarity-
 170 based messages for node i , and $\oplus$ represents the strategy of merging polarized
 171 messages. Thus, the objective function for a PMP-based GNN can be for-
 172 mulated as follows:

$$\min_{\psi} \mathbb{E}_{i \sim I_{train}, C_i \sim p(C), S_i \sim p(S)} [L(Y_i, f_\psi(C_i \oplus S_i))]. \quad (5)$$

173 Optimizing Eqs. (4) and (5) produces node representations drawn from the
 174 distribution of C and the joint distribution of C and S , respectively. Thus,
 175 analyzing the mutual information between ground-truth labels and learned
 176 representations in Eqs. (4) and (5) can quantify the expressiveness of the pro-
 177 posed PMP and paradigms adopted in conventional MPGNNs. Let $f_\psi(\cdot)$ be
 178 a function for representation learning. The following theorem demonstrates
 179 that a GNN equipped with the proposed PMP in Eq. (3) is more expressive
 180 than conventional MPGNNs.

181 **Theorem 1.** *Suppose $f_\psi(\cdot)$ is a learnable function for representation learn-*
 182 *ing adopted by PMP-based GNNs and conventional MPGNNs, $C = \{C_i | i =$*
 183 *$1, \dots, n\}$, and $S = \{S_i | i = 1, \dots, n\}$ denote the sets of similarity and dissimi-*
 184 *larity messages that are used to learn node representations, and $Y = \{Y_i | i =$*
 185 *$1, \dots, n\}$ denotes the set of node labels, where n denotes the total number of*
 186 *nodes for training. The representations learned by PMP-based GNNs and*
 187 *conventional MPGNNs are denoted as $\mathbf{z}_i^{PMP} = f_\psi(C_i \oplus S_i), i = 1, \dots, n$, and*

188 $\mathbf{z}_i^{MP} = f_\psi(C_i), i = 1, \dots, n$, respectively. For any $i \in \{1, \dots, n\}$, the following
 189 inequality holds:

$$I(Y_i; \mathbf{z}_i^{PMP}) \geq I(Y_i; \mathbf{z}_i^{MP}), \quad (6)$$

190 where $I(\cdot; \cdot)$ represents the function of mutual information.

191 We leave all proofs in Section 4 to keep the content uncluttered. The-
 192 orem 1 demonstrates that PMP-based GNNs are more expressive in learn-
 193 ing representations than conventional MPGNNs, as the mutual information
 194 between representations learned by PMP-based GNNs and ground-truth la-
 195 bels is higher than that between representations learned by conventional
 196 MPGNNs and ground-truth labels. Theorem 1 provides an information-
 197 theoretic analysis of the expressiveness gap between the proposed PMP and
 198 the message-passing paradigms used by conventional MPGNNs. Leveraging
 199 the joint effect of similarity- and dissimilarity-based messages enables PMP-
 200 based GNNs to be more likely to correctly predict the data labels, which
 201 indicates a significant performance gain of PMP-based GNNs against con-
 202 ventional MPGNNs.

203 **3. Polarized message-passing graph neural networks**

204 Using the proposed PMP paradigm (Eq. (3)), we can build GNNs (i.e.,
 205 PMP-GNNs) to effectively perform various downstream tasks. As previously
 206 mentioned, the proposed PMP paradigm offers high flexibility in constructing
 207 diverse GNNs by appropriately setting $\phi(\cdot)$. In this section, we present three
 208 novel GNNs, namely PMP-GCN, PMP-GAT, and PMP-GPN. These three
 209 GNNs are designed to evaluate the efficacy of the proposed PMP paradigm, as

demonstrated in the experiments. Moreover, our theoretical analysis demonstrates that the proposed PMP-GNNs are more expressive than conventional MPGNNs.

3.1. Polarized message-passing graph convolutional network

Graph convolutional networks (GCNs) [1, 11] have been recognized as prominent MPGNNs and are used to tackle a wide range of predictive tasks in graph-structured data. However, most existing GCNs do not consider leveraging polarized messages to learn representations. Here, we present PMP-GCN, whose layers adopt the proposed message-passing paradigm to learn more expressive representations. The output representation is learned through the following procedure for each layer in PMP-GCN:

$$\begin{aligned}
\mathbf{H} &= \mathbf{H}^l \mathbf{W}_h^l, \mathbf{P} = \mathbf{P}^l \mathbf{W}_p^l, \mathbf{E}_{ij} = \mathbf{A}_{ij} \cdot \exp(\mathbf{C}_{ij} - \beta \mathbf{S}_{ij}), \\
\mathbf{M}_i &= \frac{1}{\sqrt{\sum_{k \in \mathcal{N}_i} \mathbf{E}_{ik}}} \sum_{j \in \mathcal{N}_i} \frac{\mathbf{E}_{ij}}{\sqrt{\sum_{k \in \mathcal{N}_j} \mathbf{E}_{jk}}} \cdot \mathbf{H}_{j,:}, \\
\mathbf{H}_{i,:}^{l+1} &= \frac{\theta}{|\mathcal{N}_i|} \mathbf{H}_{i,:} + \mathbf{M}_i.
\end{aligned} \tag{7}$$

According to Eq. (7), the following two aspects differentiate the proposed PMP-GCN from conventional GCNs. First, the weights utilized by PMP-GCN to compose messages for representation learning (i.e., the graph Laplacian) are learnable rather than fixed as in most conventional GCNs. This enables PMP-GCN to effectively capture the structural information at each layer, thereby enhancing its learning capability. Second, as $\mathbf{E}_{ij}$ is computed using $\mathbf{C}$ and $\mathbf{S}$, neighbors considerably different from the central node are assigned very low weights. The importance of subsequent messages can be reduced or even negated during the generation of output representations.

Thus, the convolutional operator in PMP-GCN can learn representations from the preserved messages from strongly correlated neighbors.

3.2. Polarized message-passing graph attention network

Graph attention networks (GATs) [3, 12] are also typical MPGNNs. At each layer of a conventional GAT, the weights used to generate messages are always computed solely based on node correlations. To endow attention-based GNNs with the capability of learning representations with polarized messages, we present PMP-GAT. For each layer in PMP-GAT, the output representation for each node is learned via the following procedure:

$$\begin{aligned}\mathbf{H} &= \mathbf{H}^l \mathbf{W}_h^l, \mathbf{P} = \mathbf{P}^l \mathbf{W}_p^l, \mathbf{E}_{ij} = \mathbf{A}_{ij} \cdot \exp(\mathbf{C}_{ij} - \beta \mathbf{S}_{ij}), \\ \mathbf{M}_i &= \sum_{j \in \mathcal{N}_i} \frac{\mathbf{E}_{ij}}{\sum_{k \in \mathcal{N}_i} \mathbf{E}_{ik}} \cdot \mathbf{H}_{j,:}, \\ \mathbf{H}_{i,:}^{l+1} &= \frac{\theta}{|\mathcal{N}_i|} \mathbf{H}_{i,:} + \mathbf{M}_i.\end{aligned}\tag{8}$$

Compared with conventional GATs, PMP-GAT pays less attention to or even ignores dissimilar neighbors when aggregating features from neighbors, owing to the effect of $\mathbf{E}$. This is equivalent to reducing or negating the messages conveyed by neighbors during the generation of output representations. PMP-GAT, therefore, can learn representations by focusing on messages from truly correlated neighbors.

3.3. Polarized message-passing graph PageRank network

Graph PageRank networks (GPNs), such as APPNP [13] and GPR [14], are also powerful MPGNNs. At each layer, the representation is typically generated by summing over the initial features of each central node and features computed by conventional graph convolution. Most PageRank GNNs

do not consider leveraging polarized messages to learn representations. Here, we propose PMP-GPN to endow Graph PageRank networks with the capability to learn representations with polarized messages. PMP-GPN is mainly built upon APPNP. For each layer of PMP-GPN, the output representation of each node is generated through the following procedure:

$$\begin{aligned}
\mathbf{H}^{in} &= MLP(\mathbf{X}), \mathbf{P}^{in} = MLP(\mathbf{P}^0), \mathbf{E}_{ij} = \mathbf{A}_{ij} \cdot \exp(\mathbf{C}_{ij} - \beta \mathbf{S}_{ij}), \\
\mathbf{M}_i &= \frac{1}{\sqrt{\sum_{k \in \mathcal{N}_i} \mathbf{E}_{ik}}} \sum_{j \in \mathcal{N}_i} \frac{\mathbf{E}_{ij}}{\sqrt{\sum_{k \in \mathcal{N}_j} \mathbf{E}_{jk}}} \cdot \mathbf{H}_{j,:}, \\
\mathbf{H}_{i,:}^{l+1} &= \frac{\theta}{|\mathcal{N}_i|} \mathbf{H}_{i,:} + (1 - \alpha) \mathbf{M}_i + \alpha \mathbf{H}_{i,:}^{in},
\end{aligned} \tag{9}$$

where $\mathbf{H}^{in}$ and $\mathbf{P}^{in}$ are the node features and latent positions learned by multi-layer perception, and α is a small positive number representing the teleport probability used to control the neighborhood influence. In PMP-GPN, all layers share the same $\mathbf{E}$, which is computed using $\mathbf{H}^{in}$ and $\mathbf{P}^{in}$, for efficient purposes. In contrast to existing PageRank GNNs, the weights for composing messages for representation learning are learnable rather than fixed. Similar to PMP-GCN, neighbors that significantly differ from the central node are assigned very low weights, leading to very limited or no influence on the subsequent generation of output representations. Thus, PMP-GPN can learn representations by concentrating on multi-hop message-passing among strongly correlated neighbors.

3.4. Loss function

After message passing over multiple layers, PMP-GNNs learn the output representations, which are finally fed into the training loss to optimize the

network parameters. Two modules constitute the loss function of PMP-GNNs, including task-specific loss and latent-position loss. The loss function can be conceptually formulated as follows:

$$L = L_{task} + tr(\mathbf{P}^{outT} \mathbf{L} \mathbf{P}^{out}), \quad (10)$$

where $\mathbf{L}$ is the Laplacian matrix of $\mathbf{A}$, and $\mathbf{P}^{out}$ represents the output latent positions of all nodes in the graph.

3.5. Computational complexity of PMP-GNNs

The computational complexity of the proposed PMP-GNNs is similar to that of conventional MPGNNs. Assuming the input and output feature dimensions are d and d' , for each layer in the PMP-GNN, the complexity of feature mapping is approximately $\mathcal{O}(2Ndd')$. The complexity of computing $\mathbf{C}$ and $\mathbf{S}$ is about $\mathcal{O}(4|E|d')$. The complexity of composing polarized messages and learning representations is approximately $\mathcal{O}(|E|(1+d'))$. Therefore, the overall complexity of each layer in PMP-GCN and PMP-GAT is about $\mathcal{O}(|E|(1+5d')+2Ndd')$. The complexity of each layer in PMP-GPN is approximately $\mathcal{O}(|E|d)$ as PMP-GPN completes feature mapping and computing $\mathbf{C}$ and $\mathbf{S}$ before learning representations (i.e., $d = d'$). Considering the complexity of conventional GCN, GAT and APPNP are about $\mathcal{O}(2|E|d' + Ndd')$, $\mathcal{O}(4|E|d' + Ndd')$, and $\mathcal{O}(|E|d)$, the complexity of the proposed PMP-GNNs allows them to perform various learning tasks on massive datasets.

3.6. Expressiveness of PMP-GNNs

The expressiveness of GNNs [15–17] has drawn considerable attention recently. It concerns whether a GNN layer can acquire distinct representations

for diverse (sub-) graph structures. Previous studies have demonstrated that GNNs with greater expressiveness can perform better in downstream learning tasks [18, 19]. Accordingly, a comprehensive analysis is conducted here to demonstrate the greater expressiveness of the proposed PMP-GCN, PMP-GAT, and PMP-GPN compared to conventional MPGNNs.

For all MPGNNs, the upper limit of expressiveness is the one-dimensional Weisfeiler-Lehman test (1-WL test) [18, 20]. Similar to what the 1-WL test performs in the Weisfeiler-Lehman algorithm, this test verifies whether an aggregation function in a GNN layer can successfully differentiate all distinct (sub-) graph structures with node features. Therefore, analyzing the expressiveness of PMP-GNNs lies in whether the process of feature aggregation for representation learning can discriminate all different graph structures. In this paper, we derive the following theorem to establish the expressiveness of the proposed PMP-GNNs when they learn representations with different graph structures.

Theorem 2. *Assume the node features of two graphs are denoted as two multi-sets $X_i = (M_i, \mu_i)$ and $X_j = (M_j, \mu_j)$, $X_i, X_j \in \mathcal{X}$, where $\mathcal{X}$ represent the countable feature space, M_i is the underlying set containing the distinct elements in X_i , and $\mu_i \in \mathbb{N}^*$ stands for the number of repeating times of each element in M_i , c_i and c_j represent the features of the central nodes in these two graphs. The process of feature aggregation for c_i and c_j at each layer of PMP-GCN, PMP-GAT, and PMP-GPN is denoted as $h(c_i, X_i)$ and $h(c_j, X_j)$. For any two graphs with distinct structures, the following inequality holds:*

$$h(c_i, X_i) \neq h(c_j, X_j). \quad (11)$$

Theorem 2 shows that the proposed PMP-GNNs can discriminate all di-

verse (sub-) graph structures characterized by distinct node features. This implies that PMP-GNNs are more expressive than most MPGNNs, whose aggregation functions do not meet the criteria of the 1-WL test. The strong expressiveness of the proposed PMP-GNNs suggests they can perform superiorly in diverse downstream tasks.

