

به نام خدا

ریاست محترم پژوهشکده علوم کامپیوتر پژوهشگاه دانش‌های بنیادی

با اهدای سلام و احترام

اینجانب مصطفی صالحی، عضو هیات علمی و دانشیار دانشگاه تهران، به منظور اجراء طرح "تعبیه گراف در

شبکه‌های چندلایه"، تقاضای تنظیم قرارداد پژوهشی تک پروژه غیرمقیم را از آن پژوهشکده دارم.

سه نفر از متخصصین صاحب نظر در این حوزه پژوهشی عبارتند از:

- آقای دکتر حمیدرضا ربیعی، عضو هیات علمی دانشگاه صنعتی شریف، ایمیل: rabiee@sharif.edu
- آقای دکتر سامان هراتی‌زاده، عضو هیات علمی دانشگاه تهران، ایمیل: haratizadeh@ut.ac.ir
- آقای دکتر بهزاد اکبری، عضو هیات علمی دانشگاه تربیت مدرس، ایمیل: b.akbari@modares.ac.ir

به پیوست جزئیات طرح پیشنهادی آورده شده است.

با تشکر

مصطفی صالحی

عضو هیات علمی و دانشیار رسمی دانشگاه تهران

موسس آزمایشگاه شبکه‌های کامپیوتری و اطلاعاتی دانشگاه تهران، <https://coinlab.ut.ac.ir/>

ایمیل: mostafa_salehi@ut.ac.ir

عنوان طرح پیشنهادی

پیش‌بینی یال با استفاده از روش‌های تعبیه گراف در شبکه‌های چندلایه

مجری اصلی

مصطفی صالحی، عضو هیات علمی و دانشیار دانشگاه تهران

موسس آزمایشگاه شبکه‌های کامپیوتری و اطلاعاتی دانشگاه تهران، <https://coinlab.ut.ac.ir/>

آدرس: تهران، خ کارگر شمالی، روبروی خیابان دهم، دانشکده علوم و فنون نوین، طبقه سوم، اتاق ۳۳۷

تلفن همراه: ۰۹۱۲۳۴۰۳۹۷۷، تلفن ثابت: ۰۲۱۸۶۰۹۳۲۹۸، آدرس ایمیل: mostafa_salehi@ut.ac.ir

چکیده

امروزه شاهد مدل کردن بسیاری از سیستم‌های دنیای واقعی به عنوان شبکه‌های اطلاعاتی هستیم، هر شبکه همچون شبکه‌های اجتماعی، شبکه‌های زیستی و شبکه‌های ارتباطی، شامل مجموعه‌ای از موجودیت‌هاست که از طریق ارتباطات گوناگون می‌توانند با یکدیگر در تعامل باشند. استخراج اطلاعات از این شبکه‌ها و تحلیل آن‌ها از اهمیت ویژه‌ای برخوردار است. روش‌های زیادی برای انجام این تحلیل‌ها پیشنهاد شده که به تازگی، روش‌هایی همچون تعبیه گراف که از بازنمایی گره‌های گراف در یک فضای برداری استفاده می‌کنند، توجه جامعه پژوهشی را به خود منعطف ساخته‌است. یکی از کاربردهای تعبیه گراف، پیش‌بینی یال یا ارتباطات آتی است که در حوزه‌های مختلف از جمله پزشکی، پیشنهاد کالا به مشتریان، پیشنهاد دوست به کاربران شبکه‌های اجتماعی کاربرد دارد. حال آنکه، بسیاری از شبکه‌های جدید به صورت شبکه‌های چندلایه مدل می‌شوند، این شبکه‌ها علاوه بر اطلاعات درون‌لایه‌ای، برخی اطلاعات بین‌لایه‌ای نیز ارائه می‌کنند. مدل‌هایی که برای استفاده بر روی شبکه‌های تک‌لایه طراحی شده‌اند، به دلیل تحلیل مجزای هر لایه از شبکه، موجب ساده‌سازی بیش از حد و از دست رفتن اطلاعات بین‌لایه‌ای می‌شوند. در طرح پیشنهادی چارچوبی ارائه می‌دهیم که قابلیت استفاده بر روی شبکه‌های تک‌لایه و چندلایه را داشته‌باشد، برای لایه‌های مختلف تعبیه متناسب با گره‌ها و یال‌های آن لایه تولید کند، سپس از تعبیه‌های به دست آمده و اطلاعات بین‌لایه‌ای موجود در شبکه‌های چندلایه، برای پیش‌بینی پیوندهای آتی در لایه‌های مختلف استفاده کند. همچنین در این مدل، مشکل روش‌های موجود، در نحوه استفاده از دادگانی که اطلاعاتی از ویژگی گره‌های آن‌ها در دسترس نیست، به صورت بهیبه‌تری رفع شده و روشی جهت تولید ویژگی برای گره‌ها در این نوع از شبکه‌ها، پیشنهاد می‌گردد.

واژه‌های کلیدی:

تعبیه گراف، بازنمایی گراف، شبکه‌های چندلایه، پیش‌بینی یال، یادگیری عمیق

Title:

Graph Embedding in Multilayer Networks

Abstract:

Nowadays, many real-world systems are being modeled as information networks. Each network consists set of entities that can interact with each other through various communications. Extracting information from these networks and analyzing them is of particular importance. Graphs like social networks, biological networks, and communication networks are abundantly found in many different real-world applications. Various methods have been proposed to perform these analyses. Recently, methods like graph embedding, which uses the representation of graph nodes in vector space, have attracted the attention of the research community. Moreover, many new networks are modeled as multilayer networks. These types of networks provide some inter-layer information in addition to intra-layer information. The models designed for use on single-layer networks cause oversimplification and loss of inter-layer information, due to the separate analysis of each layer of the network. Therefore, newer models were introduced for use in multilayer networks. One of the applications of graph embedding is to predict potential and future edges or connections that do not currently exist in the graph. Today, we are witnessing the increasing progress of the need to link prediction in various fields, such as medicine, product recommendation, and offering friends to social network users. To do so, various methods have been introduced, both in the field of single-layer and multilayer networks, but the existing methods have problems such as high time complexity, on the other hand, they do not provide a suitable solution for using networks where information about the characteristics of their nodes is not available. In this project, we present a framework that can more optimally solve the problem of the unavailability of node features, generate different embedding for different layers, then by utilizing these embeddings besides inter-layer information, try to predict future links in every layer.

Keywords:

Graph Embedding, Graph Representation Learning, Multilayer Networks, Link Prediction, Deep Learning

فهرست مطالب

۱- مقدمه	۳