4. Theoretical analysis

In this section, all theorems revealing the superior expressiveness of the proposed PMP and PMP-GNNs are proved.

4.1. Proof of Theorem 1

Proof. To prove Theorem 1, the following properties of mutual information and entropy will be used:

$$\begin{aligned} I(x; y) &= H(y) - H(y|x), \\ I(x; z|y) &= H(x|y) - H(x|y, z), \\ H(x, y) &= H(x|y) + H(y), \end{aligned} \tag{12}$$

where $I(\cdot; \cdot)$ represents the mutual information between random variables such as x , y , and z , $H(\cdot)$ represents the marginal entropy, and $H(\cdot|\cdot)$ represents the conditional entropy. Let $I(Y_i; \mathbf{z}_i^{PMP})$ be the mutual information between the ground-truth label Y_i and the representation $\mathbf{z}_i^{PMP} = f_\psi(C_i \oplus S_i)$ that a PMP-GNN learns from the polarized message sets of C_i and S_i . Given Eq. (3), we know that the proposed PMP first captures the joint effect of C_i and S_i , and then leverages the shared function $\phi(\cdot)$ to compose messages for representation learning. The above procedure is equivalent to directly leveraging C_i and S_i to formulate dual messages for the subsequent

335 learning of representations. Thus, $I(Y_i; f_\psi(C_i \oplus S_i))$ can be rewritten as
 336 $I(Y_i; f_\psi(C_i), f_\psi(S_i))$, where $(f_\psi(C_i)$ and $f_\psi(S_i))$ are drawn from a joint dis-
 337 tribution. Given $I(Y_i; f_\psi(C_i), f_\psi(S_i))$, we have:

$$\begin{aligned}
 I(Y_i; f_\psi(C_i), f_\psi(S_i)) &= H(f_\psi(C_i), f_\psi(S_i)) - H(f_\psi(C_i), f_\psi(S_i)|Y_i), \\
 &= H(f_\psi(C_i), f_\psi(S_i)) + H(Y_i) - H(f_\psi(C_i), f_\psi(S_i), Y_i) \\
 &= H(Y_i) - H(Y_i|f_\psi(C_i), f_\psi(S_i)) \\
 &= H(Y_i) - H(Y_i|f_\psi(C_i)) + H(Y_i|f_\psi(C_i)) - H(Y_i|f_\psi(C_i), f_\psi(S_i)) \\
 &= I(Y_i; f_\psi(C_i)) + I(Y_i; f_\psi(S_i)|f_\psi(C_i)) \geq I(Y_i; f_\psi(C_i)),
 \end{aligned} \tag{13}$$

338 where $I(Y_i; f_\psi(C_i))$ is the mutual information between Y_i and the represen-
 339 tation learned by conventional MPGNNs ($\mathbf{z}_i^{MP}$). Mutual information is non-
 340 negative, so the above inequality always holds for any $f_\psi(\cdot)$. Since dissimi-
 341 larities widely exist in graph data and both C_i and S_i are determined by the
 342 same set of node features, $I(Y_i; f_\psi(S_i)|f_\psi(C_i))$ is probably larger than zero,
 343 leading to a greater gap between $I(Y_i; f_\psi(C_i), f_\psi(S_i))$ and $I(Y_i; f_\psi(C_i))$ for
 344 any node in the graph. $\square$

345 The proof of Theorem 1 demonstrates that PMP-GNNs are constantly
 346 more likely to correctly predict node labels in the graph than conventional
 347 MPGNNs as long as $f_\psi(\cdot)$ adopted by PMP-GNN and MPGNNs is well
 348 defined.

349 4.2. Proof of Theorem 2

350 Before proving Theorem 2, we first introduce necessary notations and
 351 preliminaries. For each central node and its neighbors, their features form
 352 a multi-set $X = (M, \mu)$, where $X \in \mathcal{X}$, $\mathcal{X}$ is a countable feature space,

353 M is the underlying set containing the distinct elements in X , and $\mu \in$
 354 $\mathbb{N}^*$ stands the number of repeating times of each element in M . To learn
 355 representations in each layer, a PMP-GNN (i.e., PMP-GCN, PMP-GAT,
 356 and PMP-GPN) aggregates the polarized messages from neighbors according
 357 to the message weights. For each central node in the graph, this learning
 358 procedure is conceptually represented as follows:

$$\begin{aligned}
 \mathbf{E}_{cx} &= \exp \{a_{cx} + b_{cx} - \beta(k_{cx} + l_{cx})\}, \\
 w_{cx} &= \phi(\mathbf{E}_{cx}), \\
 h(c, X) &= \frac{\theta}{|X|} g(c) + \sum_{x \in M} \sum_{\substack{y=x, \\ y \in X}} w_{cy} g(y),
 \end{aligned} \tag{14}$$

359 where c and x represent the features of the central node and neighbor, a_{cx} and
 360 b_{cx} represent the feature and structural similarities between the central node
 361 and a neighbor, k_{cx} and l_{cx} represent the feature and structural dissimilarities
 362 between the central node and a neighbor, $\phi(\cdot)$ is the function for computing
 363 the message weight, θ is a positive irrational independent of the message
 364 weights, $|X|$ is the size of the multi-set, and $g(\cdot)$ represents a valid mapping
 365 function in the countable feature space, respectively.

366 4.2.1. Schematic methodology to complete the proof of Theorem 2

367 To prove Theorem 2, that is, the learning procedure in Eq. (14) can
 368 discriminate all distinct graph structures, we have to validate that the learned
 369 representations for each pair of different graphs are distinct:

$$h(c_i, X_i) \neq h(c_j, X_j), \text{ if } |X_i| \neq |X_j|, \forall i \neq j. \tag{15}$$

370 In general, Theorem 2 can be verified through *proof by contradiction*, which
 371 is a formal framework widely used to verify whether a GNN layer can pass

the 1-WL test [16, 18]. Accordingly, the generic procedure for completing the proof can be described as follows. First, we can assume the statement shown in Eq. (15) is FALSE. In other words, a layer in each PMP-GNN learns an identical representation for two central nodes with different structures. Second, we identify all possible conditions that determine the aggregation function ($h(c, X)$) of the PMP-GNN to generate the representation and assume the statement shown in Eq. (15) is FALSE when these conditions are met. At last, we can complete the proof by demonstrating that when all the possible conditions are met, the assumption that Eq. (15) is FALSE can result in some easy-to-verify contradictions.

Given Eqs. (7)-(9), and (14), we can easily identify that the following two factors determine how $h(c, X)$ generates representations at each PMP-GNN layer, including the underlying set M , and the feature of the central node c . Thus, the representations of a pair of central nodes generated by each PMP layer can be analyzed in the following four cases, including $M_1 \neq M_2$ and $c_1 \neq c_2$, $M_1 \neq M_2$ but $c_1 = c_2$, $M_1 = M_2$ but $c_1 \neq c_2$, and $M_1 = M_2$ and $c_1 = c_2$. In what follows, we will prove that all the proposed PMP-GNNs can pass the 1-WL test by showing that the assumption that Eq. (15) is FALSE can result in contradictions in all the previously enumerated cases.

4.2.2. The equivalence between the PMP-GCN layer and the 1-WL test

Proof. We can complete the proof by following the procedure introduced in Section 4.2.1. Let us first assume the statement regarding PMP-GCN in Theorem 2 is FALSE, that is, there exists a pair of graph structures X_1 and X_2 with central node features c_1 and c_2 , that the PMP-GCN layer can NOT

discriminate. Based on Eq. (14), this assumption can be written as follows:

$$\begin{aligned} \mathbf{E}_{c_i x} &= \exp \{a_{c_i x} + b_{c_i x} - \beta(k_{c_i x} + l_{c_i x})\}, \\ w_{c_i x} &= \phi(\mathbf{E}_{c_i x}) = \frac{\mathbf{E}_{c_i x}}{\sqrt{\sum_{x \in X_i} \mathbf{E}_{c_i x} \cdot \sum_{z \in X_y} \mathbf{E}_{yz}}}, \\ h(c_1, X_1) &= h(c_2, X_2), |X_1| \neq |X_2|, \end{aligned} \quad (16)$$

where $h(\cdot)$ is defined as Eq. (7) shows. As mentioned previously, we know that the representations for c_1 and c_2 generated by a PMP-GCN layer can be analyzed in four different cases, including $M_1 \neq M_2$ and $c_1 \neq c_2$, $M_1 \neq M_2$ but $c_1 = c_2$, $M_1 = M_2$ but $c_1 \neq c_2$, and $M_1 = M_2$ and $c_1 = c_2$. We will then demonstrate that $h(c_1, X_1) = h(c_2, X_2)$ can not hold as possible contradictions are identified under all these four cases.

$M_1 \neq M_2$ and $c_1 \neq c_2$, or $M_1 \neq M_2$ but $c_1 = c_2$. Here, we analyze the representations generated by a PMP-GCN layer when $M_1 \neq M_2$. If $h(c_1, X_1) = h(c_2, X_2)$, we have:

$$\begin{aligned} h(c_1, X_1) &= \frac{\theta}{|X_1|} g(c_1) + \sum_{x \in M_1} \sum_{\substack{y=x, \\ y \in X_1}} w_{c_1 y} g(y) \\ &= \frac{\theta}{|X_2|} g(c_2) + \sum_{x \in M_2} \sum_{\substack{y=x, \\ y \in X_2}} w_{c_2 y} g(y) = h(c_2, X_2). \end{aligned} \quad (17)$$

As each $w_{c_i x}$ is positive, each distinct element in either M_1 or M_2 will contribute to the representations of c_1 and c_2 . We deduce $M_1 = M_2$ if $h(c_1, X_1) = h(c_2, X_2)$ and reach a contradiction. Thus, $h(c_1, X_1) = h(c_2, X_2)$, $|X_1| \neq |X_2|$ does not hold when $M_1 \neq M_2$.

$M_1 = M_2$ but $c_1 \neq c_2$. Next, we analyze the representations learned by the PMP-GCN layer when $M_1 = M_2$ but $c_1 \neq c_2$. Accordingly, we assume

412 $M_1 = M_2 = M$, $c_1 \neq c_2$, and $h(c_1, X_1) = h(c_2, X_2)$, $|X_1| \neq |X_2|$. Thus, we
 413 have the following:

$$\begin{aligned}
 w_{c_i x} &= \phi(\mathbf{E}_{c_i x}) = \frac{\mathbf{E}_{c_i x}}{\sqrt{\sum_{x \in X_i} \mathbf{E}_{c_i x} \cdot \sum_{z \in X_y} \mathbf{E}_{yz}}}, \\
 h(c_1, X_1) &= \frac{\theta}{|X_1|} g(c_1) + \sum_{y \in X_1} w_{c_1 y} g(y) \\
 &= \frac{\theta}{|X_2|} g(c_2) + \sum_{y \in X_2} w_{c_2 y} g(y) = h(c_2, X_2).
 \end{aligned} \tag{18}$$

414 Given $M_1 = M_2 = M$, the message weights of each $x \in M$ composed of
 415 the representations of c_1 and c_2 should be identical if $h(c_1, X_1) = h(c_2, X_2)$.
 416 Denote the underlying sets for X_1 and X_2 as $M = \{c_1, c_2, \dots\}$, and central
 417 node features as c_1 and c_2 . For $x = c_1$, we have the following:

$$\begin{aligned}
 w_{c_i x} &= \phi(\mathbf{E}_{c_i x}) = \frac{\mathbf{E}_{c_i x}}{\sqrt{\sum_{x \in X_i} \mathbf{E}_{c_i x} \cdot \sum_{z \in X_y} \mathbf{E}_{yz}}}, \\
 \frac{\theta}{|X_1|} g(c_1) + \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} g(y) &= \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y} g(y).
 \end{aligned} \tag{19}$$

418 If $\sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} = \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y}$, we have:

$$\frac{\theta}{|X_1|} = 0, \tag{20}$$

419 which contradicts $\frac{\theta}{|X_1|} > 0$. If $\sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} \neq \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y}$, we have:

$$|X_1| = \frac{\theta}{\sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y} - \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y}}. \tag{21}$$

420 As θ is independent of message weights, the RHS of the equation above can
 421 be an irrational number, which contradicts the LHS as a positive integer.
 422 Thus, $h(c_1, X_1) = h(c_2, X_2)$, $|X_1| \neq |X_2|$ does not hold when $M_1 = M_2$ but
 423 $c_1 \neq c_2$.

424 $M_1 = M_2$ and $c_1 = c_2$. At last, we analyze the representations learned
 425 by the PMP-GCN layer when $M_1 = M_2$ and $c_1 = c_2$. Similarly, we can
 426 assume that $h(c_1, X_1) = h(c_2, X_2)$, which is then shown to be FALSE due to
 427 contradictions. If $h(c_1, X_1) = h(c_2, X_2)$, the message weights of each $x \in M$
 428 that are used to generate the representations of c_1 and c_2 should be the same.
 429 For an x in M , say $x = c$, we have the following:

$$\begin{aligned}
 w_{cx} = \phi(\mathbf{E}_{cx}) &= \frac{\mathbf{E}_{cx}}{\sqrt{\sum_{x \in X_i} \mathbf{E}_{cx} \cdot \sum_{z \in X_y} \mathbf{E}_{yz}}}, \\
 \frac{\theta}{|X_1|} g(c) + \sum_{\substack{y=c, \\ y \in X_1}} w_{cy} g(y) &= \frac{\theta}{|X_2|} g(c) + \sum_{\substack{y=c, \\ y \in X_2}} w_{cy} g(y).
 \end{aligned} \tag{22}$$

430 Accordingly, Eq. (22) can be rewritten as follows:

$$\frac{|X_1|}{|X_2|} = 1 + \frac{|X_1|}{\theta} \left[\sum_{\substack{y=c, \\ y \in X_1}} w_{cy} - \sum_{\substack{y=c, \\ y \in X_2}} w_{cy} \right]. \tag{23}$$

431 The LHS of Eq. (23) is a positive rational number that is not equal to 1,
 432 which contradicts that the RHS is either 1 or an irrational number. Under
 433 all possible conditions, Eq. (23) does not hold. Combining the conducted
 434 analysis of all four cases, the assumption shown in Eq. (16) is FALSE. In
 435 other words, there does NOT exist a pair of graph structures that the PMP-
 436 GCN layer cannot distinguish. Therefore, any layer in a PMP-GCN can pass
 437 the 1-WL test. $\square$

438 4.2.3. The equivalence between the PMP-GAT layer and the 1-WL test

439 *Proof.* Similarly, we can validate that any layer in a PMP-GAT can pass
 440 the 1-WL test via *proof by contradiction*. First, assume that Eq. (14) for
 441 PMP-GAT is FALSE, i.e., there exists a pair of graphs that the PMP-GAT

layer cannot distinguish. This assumption can be written as follows:

$$\begin{aligned}
\mathbf{E}_{c_i x} &= \exp \{a_{c_i x} + b_{c_i x} - \beta(k_{c_i x} + l_{c_i x})\}, \\
w_{c_i x} &= \phi(\mathbf{E}_{c_i x}) = \frac{\mathbf{E}_{c_i x}}{\sum_{x \in X_i} \mathbf{E}_{c_i x}}, \\
h(c_1, X_1) &= h(c_2, X_2), |X_1| \neq |X_2|,
\end{aligned} \tag{24}$$

where $h(\cdot)$ is defined as Eq. (8) shows. We know that the representations for c_1 and c_2 generated by a PMP-GAT layer can be analyzed in four different cases, including $M_1 \neq M_2$ and $c_1 \neq c_2$, $M_1 \neq M_2$ but $c_1 = c_2$, $M_1 = M_2$ but $c_1 \neq c_2$, and $M_1 = M_2$ and $c_1 = c_2$. We will then demonstrate that $h(c_1, X_1) = h(c_2, X_2)$ can not hold as possible contradictions are identified under all these four cases.

$M_1 \neq M_2$ and $c_1 \neq c_2$ or $M_1 \neq M_2$ but $c_1 = c_2$. Let us first analyze the representations learned by a PMP-GAT layer when $M_1 \neq M_2$. If $h(c_1, X_1) = h(c_2, X_2)$, we have:

$$\begin{aligned}
h(c_1, X_1) &= \frac{\theta}{|X_1|} g(c_1) + \sum_{x \in M_1} \sum_{\substack{y=x, \\ y \in X_1}} w_{c_1 y} g(y) \\
&= \frac{\theta}{|X_2|} g(c_2) + \sum_{x \in M_2} \sum_{\substack{y=x, \\ y \in X_2}} w_{c_2 y} g(y) = h(c_2, X_2).
\end{aligned} \tag{25}$$

As each $w_{c_i x}$ is positive, each distinct element in either M_1 or M_2 will contribute to the representations of c_1 and c_2 . We deduce $M_1 = M_2$ if $h(c_1, X_1) = h(c_2, X_2)$, which contradicts $M_1 \neq M_2$. Thus, $h(c_1, X_1) = h(c_2, X_2)$, $|X_1| \neq |X_2|$ does not hold when $M_1 \neq M_2$.

$M_1 = M_2$ but $c_1 \neq c_2$. Next, we analyze the representations learned by a PMP-GAT layer when $M_1 = M_2$ but $c_1 \neq c_2$. Accordingly, we assume

458 $M_1 = M_2 = M$, $c_1 \neq c_2$, and $h(c_1, X_1) = h(c_2, X_2)$. According to Eq. (24),
 459 we have the following:

$$\begin{aligned}
 w_{c_i x} &= \phi(\mathbf{E}_{c_i x}) = \frac{\mathbf{E}_{c_i x}}{\sum_{x \in X_i} \mathbf{E}_{c_i x}}, \\
 h(c_1, X_1) &= \frac{\theta}{|X_1|} g(c_1) + \sum_{x \in M} \sum_{\substack{y=x, \\ y \in X_1}} w_{c_1 y} g(y) \\
 &= \frac{\theta}{|X_2|} g(c_2) + \sum_{x \in M} \sum_{\substack{y=x, \\ y \in X_2}} w_{c_2 y} g(y) = h(c_2, X_2).
 \end{aligned} \tag{26}$$

460 Given $M_1 = M_2 = M$, the message weights of each $x \in M$ composed of
 461 the representations of c_1 and c_2 should be identical if $h(c_1, X_1) = h(c_2, X_2)$.
 462 Denote the underlying sets for X_1 and X_2 as $M = \{c_1, c_2, \dots\}$, and central
 463 node features as c_1 and c_2 . For $x = c_1$, we have the following:

$$\begin{aligned}
 w_{c_i x} &= \phi(\mathbf{E}_{c_i x}) = \frac{\mathbf{E}_{c_i x}}{\sum_{x \in X_i} \mathbf{E}_{c_i x}}, \\
 \frac{\theta}{|X_1|} g(c_1) + \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} g(y) &= \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y} g(y).
 \end{aligned} \tag{27}$$

464 If $\sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} = \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y}$, we have:

$$\frac{\theta}{|X_1|} = 0, \tag{28}$$

465 which contradicts $\frac{\theta}{|X_1|} > 0$. If $\sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} \neq \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y}$, we have:

$$|X_1| = \frac{\theta}{\sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y} - \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y}}. \tag{29}$$

466 The RHS of the equation above can be an irrational number, which contra-
 467 dicts the LHS as a positive integer. Thus, $h(c_1, X_1) = h(c_2, X_2)$, $|X_1| \neq |X_2|$
 468 does not hold when $M_1 = M_2$ but $c_1 \neq c_2$.

469 $M_1 = M_2$ and $c_1 = c_2$. At last, we analyze the representations learned by the
 470 PMP-GAT layer when $M_1 = M_2$ and $c_1 = c_2$. If $h(c_1, X_1) = h(c_2, X_2)$, the
 471 message weights of each $x \in M$ that are used to generate the representations
 472 of c_1 and c_2 should be the same. Thus, for $x = c$ and $x \neq c$, we have the
 473 follows:

$$\begin{aligned} \frac{\theta}{|X_1|}g(c) + \sum_{\substack{y=c, \\ y \in X_1}} \frac{\mathbf{E}_{cy}}{\sum_{x \in X_1} \mathbf{E}_{cx}}g(y) &= \frac{\theta}{|X_2|}g(c) + \sum_{\substack{y=c, \\ y \in X_2}} \frac{\mathbf{E}_{cy}}{\sum_{x \in X_2} \mathbf{E}_{cx}}g(y), \\ \sum_{\substack{x \neq c, \\ x \in M}} \sum_{\substack{y=x, \\ y \in X_1}} \frac{\mathbf{E}_{cy}}{\sum_{x \in X_1} \mathbf{E}_{cx}}g(y) &= \sum_{\substack{x \neq c, \\ x \in M}} \sum_{\substack{y=x, \\ y \in X_2}} \frac{\mathbf{E}_{cy}}{\sum_{x \in X_2} \mathbf{E}_{cx}}g(y). \end{aligned} \quad (30)$$

474 Since the sum of attention coefficients equals 1, we have the following:

$$\begin{aligned} \sum_{\substack{y=c, \\ y \in X_1}} \frac{\mathbf{E}_{cy}}{\sum_{x \in X_1} \mathbf{E}_{cx}} &= 1 - \sum_{\substack{x \neq c, \\ x \in M}} \sum_{\substack{y=x, \\ y \in X_1}} \frac{\mathbf{E}_{cy}}{\sum_{x \in X_1} \mathbf{E}_{cx}} \\ &= 1 - \sum_{\substack{x \neq c, \\ x \in M}} \sum_{\substack{y=x, \\ y \in X_2}} \frac{\mathbf{E}_{cy}}{\sum_{x \in X_2} \mathbf{E}_{cx}} = \sum_{\substack{y=c, \\ y \in X_2}} \frac{\mathbf{E}_{cy}}{\sum_{x \in X_2} \mathbf{E}_{cx}}. \end{aligned} \quad (31)$$

475 Accordingly, for $x = c$, Eq. (30) can be rewritten as follows:

$$\frac{\theta}{|X_1|} = \frac{\theta}{|X_2|}. \quad (32)$$

476 This equation does not hold as it contradicts $|X_1| \neq |X_2|$. Combining the
 477 conducted analysis of all four cases, the assumption shown in Eq. (24) is
 478 FALSE. In other words, there does NOT exist a pair of graph structures that
 479 the PMP-GAT layer cannot distinguish. Therefore, any layer in a PMP-GAT
 480 can pass the 1-WL test. $\square$

481 4.2.4. The equivalence between the PMP-GPN layer and the 1-WL test

482 *Proof.* Like the proof for PMP-GCN and PMP-GAT, we can prove that any
 483 layer in a PMP-GAT can pass the 1-WL test via *proof by contradiction*.

484 First, assume that Eq. (14) for PMP-GPN is FALSE, i.e., there exists a pair
 485 of graphs that the PMP-GPN layer cannot distinguish. This assumption can
 486 be written as follows:

$$\begin{aligned} \mathbf{E}_{c_i x} &= \exp \{a_{c_i x} + b_{c_i x} - \beta(k_{c_i x} + l_{c_i x})\}, \\ w_{c_i x} &= \phi(\mathbf{E}_{c_i x}) = \frac{\mathbf{E}_{c_i x}}{\sqrt{\sum_{x \in X_i} \mathbf{E}_{c_i x} \cdot \sum_{z \in X_y} \mathbf{E}_{yz}}}, \\ h(c_1, X_1) &= h(c_2, X_2), |X_1| \neq |X_2|, \end{aligned} \quad (33)$$

487 where $h(\cdot)$ is defined as Eq. (9) shows. We know that the representations for
 488 c_1 and c_2 generated by a PMP-GPN layer can be analyzed in four different
 489 cases, including $M_1 \neq M_2$ and $c_1 \neq c_2$, $M_1 \neq M_2$ but $c_1 = c_2$, $M_1 = M_2$
 490 but $c_1 \neq c_2$, and $M_1 = M_2$ and $c_1 = c_2$. We will then demonstrate that
 491 $h(c_1, X_1) = h(c_2, X_2)$ can not hold as possible contradictions are identified
 492 under all these four cases.

493 $M_1 \neq M_2$ and $c_1 \neq c_2$ or $M_1 \neq M_2$ but $c_1 = c_2$. Let us first analyze the
 494 representations learned by a PMP-GPN layer when $M_1 \neq M_2$. If $h(c_1, X_1) =$
 495 $h(c_2, X_2)$, we have:

$$\begin{aligned} h(c_1, X_1) &= \frac{\theta}{|X_1|} c_1 + (1 - \alpha) \sum_{x \in M_1} \sum_{\substack{y=x, \\ y \in X_1}} w_{c_1 y} y + \alpha c_1^{in} \\ &= \frac{\theta}{|X_2|} c_2 + (1 - \alpha) \sum_{x \in M_2} \sum_{\substack{y=x, \\ y \in X_2}} w_{c_2 y} y + \alpha c_2^{in} = h(c_2, X_2), \end{aligned} \quad (34)$$

496 where c_1^{in} and c_2^{in} represent the features of the two central nodes before the
 497 message-passing of the first PMP-GPN layer. As each $w_{c_i x}$ is positive, each
 498 distinct element in either M_1 or M_2 will contribute to the representations of
 499 c_1 and c_2 . We deduce $M_1 = M_2$ if $h(c_1, X_1) = h(c_2, X_2)$, which contradicts

500 $M_1 \neq M_2$. Thus, $h(c_1, X_1) = h(c_2, X_2)$, $|X_1| \neq |X_2|$ does not hold when
 501 $M_1 \neq M_2$.

502 $M_1 = M_2$ but $c_1 \neq c_2$. Next, we analyze the representations learned by
 503 a PMP-GPN layer when $M_1 = M_2$ but $c_1 \neq c_2$. Accordingly, we assume
 504 $M_1 = M_2 = M$, $c_1 \neq c_2$, and $h(c_1, X_1) = h(c_2, X_2)$. According to Eq. (33),
 505 we have the following:

$$\begin{aligned} w_{c_i x} &= \phi(\mathbf{E}_{c_i x}) = \frac{\mathbf{E}_{c_i x}}{\sqrt{\sum_{x \in X_i} \mathbf{E}_{c_i x} \cdot \sum_{z \in X_y} \mathbf{E}_{yz}}}, \\ h(c_1, X_1) &= \frac{\theta}{|X_1|} c_1 + (1 - \alpha) \sum_{x \in M} \sum_{\substack{y=x, \\ y \in X_1}} w_{c_1 y} y + \alpha c_1^{in}, \\ h(c_2, X_2) &= \frac{\theta}{|X_2|} c_2 + (1 - \alpha) \sum_{x \in M} \sum_{\substack{y=x, \\ y \in X_2}} w_{c_2 y} y + \alpha c_2^{in}, \end{aligned} \quad (35)$$

506 where c_1^{in} and c_2^{in} represent the features of the two central nodes before the
 507 message-passing of the first PMP-GPN layer. Given $M_1 = M_2 = M$, the
 508 message weights of each $x \in M$ composed of the representations of c_1 and c_2
 509 should be identical if $h(c_1, X_1) = h(c_2, X_2)$. Denote the underlying sets for
 510 X_1 and X_2 as $M = \{c_1, c_2, \dots\}$, and central node features as c_1 and c_2 . For
 511 $x = c_1$ at the first PMP-GPN layer, we have the following:

$$\begin{aligned} w_{c_i x} &= \phi(\mathbf{E}_{c_i x}) = \frac{\mathbf{E}_{c_i x}}{\sqrt{\sum_{x \in X_i} \mathbf{E}_{c_i x} \cdot \sum_{z \in X_y} \mathbf{E}_{yz}}}, \\ \frac{\theta}{|X_1|} c_1 + (1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} y + \alpha c_1 &= (1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y} y. \end{aligned} \quad (36)$$

512 If $(1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} y + \alpha c_1 = (1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y} y$, we have:

$$\frac{\theta}{|X_1|} = 0, \quad (37)$$

513 which contradicts $\frac{\theta}{|X_1|} > 0$. If $(1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} + \alpha \neq (1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y}$,
 514 we have:

$$|X_1| = \frac{\theta}{(1 - \alpha) [\sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y} - \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y}] - \alpha}. \quad (38)$$

515 The RHS of the equation above can be an irrational number, which contra-
 516 dicts the LHS as a positive integer. For $x = c_1$ at a higher PMP-GPN layer,
 517 we have the following:

$$w_{c_i x} = \phi(\mathbf{E}_{c_i x}) = \frac{\mathbf{E}_{c_i x}}{\sqrt{\sum_{x \in X_i} \mathbf{E}_{c_i x} \cdot \sum_{z \in X_y} \mathbf{E}_{yz}}}, \quad (39)$$

$$\frac{\theta}{|X_1|} c_1 + (1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} y + \alpha c_1^{in} = (1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y} y.$$

518 The equation above can be further rewritten as:

$$\frac{\theta}{|X_1|} + (1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} y + \alpha \frac{c_1^{in}}{c_1} = (1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y}. \quad (40)$$

519 $\frac{c_1^{in}}{c_1}$ is a non-zero number as c_1 and c_1^{in} are countable features. If $(1 -$
 520 $\alpha) \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} + \alpha \frac{c_1^{in}}{c_1} = (1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y}$, we have $\frac{\theta}{|X_1|} = 0$, which con-
 521 tradicts $\frac{\theta}{|X_1|} > 0$. If $(1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y} + \alpha \frac{c_1^{in}}{c_1} \neq (1 - \alpha) \sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y}$, we
 522 have:

$$|X_1| = \frac{\theta}{(1 - \alpha) [\sum_{\substack{y=c_1, \\ y \in X_2}} w_{c_2 y} - \sum_{\substack{y=c_1, \\ y \in X_1}} w_{c_1 y}] - \alpha \frac{c_1^{in}}{c_1}}. \quad (41)$$

523 The RHS of the equation above can be an irrational number, which contra-
 524 dicts the LHS as a positive integer. Thus, $h(c_1, X_1) = h(c_2, X_2)$, $|X_1| \neq |X_2|$
 525 does not hold when $M_1 = M_2$ but $c_1 \neq c_2$.

526 $M_1 = M_2$ and $c_1 = c_2$. At last, we analyze the representations learned by the
 527 PMP-GPN layer when $M_1 = M_2$ and $c_1 = c_2$. If $h(c_1, X_1) = h(c_2, X_2)$, the

528 message weights of each $x \in M$ that are used to generate the representations
 529 of c_1 and c_2 should be the same. For an x in M , say $x = c$, we have the
 530 following:

$$w_{cx} = \phi(\mathbf{E}_{cx}) = \frac{\mathbf{E}_{cx}}{\sqrt{\sum_{x \in X_i} \mathbf{E}_{cx} \cdot \sum_{z \in X_y} \mathbf{E}_{yz}}},$$

$$\frac{\theta}{|X_1|}c + (1 - \alpha) \sum_{\substack{y=c, \\ y \in X_1}} w_{cy}y + \alpha c^{in} = \frac{\theta}{|X_2|}c + (1 - \alpha) \sum_{\substack{y=c, \\ y \in X_2}} w_{cy}y + \alpha c^{in}. \quad (42)$$

531 Accordingly, Eq. (42) can be rewritten as follows:

$$\frac{1}{|X_2|} - \frac{1}{|X_1|} = \frac{(1 - \alpha)}{\theta} \left[\sum_{\substack{y=c, \\ y \in X_1}} w_{cy} - \sum_{\substack{y=c, \\ y \in X_2}} w_{cy} \right]. \quad (43)$$

532 The LHS of Eq. (43) is a non-zero rational number, which contradicts that
 533 the RHS of Eq. (43) can be zero, or an irrational number. Combining the
 534 conducted analysis of all four cases, the assumption shown in Eq. (33) is
 535 FALSE. In other words, there does NOT exist a pair of graph structures that
 536 the PMP-GPN layer cannot distinguish. Therefore, any layer in a PMP-GPN
 537 can pass the 1-WL test. $\square$

538 Combining the theoretical analysis conducted in Sections 4.2.2-4.2.4, The-
 539 orem 2 is proved. Previous studies have demonstrated that most conventional
 540 MPGNNs cannot pass the 1-WL test [17, 18, 21]. Therefore, the theoretical
 541 results shown in Theorems 1 and 2 establish the constantly higher expres-
 542 siveness of the proposed PMP-GNNs compared with conventional MPGNNs,
 543 additionally providing a theoretical guarantee of the robust performance of
 544 PMP-GNNs in diverse learning tasks.

545 5. Related work

546 In this section, previous works that are closely related to the proposed
547 PMP-GNNs are briefly reviewed.

548 5.1. Graph neural networks

549 GNNs [8, 18, 22] are recognized as powerful tools for analyzing graph-
550 structured data. They have been successfully adapted for diverse real-world
551 applications, including scientific document classification [23], coauthor anal-
552 ysis [5], conversation generation [24], social community detection [21], and
553 text classification [25]. Typically, the low-dimensional representations tai-
554 lored for downstream tasks are learned using a GNN through multi-layer
555 node feature aggregation and projection. Consequently, the aggregation of
556 neighbor features for each central node is crucial for GNNs to learn mean-
557 ingful representations.

558 5.2. Message-passing graph neural networks

559 Most present GNNs can be conceptualized as MPGNNs [26, 27]. In each
560 layer of an MPGNN, the output representation is learned through a two-
561 step approach. Messages from the neighbors of the central node are first
562 composed through the multiplication of neighbor features with weights that
563 indicate their importance. Then the messages are conveyed to each central
564 node. The output representations are generated by updating the features of
565 the central node and the received messages, through aggregation or max or
566 min operations.

567 The strategy of composing messages, that is, the calculation of weights
568 indicating the importance of neighbor’s features, plays a central role in

MPGNNs. Various message-passing paradigms for GNNs emerge from different selections of weight calculation strategies. Several approaches for computing message weights have been proposed. For instance, some GNNs leverage an attention mechanism [3, 21, 28] that quantifies the feature correlation or similarity between pairwise nodes to formulate messages. Correlations induced by the graph structure, such as the graph Laplacian [1, 17, 29] and graph diffusion [13, 30], can also be considered measures of importance, which are subsequently used to convey messages to central nodes in the graph. To improve the efficiency of MPGNNs, sampling strategies [8, 22, 31] can be implemented to select a subset of similar neighbors, enabling the transmission of messages from the neighbors to the central node for representation learning.

Existing MPGNNs concentrate on learning representations from messages derived from neighbor similarity or correlation. Messages from the opposite side, that is, the side indicating the dissimilarity between pairwise nodes in the graph, are unexplored. Dissimilarity among data has been investigated in other domains, such as multi-view and manifold learning [32, 33]. However, these approaches cannot be readily adapted to existing MPGNNs. Therefore, we propose PMP to enable GNNs to simultaneously leverage messages related to both similarity and dissimilarity to learn more expressive representations suitable for various downstream tasks.

6. Experimental setup

In this section, we introduce the setup of our experiments to validate the proposed PMP-GCN, PMP-GAT, and PMP-GPN to reveal the effectiveness

Table 1: Statistics of the testing datasets.

	Cora	Cite	PubMed	Blog	Flickr	CoauthorCS	CoauthorPH	Sport	BBC	IrishTimes	Guardian	Wikipedia
N	2708	3327	19717	5196	7575	18333	34493	737	2225	3246	6521	5739
$ E $	5429	4732	44338	171743	239738	327576	991848	17242	9127	490325	96279	524692
D	1433	3703	500	8189	12047	6805	8415	966	3121	4823	10801	17311
C	7	6	3	6	9	15	5	5	5	7	6	6
Type	Scientific article			Blog	Text-based image	Collaboration		Rich text				

of the PMP paradigm.

6.1. Comparison baselines

To validate the effectiveness of the proposed PMP-GCN, PMP-GAT, and PMP-GPN, we selected twelve state-of-the-art GNNs as comparison baselines: GCN [1], MoNet [34], GraphSAGE [8], ARMA [29], GAT [3], GATv2 [35], HardGAT [31], MAGNA [4], JKNet [30], APPNP [5], SGC [11], and GPR [14]. Specifically, GCN, MoNet, GraphSAGE, and ARMA are four strong GNNs that are built upon various graph convolutional layers. GAT, GATv2, HardGAT, and MAGNA are four representative attention-based GNNs. JKNet, APPNP, SGC, and GPR are four strong GNNs whose layers are constructed according to PageRank (graph diffusion). Comparing PMP-GCN, PMP-GAT, and PMP-GPN with baselines whose message for representation learning is solely based on neighbor correlation or similarity can demonstrate the effectiveness of both PMP-GNNs.

6.2. Datasets and real-world analytical tasks

Twelve real-world datasets are selected to experimentally test the effectiveness of all considered approaches. The datasets are Cora, Cite, PubMed [36], Blog, Flickr [37], CoauthorCS, CoauthorPH [38], Sport, BBC, IrishTimes, Guardian, and Wikipedia [39]. The details of the test datasets are presented below.

613 *Cora, Cite, PubMed.* Cora, Cite, and PubMed are three classical graph
614 datasets representing scientific article citations. The nodes, edges, and node
615 features represent the scientific articles, article-article citations, and article
616 keywords, respectively. Cora consists of 2,708 nodes, 5,429 edges, and 1,433
617 features. The Cite dataset consists of 3,327 nodes, 4,732 edges, and 3,703 fea-
618 tures. PubMed consists of 19,717 nodes, 44,338 edges, and 500 features. All
619 articles (nodes) in these three datasets are assigned class labels representing
620 the research domains to which they belong. For Cora, Cite, and PubMed,
621 there are seven, six, and three ground-truth classes, respectively.

622 *Blog.* Blog is a social graph dataset extracted from blogcatalog.com. It con-
623 sists of 5,196 nodes, 171,743 edges, and 8,189 features, representing blog
624 sites, blog-blog hyperlinks, and contextual descriptions, respectively. Each
625 blog site in this dataset is assigned one of the six class labels representing
626 blog categories.

627 *Flickr.* Flickr is a graph dataset collected from flickr.com. It consists of
628 7,575 nodes, 239,738 edges, and 12,047 features. The nodes, edges, and
629 features in this dataset represent images, image-image links, and the content
630 descriptions of images. Each of the 7,575 images is assigned to one of nine
631 ground-truth classes.

632 *CoauthorCS and CoauthorPH.* CoauthorCS and CoauthorPH are collabo-
633 ration graphs extracted from the computer science and physics societies of
634 arxiv.com. In both datasets, each node represents a researcher, an edge be-
635 tween nodes indicates that the researchers have coauthored a paper published
636 on Arxiv, and features represent the research topics on which the researchers

637 worked. CoauthorCS consists of 18,333 nodes divided into 15 research do-
638 mains of computer science, 327,576 edges, and 6,805 features. CoauthorPH
639 consists of 34,493 nodes assigned to five research domains of physics, 991,848
640 edges, and 8,415 features.

641 *Sport, BBC, IrishTimes, Guardian, and Wikipedia.* Unlike the abovementioned graph datasets, the Sport, BBC, IrishTimes, Guardian, and Wikipedia
642 datasets are collected from rich-text media, namely BBC Sports news, BBC
643 news, news from The Times (Ireland), news from The Guardian, and Wikipedia
644 content. In these five datasets, nodes, features, and edges represent pieces
645 of news or wiki pages, textual contents, and news-news or page-page simi-
646 larities, respectively. The Sport dataset consists of 737 nodes belonging to
647 five news classes, 17,242 edges, and 966 features. The BBC dataset consists
648 of 2,225 nodes belonging to five disjoint news classes, 9,127 edges, and 3,121
649 features. IrishTimes consists of 3,246 nodes belonging to seven news classes,
650 490,325 edges, and 4,823 features. Guardian consists of 6,521 nodes belong-
651 ing to six news classes, 96,279 edges, and 10,801 features. The Wikipedia
652 dataset consists of 5,739 nodes belonging to six classes, 524,692 edges, and
653 17,311 features.

655 As these test datasets are collected from five distinct real-world sce-
656 narios, five diverse analytical tasks, namely scientific article classification
657 (Cora, Cite, and PubMed), blog clustering (Blog), text-based image cluster-
658 ing (Flickr), coauthor analysis (CoauthorCS and CoauthorPH), and rich-text
659 classification (Sport, BBC, IrishTimes, Guardian, and Wikipedia), are con-
660 sidered in the experiment to test the effectiveness of all approaches. Through
661 these diverse downstream tasks, the effectiveness of different GNNs can be

comprehensively validated. The characteristics of the 12 abovementioned datasets are summarized in Table 1. Besides, the accumulative distribution of normalized feature and structure distances among neighbors on all test datasets is plotted in Fig. A.3. It is observed that most neighbors on all test datasets are very distant, which is consistent with previous studies showing real-world graphs are created differently [6, 7]. Conducting representation learning tasks in such real-world graphs allows the proposed PMP-GNNs to perform robustly in diverse downstream tasks.

6.3. Experimental settings

Configurations of comparison baselines. To conduct fair comparisons, all GNNs generally employ a two-layer network structure to perform all downstream tasks mentioned in Section 6.2. In other words, each GNN in the experiment contains one hidden layer, followed by the output layer. The official source codes of all baseline GNNs are used in the experiment and configured using the recommended settings. Specifically, the hidden layer dimension and learning rate of GCN, MoNet, GraphSAGE, and ARMA are set to 32 and 0.01, respectively. The dropout rate of GCN, MoNet, and GraphSAGE is set to 0.5, while that of ARMA is 0.75. For attention-based GNN baselines, including GAT, GATv2, and MAGNA, the learning rate is set to 0.005 in all analytical tasks, while that of HardGAT is set to 0.01. The dropout rates of GAT, GATv2, and HardGAT are set to 0.6, while that of MAGNA is set to 0.5 in all analytical tasks. Eight attention heads are adopted by the four attention-based GNN baselines to construct the hidden layer, the dimension of each head adopted by GAT, GATv2, and HardGAT is set to 8 in the scientific article classification task, and 32 in the remaining

687 tasks. The hidden dimension of each head adopted by MAGNA is set to 32
688 and 512 in the scientific article classification task and the remaining tasks.
689 One attention head is adopted by the four attention-based GNNs to build the
690 output layer. For JKNet, APPNP, and GPR, the dropout rate is set to 0.5,
691 while that of SGC is set to 0 as recommended. The hidden layer dimension
692 is set to 32 for JKNet, 64 for APPNP and GPR, and the class number of
693 each dataset for SGC. The learning rate of JKNet and SGC is set to 0.005,
694 and that of APPNP and GPR is set to 0.01.

695 *Configurations of PMP-GCN, PMP-GAT and PMP-GPN.* The settings of
696 PMP-GCN, PMP-GAT, and PMP-GPN are similar to those of GNNs based
697 on convolutional, attention, and PageRank layers. Specifically, the hidden
698 layer dimension and dropout rate are 32 and 0.75 for PMP-GCN and 0.6 and
699 64 for PMP-GAT, respectively. The hidden layer dimension and dropout rate
700 are 64 and 0.5 for PMP-GPN, respectively. The learning rates are set to 0.01
701 for PMP-GCN and PMP-GPN and 0.005 for PMP-GAT in all downstream
702 tasks. For PMP-GAT, we only use one attention head to construct the
703 hidden layer. All GNNs are executed on a workstation installed with an
704 NVIDIA RTX3090 graphics card, Python 3.8.5, and CUDA 11.1. All GNNs
705 are trained using the Adam optimizer [40] and run 10 times in each learning
706 task to obtain the average performance with the standard deviation. To
707 evaluate the performance of all GNNs, we use *Accuracy* (*Acc*), a widely
708 accepted metric in various real-world learning tasks.

709 *Data split in all learning tasks.* The data split for training, validation, and
710 testing significantly differs among the five considered learning tasks. Specif-
711 ically, for scientific article classification, established data splits [1, 3, 41] are

Table 2: Comparison of average performance in scientific article classification between convolutional GNNs and PMP-GCN. The best performance on each dataset is highlighted in bold.

	GCN	MoNet	GraphSAGE	ARMA	PMP-GCN
Cora	0.814 ± 0.019	0.820 ± 0.005	0.819 ± 0.001	0.791 ± 0.009	0.841 ± 0.002
Cite	0.716 ± 0.007	0.642 ± 0.002	0.704 ± 0.012	0.700 ± 0.016	0.727 ± 0.003
PubMed	0.797 ± 0.004	0.798 ± 0.003	0.788 ± 0.009	0.801 ± 0.008	0.823 ± 0.002

712 adopted. For each test dataset, 20 nodes with labels from each ground-truth
 713 class are used for model training. Thus, Cora, Cite, and PubMed consist
 714 of 140, 120, and 60 nodes for model training. For these three datasets, 500
 715 and 1,000 nodes are used for validation and testing, respectively. For the
 716 blog clustering, text-based image clustering, and coauthor analysis tasks,
 717 five sets of training and validation splits are randomly generated, as there
 718 are no established data splits. For each dataset, the size of the training,
 719 validation, and test splits is equal to the number of ground-truth classes
 720 multiplied by 20, 500, and all nodes, respectively. Consequently, 120, 180,
 721 300, and 100 nodes are sampled for training on Blog, Flickr, CoauthorCS,
 722 and CoauthorPH, respectively. For the rich-text classification task on each
 723 test dataset, 60%, 20%, and 20% of nodes are sampled as training, validation,
 724 and testing splits, respectively, and five sets of splits are generated to test
 725 the performance of all GNNs.

7. Results and analysis

727 In this section, the predictive performances of the proposed PMP-GCN,
 728 PMP-GAT, and PMP-GPN are compared with those of the well-established
 729 GNN baselines mentioned in Section 6 on five learning tasks based on 12 real-

Table 3: Comparison of average performance in scientific article classification between attention-based GNNs and PMP-GAT. The best performance on each dataset is highlighted in bold.

	GAT	GATv2	MAGNA	HardGAT	PMP-GAT
Cora	0.835 ± 0.003	0.841 ± 0.008	0.829 ± 0.002	0.794 ± 0.005	0.850 ± 0.003
Cite	0.708 ± 0.003	0.714 ± 0.005	0.672 ± 0.008	0.719 ± 0.005	0.721 ± 0.004
PubMed	0.815 ± 0.005	0.819 ± 0.003	0.813 ± 0.007	0.810 ± 0.003	0.833 ± 0.004

world datasets. In addition, the significant difference between the proposed PMP-GNNs and other MPGNNs is discussed according to the corresponding results.

7.1. Classification performance on scientific articles

Scientific article classification is a widely accepted testbed for evaluating GNN performance. Three classic datasets, Cora, Cite, and PubMed, are used in our experiment to reveal the effectiveness of different GNNs. The comparisons regarding classification performances are summarized in Tables 2, 3, and 4. The proposed PMP-GCN outperforms all GNN baselines built with various convolutional layers (Table 2). Specifically, PMP-GCN achieves a performance improvement of 2.56% on Cora, 1.53% on Cite, and 2.75% on PubMed. The proposed PMP-GAT also outperforms all baselines of attention-based GNNs. PMP-GAT achieves a performance gain of 1.07%, 0.28%, and 1.71% on Cora, Cite, and PubMed, respectively (Table 3). Compared with PageRank GNNs, the GNN based on the proposed PMP paradigm still performs robustly. The proposed PMP-GPN achieves a performance gain of 2.55%, 0.42%, and 1.09% on Cora, Cite, and PubMed, respectively (Table 4). The experimental results presented in Tables 2, 3, and 4 demonstrate

Table 4: Comparison of average performance in scientific article classification between PageRank GNNs and PMP-GPN. The best performance on each dataset is highlighted in bold.

	SGC	GPR	JKNet	APPNP	PMP-GPN
Cora	0.814 ± 0.001	0.818 ± 0.005	0.794 ± 0.009	0.825 ± 0.008	0.846 ± 0.003
Cite	0.713 ± 0.000	0.720 ± 0.006	0.651 ± 0.008	0.720 ± 0.005	0.723 ± 0.002
PubMed	0.773 ± 0.000	0.827 ± 0.004	0.807 ± 0.006	0.821 ± 0.005	0.836 ± 0.002

that the proposed PMP paradigm is very effective for scientific article classification. GNNs based on the proposed message-passing paradigm can learn more expressive representations from strongly correlated articles.

7.2. Clustering performance on blogs and text-based images

Blog clustering and text-based image clustering are two challenging tasks, as all nodes in the datasets are classified with very limited supervision. In our experiment, Blog and Flickr are used to test the effectiveness of all GNNs in these two challenging tasks. The corresponding results are summarized in Tables 5, 6 and 7. The proposed PMP-based GNNs still perform robustly compared with the various types of GNN baselines. PMP-GCN outperforms MoNet in clustering *Acc* on Blog and Flickr datasets by 11.08% and 4.20%, respectively. PMP-GAT outperforms attention-based GNNs by 32.84% on the Blog dataset and 11.32% on the Flickr dataset. PMP-GPN outperforms PageRank GNNs by 13.73% and 10.99% on Blog and Flickr datasets, respectively. Compared with the datasets of scientific articles, Blog and Flickr are denser, potentially resulting in a larger number of highly dissimilar nodes that are connected. The proposed PMP paradigm can significantly improve the learning performance of GNNs on these dense datasets by considerably

Table 5: Comparison of average performance in blog and text-based image clustering, and coauthor analysis between convolutional GNNs and PMP-GCN. The best performance in each dataset is highlighted in bold.

	GCN	MoNet	GraphSAGE	ARMA	PMP-GCN
Blog	0.654 ± 0.005	0.713 ± 0.012	0.572 ± 0.007	0.625 ± 0.007	0.792 ± 0.007
Flickr	0.471 ± 0.009	0.524 ± 0.007	0.370 ± 0.005	0.408 ± 0.003	0.546 ± 0.009
CoauthorCS	0.887 ± 0.007	0.875 ± 0.009	0.899 ± 0.009	0.824 ± 0.003	0.905 ± 0.005
CoauthorPH	0.923 ± 0.003	0.918 ± 0.003	0.936 ± 0.002	0.917 ± 0.007	0.927 ± 0.007

Table 6: Comparison of average performance in blog and text-based image clustering, and coauthor analysis between attention-based GNNs and PMP-GAT. The best performance in each dataset is highlighted in bold.

	GAT	GATv2	MAGNA	HardGAT	PMP-GAT
Blog	0.489 ± 0.004	0.606 ± 0.007	0.570 ± 0.006	0.447 ± 0.006	0.805 ± 0.007
Flickr	0.468 ± 0.005	0.357 ± 0.004	0.365 ± 0.007	0.420 ± 0.003	0.521 ± 0.005
CoauthorCS	0.880 ± 0.004	0.884 ± 0.004	0.857 ± 0.006	0.890 ± 0.005	0.905 ± 0.002
CoauthorPH	0.922 ± 0.003	0.921 ± 0.004	0.908 ± 0.004	0.927 ± 0.002	0.939 ± 0.002

766 reducing the importance of messages from dissimilar neighbors during the
767 generation of output representations for each node. This explains the signifi-
768 cant performance gains of PMP-GCN, PMP-GAT, and PMP-GPN compared
769 with other MPGNNs.

770 7.3. Analytical performance in coauthor graphs

771 Coauthor analysis is one of the major tasks performed by modern GNNs.
772 It usually involves leveraging limited supervision to correctly categorize au-
773 thors who frequently collaborate or work within correlated research domains.
774 In our experiment, we use two large collaborative graphs, CoauthorCS and
775 CoauthorPH, to test the performance of different GNNs. The comparative
776 results are summarized in Tables 5, 6, and 7. The proposed PMP-GCN,

Table 7: Comparison of average performance in blog and text-based image clustering, and coauthor analysis between PageRank GNNs and PMP-GPN. The best performance in each dataset is highlighted in bold.

	SGC	GPR	JKNet	APNP	PMP-GPN
Blog	0.501 ± 0.007	0.675 ± 0.004	0.714 ± 0.007	0.686 ± 0.009	0.812 ± 0.006
Flickr	0.352 ± 0.001	0.441 ± 0.006	0.430 ± 0.003	0.473 ± 0.008	0.525 ± 0.007
CoauthorCS	0.885 ± 0.003	0.854 ± 0.001	0.879 ± 0.005	0.887 ± 0.010	0.903 ± 0.002
CoauthorPH	0.919 ± 0.001	0.927 ± 0.001	0.932 ± 0.002	0.922 ± 0.006	0.934 ± 0.001

PMP-GAT, and PMP-GPN demonstrate a robust analytical performance
 compared with the other state-of-the-art GNNs. PMP-GCN outperforms
 the best convolutional GNN baseline (GraphSAGE) by 0.67% on the Coau-
 thorCS dataset and performs second best on the CoauthorPH dataset (Table
 5). After investigating the data characteristics and comparing the message-
 passing paradigms of PMP-GCN and GraphSAGE, we may deduce the rea-
 sons leading PMP-GCN not to perform the best on the CoauthorPH dataset.
 It is observed from Fig. A.3 that the portion of very large distances (i.e.,
 very dissimilar neighbors pertaining to structure or feature) on the Coau-
 thorPH dataset is relatively lower compared with other test datasets. There-
 fore, a GNN is expected to perform better if it leverages more information
 from neighbors to learn representations. As the proposed PMP-GCN pur-
 sues learning representations from sparse and strongly correlated neighbors,
 under the experimental settings for all test datasets, PMP-GCN still leans
 toward learning zero or very-close-to-zero message weights for a large num-
 ber of neighbors on the CoauthorPH dataset (See Figs. A.5). As a result,
 PMP-GCN might not involve a small portion of neighbors that might be
 meaningful to generate representations. In contrast, the message weights

Table 8: Comparison of average performance in rich-text classification between convolutional GNNs and PMP-GCN. The best performance in each dataset is highlighted in bold.

	GCN	MoNet	GraphSAGE	ARMA	PMP-GCN
Sport	0.966 ± 0.006	0.961 ± 0.006	0.955 ± 0.005	0.972 ± 0.007	0.991 ± 0.003
BBC	0.956 ± 0.004	0.964 ± 0.003	0.973 ± 0.003	0.938 ± 0.010	0.981 ± 0.002
IrishTimes	0.898 ± 0.004	0.684 ± 0.005	0.927 ± 0.004	0.868 ± 0.005	0.937 ± 0.002
Guardian	0.951 ± 0.002	0.795 ± 0.005	0.936 ± 0.004	0.934 ± 0.004	0.959 ± 0.001
Wikipedia	0.883 ± 0.003	0.659 ± 0.008	0.880 ± 0.008	0.882 ± 0.002	0.904 ± 0.005

adopted by GraphSAGE are defined as $\frac{1}{d_i}$, where d_i is the degree of central node i . GraphSAGE performs slightly better than the proposed PMP-GCN, possibly because GraphSAGE leverages more neighbors to learn representations on the CoauthorPH dataset.

In contrast to PMP-GCN, the other two PMP-GNNs (PMP-GAT and PMP-GPN) exhibit the best performance on both datasets. As shown in Table 6, PMP-GAT outperforms HardGAT by 1.69% and 1.29% on CoauthorCS and CoauthorPH. PMP-GPN outperforms APPNP by 1.80% on CoauthorCS, and is better than JKNet by 0.21% on CoauthorPH (Table 7). The results demonstrate that leveraging polarized messages to learn representations for coauthor analysis tasks allows GNNs to reduce the influence of dissimilar coauthors. The subsequent representations are learned by concentrating on the strongly correlated coauthors.

7.4. Performance in rich-text classification

Rich-text classification is a fundamental task in natural language processing. It involves the accurate categorization of rich-content documents, such as news and wiki pages. In our experiment, we select five representative

Table 9: Comparison of average performance in rich-text classification between attention-based GNNs and PMP-GAT. The best performance in each dataset is highlighted in bold.

	GAT	GATv2	MAGNA	HardGAT	PMP-GAT
Sport	0.965 ± 0.008	0.964 ± 0.007	0.954 ± 0.006	0.923 ± 0.001	0.984 ± 0.003
BBC	0.963 ± 0.003	0.962 ± 0.003	0.959 ± 0.006	0.965 ± 0.008	0.977 ± 0.001
IrishTimes	0.880 ± 0.004	0.898 ± 0.004	0.885 ± 0.005	0.780 ± 0.003	0.925 ± 0.004
Guardian	0.940 ± 0.003	0.929 ± 0.003	0.925 ± 0.002	0.923 ± 0.007	0.949 ± 0.002
Wikipedia	0.876 ± 0.002	0.874 ± 0.004	0.869 ± 0.008	0.844 ± 0.002	0.889 ± 0.002

Table 10: Comparison of average performance in rich-text classification between PageRank GNNs and PMP-GPN. The best performance in each dataset is highlighted in bold.

	SGC	GPR	JKNet	APPNP	PMP-GPN
Sport	0.973 ± 0.000	0.936 ± 0.003	0.955 ± 0.009	0.969 ± 0.009	0.993 ± 0.001
BBC	0.962 ± 0.000	0.967 ± 0.003	0.944 ± 0.008	0.963 ± 0.004	0.974 ± 0.002
IrishTimes	0.748 ± 0.000	0.729 ± 0.002	0.825 ± 0.007	0.872 ± 0.006	0.932 ± 0.005
Guardian	0.920 ± 0.000	0.926 ± 0.006	0.881 ± 0.003	0.930 ± 0.009	0.948 ± 0.001
Wikipedia	0.814 ± 0.000	0.869 ± 0.005	0.729 ± 0.006	0.858 ± 0.006	0.905 ± 0.004

812 text datasets, namely Sport, BBC, IrishTimes, Guardian, and Wikipedia,
813 to verify the effectiveness of different GNNs. The corresponding results
814 are summarized in Tables 8, 9, and 10. The proposed PMP-GCN, PMP-
815 GAT, and PMP-GPN outperform various types of MPGNNs in the task of
816 rich-text classification. PMP-GCN outperforms convolution-based GNNs on
817 all datasets. For example, it outperforms ARMA by 1.95% on the Sport
818 dataset. PMP-GCN outperforms GraphSAGE by 0.82% on BBC, and 1.08%
819 on IrishTimes. PMP-GCN outperforms GCN on Guardian and Wikipedia
820 datasets by 0.84% and 2.38%, respectively. The proposed PMP-GAT ex-
821 hibits a robust performance compared with attention-based GNNs. PMP-
822 GAT outperforms GAT by 1.97%, 0.96%, and 1.48% on Sport, Guardian,

Table 11: Performance comparisons between PMP-based GNNs and their variants.

Graph Convolution												
	Cora	Cite	PubMed	Blog	Flickr	CoauthorCS	CoauthorPH	Sport	BBC	IrishTimes	Guardian	Wikipedia
GCN	0.814	0.716	0.797	0.654	0.471	0.887	0.923	0.966	0.956	0.898	0.951	0.883
PMP-GCN/d	0.838	0.718	0.809	0.767	0.530	0.899	0.922	0.963	0.966	0.902	0.952	0.903
PMP-GCN	0.841	0.727	0.823	0.792	0.546	0.905	0.927	0.991	0.981	0.937	0.959	0.904
Graph Attention												
	Cora	Cite	PubMed	Blog	Flickr	CoauthorCS	CoauthorPH	Sport	BBC	IrishTimes	Guardian	Wikipedia
GAT	0.835	0.708	0.815	0.489	0.468	0.880	0.922	0.965	0.963	0.880	0.940	0.876
PMP-GAT/d	0.836	0.709	0.818	0.802	0.497	0.893	0.915	0.979	0.962	0.921	0.942	0.885
PMP-GAT	0.850	0.721	0.833	0.805	0.521	0.905	0.939	0.984	0.977	0.925	0.949	0.889
Graph PageRank												
	Cora	Cite	PubMed	Blog	Flickr	CoauthorCS	CoauthorPH	Sport	BBC	IrishTimes	Guardian	Wikipedia
APNP	0.825	0.720	0.821	0.686	0.473	0.887	0.922	0.969	0.963	0.872	0.930	0.858
PMP-GPN/d	0.835	0.715	0.825	0.674	0.474	0.903	0.925	0.987	0.964	0.922	0.931	0.899
PMP-GPN	0.846	0.723	0.836	0.812	0.525	0.908	0.934	0.993	0.974	0.932	0.948	0.905

and Wikipedia, respectively. PMP-GAT outperforms HardGAT by 1.24% on BBC, and GATv2 by 3.01% on IrishTimes. PMP-GPN also performs the best when compared with PageRank GNNs. PMP-GPN outperforms SGC by 2.06% on Sport. PMP-GPN outperforms APPNP by 6.88% and 1.94% on IrishTimes and Guardian, respectively. PMP-GPN outperforms GPR by 0.72% and 4.14% on BBC and Wikipedia, respectively. In contrast to the existing GNN learning representations that solely rely on messages related to correlation, PMP-GCN, PMP-GAT, and PMP-GPN can identify dissimilar documents. PMP-GNNs can thus utilize messages from truly correlated neighbors to learn more expressive representations.

7.5. Comparisons of similarity-, dissimilarity-, and PMP-based GNNs

As the proposed PMP-GNNs leverage similarity and dissimilarity between neighbors to learn representations, comparing them with GNN variants considering either similarity or dissimilarity can further demonstrate the effectiveness of the proposed methods. Specifically, we let the learnable matrix

838 $\mathbf{C} = \mathbf{0}$, leading the proposed PMP-based GNNs to learn representations
 839 only considering neighbor dissimilarity (PMP-GCN/d, PMP-GAT/d, and
 840 PMP-GPN/d). Taking conventional GCN, GAT, and APPNP as GNN vari-
 841 ants only considering neighbor similarity for representation learning, PMP-
 842 GCN/d, PMP-GAT/d, and PMP-GPN/d as GNN variants only concerning
 843 neighbor dissimilarity, we compare their predictive performances with the
 844 proposed PMP-based GNNs on all learning tasks. The corresponding results
 845 have been summarized in Table 11. As the table shows, PMP-GCN/d, PMP-
 846 GAT/d, and PMP-GPN/d still can obtain competitive performance on most
 847 testing datasets. Considering that a very large portion of neighbors are shown
 848 to be very distant in Fig. A.3, we can conclude that solely leveraging dissim-
 849 ilarity widely existing in real-world graph data to learn message weights is
 850 an auspicious way to guide message-passing GNNs to learn expressive rep-
 851 resentations. Nevertheless, the best performances obtained by PMP-based
 852 GNNs demonstrate that considering both similarity and dissimilarity can
 853 lead message-passing GNNs to learn more expressive representations com-
 854 pared with those considering similarity or dissimilarity only.

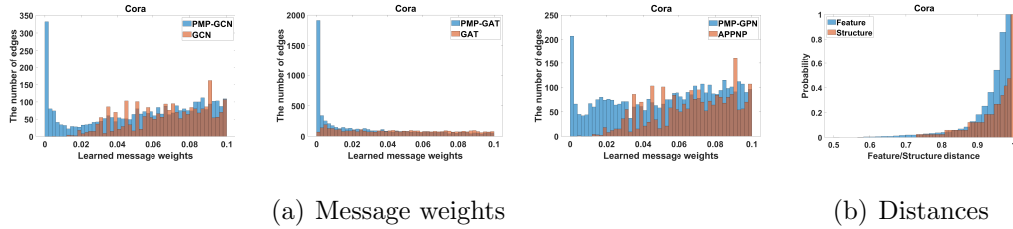


Figure 2: Message weights learned by PMP-GNNs and conventional message-passing GNNs, and feature/structure distances on Cora dataset.

855 7.6. Comparisons of learned message weights

856 As previously mentioned, the proposed PMP paradigm enables GNNs to
857 learn to reduce and negate messages transmitted by dissimilar neighbors so
858 as to preserve messages from strongly correlated neighbors. To demonstrate
859 this advantageous capability, we conduct a comparative analysis of the small
860 message weights learned by PMP-GCN, PMP-GAT, and PMP-GPN, as well
861 as the conventional GCN, GAT, and APPNP. The corresponding results on
862 the Cora dataset are exemplified in Fig. 2(a), and the results on other datasets
863 are depicted in Figs. A.4-A.6 in Appendix. Besides, the accumulative distri-
864 bution of normalized neighbor distances regarding node features and graph
865 structure is depicted in Fig. 2(b) to show the strong correlations between
866 message weights learned by PMP-GNNs and the practical data context.

867 The notable contrast regarding the message weights of PMP-GNNs and
868 conventional message-passing GNNs is shown in Fig. 2 and Figs. A.4-A.6
869 in Appendix. Taking the result on the Cora dataset as an example, PMP-
870 GCN and PMP-GPN acquire $\sim 1,000$ and ~ 700 messages whose weights are
871 extremely low (≤ 0.02), whereas conventional GCN and APPNP identify ~ 50
872 low-weight messages. These comparative results strongly indicate that PMP-
873 GCN and PMP-GPN are much better than conventional GCN and APPNP
874 in learning to reduce or negate messages conveyed by dissimilar neighbors
875 while preserving messages from a significantly reduced number of correlated
876 neighbors. Similarly, PMP-GAT can learn over 3,000 messages with very low
877 weights (≤ 0.02), whereas GAT identifies only $\sim 1,000$ low-weight messages.
878 As message weights (attention scores) for each central node are summed to
879 one, PMP-GAT can concentrate more on truly correlated neighbors in the

880 subsequent feature aggregation stage.

881 Considering a large portion of neighbors are very distant regarding node
882 features or structure (Fig. 2(b) and Fig. A.3), the results depicted in Fig. 2
883 and Figs. A.4-A.6 demonstrate that the proposed PMP paradigm enables
884 GNNs to better capture the sparseness of real-world graph data. Therefore,
885 the proposed PMP paradigm empowers GNNs to reduce or negate more mes-
886 sages from dissimilar neighbors, enabling them to learn more expressive rep-
887 resentations by appropriately preserving messages conveyed by much fewer
888 but highly similar neighbors.

889 8. Conclusion

890 In this paper, we have proposed Polarized message-passing (PMP), a
891 novel paradigm that can lead to novel designs of GNNs. Different from exist-
892 ing strategies, PMP leverages similarity and dissimilarity to acquire neigh-
893 boring messages, which enables GNNs to learn expressive representations
894 with a significantly reduced number of highly correlated neighbors. Three
895 novel GNNs based on the proposed PMP, namely PMP-GCN, PMP-GAT,
896 and PMP-GPN, are constructed. Theoretical analysis is further conducted to
897 verify the strong expressiveness of PMP-GCN, PMP-GAT, and PMP-GPN.
898 An empirical study involving five diverse downstream tasks from 12 datasets
899 is also conducted to reveal the superior performances of PMP-GCN, PMP-
900 GAT, and PMP-GPN. The results demonstrate that PMP-GCN, PMP-GAT,
901 and PMP-GPN outperform established MPGNNs in all downstream tasks,
902 indicating the effectiveness of the proposed PMP paradigm. In the future,
903 we will further improve the proposed PMP by designing and incorporating

904 novel measures for generating message weights that can reveal the intrinsic
905 properties of graphs. Moreover, we will devote efforts toward adapting PMP
906 to solve more challenging tasks, such as multi-view graph learning.

907 **Acknowledgments**

908 This work is supported in part by A*STAR I&E GAP funding (NO.
909 I23D1AG080), and the Center for Frontier AI Research (CFAR), Agency for
910 Science, Technology and Research (A*STAR).

911 **References**

- 912 [1] T. N. Kipf, M. Welling, Semi-supervised classification with graph con-
913 volutional networks, in: International Conference on Learning Repre-
914 sentations, 2017.
- 915 [2] J. Gilmer, S. S. Schoenholz, P. F. Riley, O. Vinyals, G. E. Dahl, Mes-
916 sage passing neural networks, Machine learning meets quantum physics
917 (2020) 199–214.
- 918 [3] P. Veličković, G. Cucurull, A. Casanova, A. Romero, P. Liò, Y. Bengio,
919 Graph attention networks, in: International Conference on Learning
920 Representations, 2018.
- 921 [4] G. Wang, R. Ying, J. Huang, J. Leskovec, Multi-hop attention graph
922 neural networks (2021) 3089–3096.
- 923 [5] J. Gasteiger, A. Bojchevski, S. Günnemann, Predict then propagate:
924 Graph neural networks meet personalized pagerank, arXiv preprint
925 arXiv:1810.05997 (2018).

- 926 [6] S. L. Feld, Why your friends have more friends than you do, *American*
927 *journal of sociology* 96 (1991) 1464–1477.
- 928 [7] N. Alipourfard, B. Nettasinghe, A. Abeliuk, V. Krishnamurthy, K. Lerman,
929 Friendship paradox biases perceptions in directed networks, *Nature*
930 *communications* 11 (2020) 707.
- 931 [8] W. Hamilton, Z. Ying, J. Leskovec, Inductive representation learning
932 on large graphs, in: *Advances in neural information processing systems*,
933 2017, pp. 1024–1034.
- 934 [9] J. You, R. Ying, J. Leskovec, Position-aware graph neural networks, in:
935 *International conference on machine learning*, PMLR, 2019, pp. 7134–
936 7143.
- 937 [10] F. Li, Y. Qian, J. Wang, C. Dang, L. Jing, Clustering ensemble based
938 on sample’s stability, *Artificial Intelligence* 273 (2019) 37–55.
- 939 [11] F. Wu, A. Souza, T. Zhang, C. Fifty, T. Yu, K. Weinberger, Simplifying
940 graph convolutional networks, in: *International conference on machine*
941 *learning*, PMLR, 2019, pp. 6861–6871.
- 942 [12] G. Wang, R. Ying, J. Huang, J. Leskovec, Improving graph atten-
943 tion networks with large margin-based constraints, *arXiv preprint*
944 *arXiv:1910.11945* (2019).
- 945 [13] J. Klicpera, S. Weißenberger, S. Günnemann, Diffusion improves graph
946 learning, in: *Advances in Neural Information Processing Systems*, 2019,
947 pp. 13354–13366.

- 948 [14] E. Chien, J. Peng, P. Li, O. Milenkovic, Adaptive universal general-
949 ized pagerank graph neural network, in: International Conference on
950 Learning Representations, 2021.
- 951 [15] G. Corso, L. Cavalleri, D. Beaini, P. Liò, P. Veličković, Principal neigh-
952 bourhood aggregation for graph nets, arXiv preprint arXiv:2004.05718
953 (2020).
- 954 [16] S. Zhang, L. Xie, Improving attention mechanism in graph neural net-
955 works via cardinality preservation, in: IJCAI: proceedings of the con-
956 ference, volume 2020, NIH Public Access, 2020, p. 1395.
- 957 [17] A. Wijesinghe, Q. Wang, A new perspective on "how graph neural
958 networks go beyond weisfeiler-lehman?", in: International Conference
959 on Learning Representations, 2022.
- 960 [18] K. Xu, W. Hu, J. Leskovec, S. Jegelka, How powerful are graph neural
961 networks?, arXiv preprint arXiv:1810.00826 (2018).
- 962 [19] C. Morris, F. Geerts, J. Tönshoff, M. Grohe, Wl meet vc, arXiv preprint
963 arXiv:2301.11039 (2023).
- 964 [20] B. Weisfeiler, A. Leman, The reduction of a graph to canonical form
965 and the algebrgra which appears therein, NTI, Series 2 (1968).
- 966 [21] T. He, Y. S. Ong, L. Bai, Learning conjoint attentions for graph neural
967 nets, Advances in Neural Information Processing Systems 34 (2021)
968 2641–2653.

- 969 [22] A. Hasanzadeh, E. Hajiramezanali, S. Boluki, M. Zhou, N. Duffield,
970 K. Narayanan, X. Qian, Bayesian graph neural networks with adaptive
971 connection sampling, in: International conference on machine learning,
972 2020.
- 973 [23] K. Yao, J. Liang, J. Liang, M. Li, F. Cao, Multi-view graph convo-
974 lutional networks with attention mechanism, Artificial Intelligence 307
975 (2022) 103708.
- 976 [24] Y. Liang, F. Meng, Y. Zhang, Y. Chen, J. Xu, J. Zhou, Emotional con-
977 versation generation with heterogeneous graph neural network, Artificial
978 Intelligence 308 (2022) 103714.
- 979 [25] L. Huang, D. Ma, S. Li, X. Zhang, H. Wang, Text level graph neural
980 network for text classification, in: Proceedings of the 2019 Conference
981 on Empirical Methods in Natural Language Processing and the 9th In-
982 ternational Joint Conference on Natural Language Processing (EMNLP-
983 IJCNLP), 2019, pp. 3444–3450.
- 984 [26] J. You, J. M. Gomes-Selman, R. Ying, J. Leskovec, Identity-aware graph
985 neural networks, in: Proceedings of the AAAI conference on artificial
986 intelligence, volume 35, 2021, pp. 10737–10745.
- 987 [27] V. Garg, S. Jegelka, T. Jaakkola, Generalization and representational
988 limits of graph neural networks, in: International Conference on Ma-
989 chine Learning, PMLR, 2020, pp. 3419–3430.
- 990 [28] D. Chen, L. O’Bray, K. Borgwardt, Structure-aware transformer for

- graph representation learning, in: International Conference on Machine Learning, PMLR, 2022, pp. 3469–3489.
- [29] F. M. Bianchi, D. Grattarola, L. Livi, C. Alippi, Graph neural networks with convolutional arma filters, *IEEE Transactions on Pattern Analysis and Machine Intelligence* 44 (2022) 3496–3507.
- [30] K. Xu, C. Li, Y. Tian, T. Sonobe, K.-i. Kawarabayashi, S. Jegelka, Representation learning on graphs with jumping knowledge networks, in: International Conference on Machine Learning, PMLR, 2018, pp. 5453–5462.
- [31] H. Gao, S. Ji, Graph representation learning via hard and channel-wise attention networks, in: Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining, 2019, pp. 741–749.
- [32] X. He, L. Li, D. Roqueiro, K. Borgwardt, Multi-view spectral clustering on conflicting views, in: Machine Learning and Knowledge Discovery in Databases: European Conference, ECML PKDD 2017, Skopje, Macedonia, September 18–22, 2017, Proceedings, Part II 10, Springer, 2017, pp. 826–842.
- [33] Y. Liu, Y. Liu, K. Chan, Ordinal regression via manifold learning, in: Proceedings of the AAAI Conference on Artificial Intelligence, volume 25, 2011, pp. 398–403.
- [34] F. Monti, D. Boscaini, J. Masci, E. Rodola, J. Svoboda, M. M. Bronstein, Geometric deep learning on graphs and manifolds using mixture

- 1014 model cnns, in: Proceedings of the IEEE Conference on Computer Vi-
1015 sion and Pattern Recognition, 2017, pp. 5115–5124.
- 1016 [35] S. Brody, U. Alon, E. Yahav, How attentive are graph attention net-
1017 works?, in: International Conference on Learning Representations, 2022.
- 1018 [36] P. Sen, G. Namata, M. Bilgic, L. Getoor, B. Galligher, T. Eliassi-Rad,
1019 Collective classification in network data, *AI magazine* 29 (2008) 93–93.
- 1020 [37] X. Huang, J. Li, X. Hu, Label informed attributed network embedding,
1021 in: Proceedings of the tenth ACM international conference on web search
1022 and data mining, 2017, pp. 731–739.
- 1023 [38] O. Shchur, M. Mumme, A. Bojchevski, S. Günnemann, Pitfalls of graph
1024 neural network evaluation, *arXiv preprint arXiv:1811.05868* (2018).
- 1025 [39] D. Greene, D. O’Callaghan, P. Cunningham, How Many Topics? Sta-
1026 bility Analysis for Topic Models, in: *Proc. European Conference on*
1027 *Machine Learning (ECML’14)*, 2014.
- 1028 [40] D. P. Kingma, J. Ba, Adam: A method for stochastic optimization,
1029 *arXiv preprint arXiv:1412.6980* (2014).
- 1030 [41] Z. Yang, W. Cohen, R. Salakhudinov, Revisiting semi-supervised learn-
1031 ing with graph embeddings, in: *International conference on machine*
1032 *learning*, PMLR, 2016, pp. 40–48.
- 1033 [42] X. Liang, Y. Qian, Q. Guo, H. Cheng, J. Liang, Af: An association-
1034 based fusion method for multi-modal classification, *IEEE Transactions*
1035 *on Pattern Analysis and Machine Intelligence* 44 (2021) 9236–9254.

1036 **Appendix A. Complete results of data characteristics and learned**
1037 **message weights**

1038 *Appendix A.1. Feature and structure distances of all test datasets*

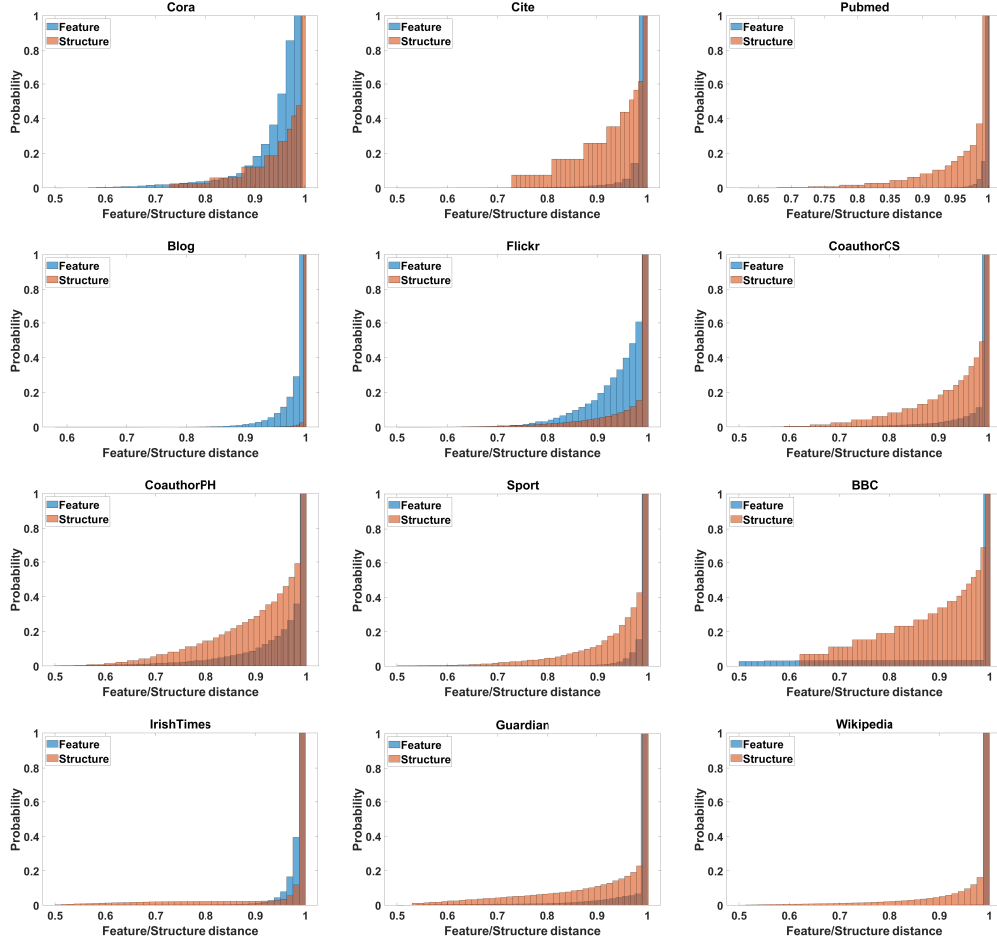


Figure A.3: Feature and structure distances of all test datasets.

1039 *Appendix A.2. Message weights learned by PMP-based GNNs and conven-*
1040 *tional message-passing GNNs*

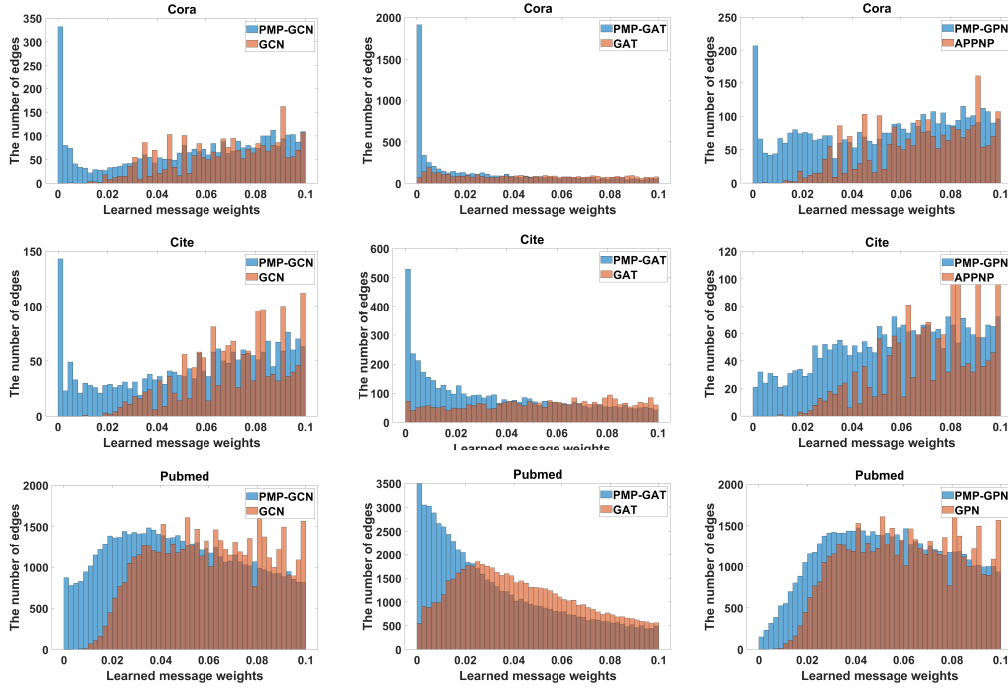


Figure A.4: Small message weights learned by PMP-based GNNs and conventional message-passing GNNs on article citation graphs.

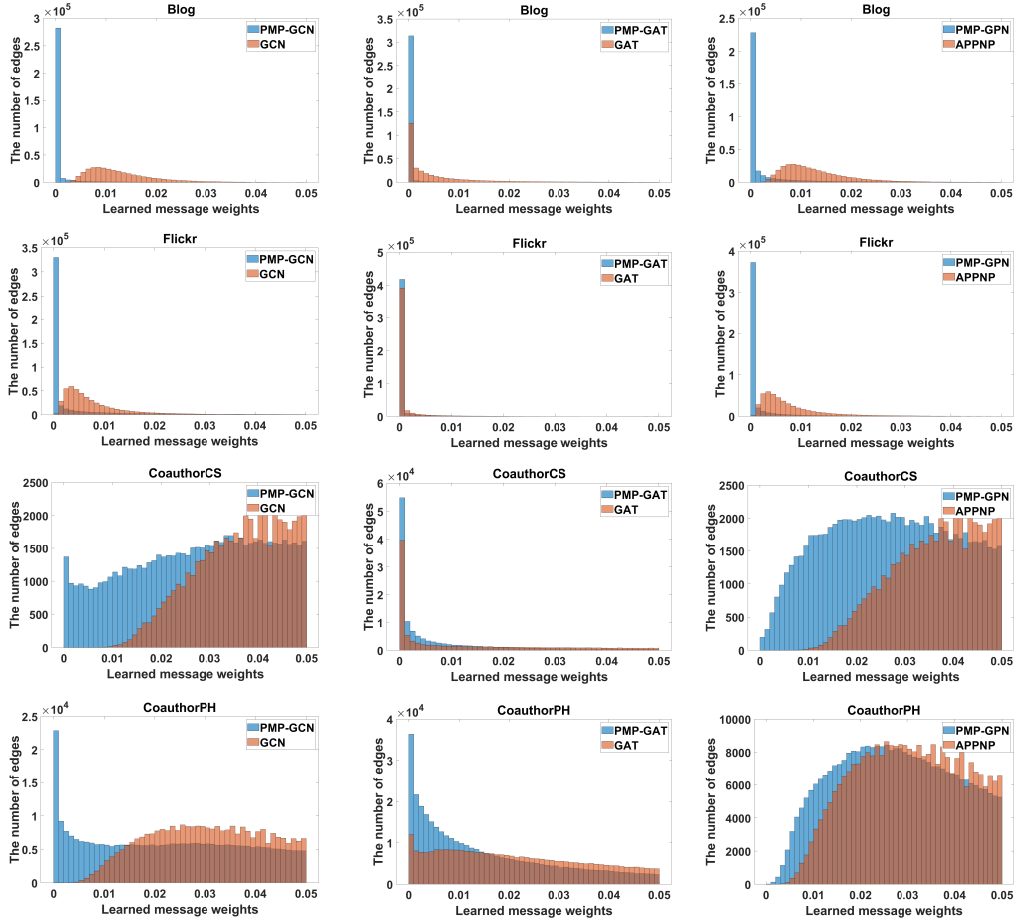


Figure A.5: Small message weights learned by PMP-based GNNs and conventional message-passing GNNs on social and collaboration graphs.

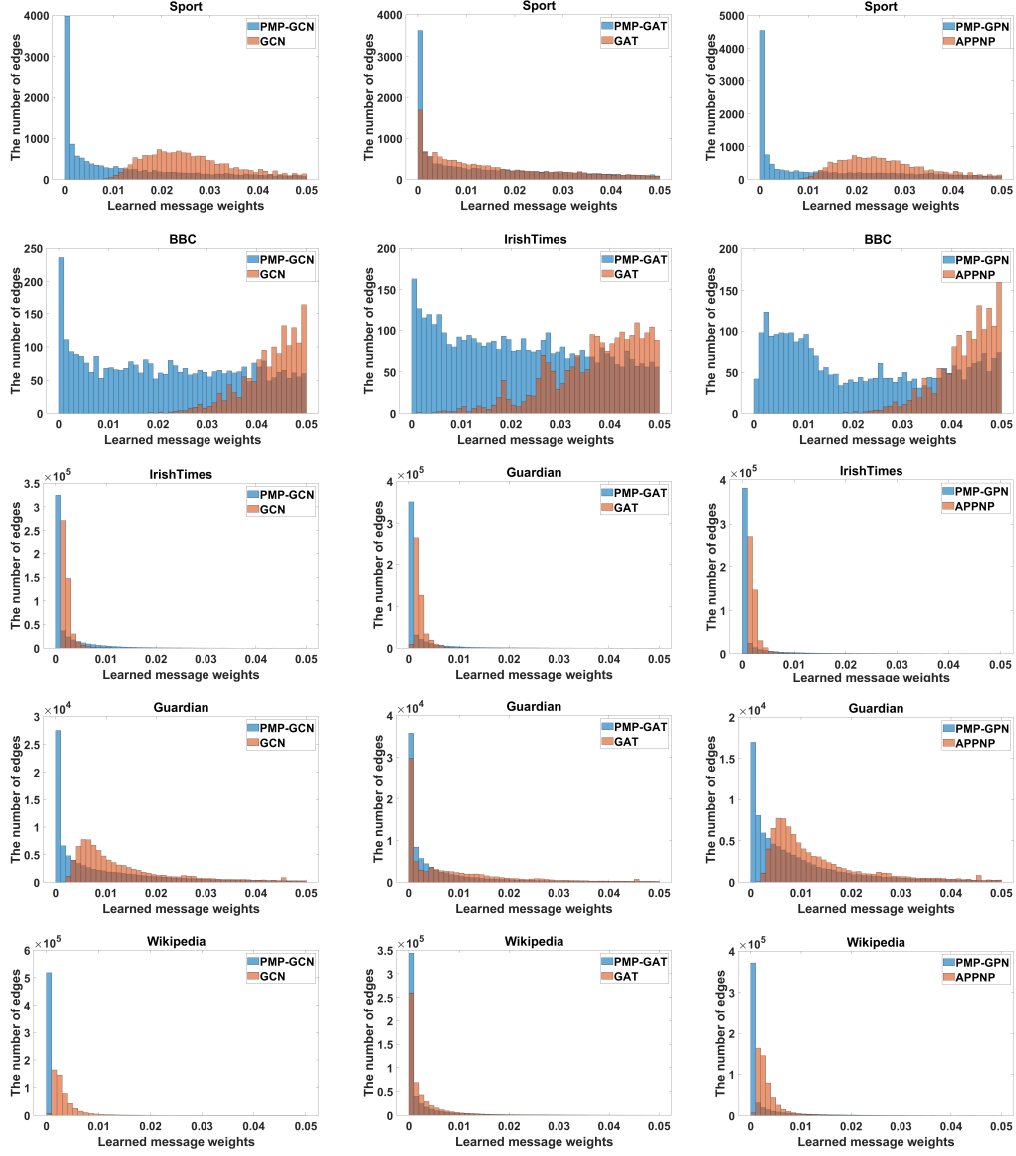


Figure A.6: Small message weights learned by PMP-based GNNs and conventional message-passing GNNs on rich-text graphs.

1041 Appendix B. Extension of Theorem 1

1042 When multiple data sources or functions are used to compute similarity
 1043 and dissimilarity between pairwise nodes, the proposed PMP is still more
 1044 expressive than conventional message-passing GNNs. Several well-defined
 1045 measures, such as normalized mutual information, distance correlation, and
 1046 maximal information coefficient [42], can be adapted to reveal the previously
 1047 mentioned property of the proposed PMP. In this paper, we achieve this by
 1048 further extending the proof of Theorem 1. Assume that C and S represent the
 1049 polarized message sets that PMP has already used to learn representations.
 1050 And let C' and S' represent new polarized message sets computed according
 1051 to new data sources or functions. When $f_\psi(\cdot)$ stays unchanged, for a node
 1052 label Y , we can verify the following:

$$\begin{aligned}
 & I(Y; [f_\psi(C), f_\psi(S)], [f_\psi(C'), f_\psi(S')]) \\
 &= H([f_\psi(C), f_\psi(S)], [f_\psi(C'), f_\psi(S')]) \\
 &\quad - H([f_\psi(C), f_\psi(S)], [f_\psi(C'), f_\psi(S')] | Y), \\
 &= H([f_\psi(C), f_\psi(S)], [f_\psi(C'), f_\psi(S')]) + H(Y) \\
 &\quad - H([f_\psi(C), f_\psi(S)], [f_\psi(C'), f_\psi(S')], Y) \\
 &= H(Y) - H(Y | [f_\psi(C), f_\psi(S)], [f_\psi(C'), f_\psi(S')]) \\
 &= H(Y) - H(Y | [f_\psi(C), f_\psi(S)]) + H(Y | [f_\psi(C), f_\psi(S)]) \\
 &\quad - H(Y | [f_\psi(C), f_\psi(S)], [f_\psi(C'), f_\psi(S')]) \\
 &= I(Y; [f_\psi(C), f_\psi(S)]) + I(Y; [f_\psi(C'), f_\psi(S')] | [f_\psi(C), f_\psi(S)]) \\
 &\geq I(Y; [f_\psi(C), f_\psi(S)]).
 \end{aligned} \tag{B.1}$$

1053 The analysis above shows that the learning capability of the proposed PMP
 1054 can be further improved as long as the new data sources or functions are

1055 appropriately used to compute polarized messages for the subsequent learning
1056 of representations (i.e., $I(Y; [f_\psi(C'), f_\psi(S')][f_\psi(C), f_\psi(S)]) > 0$).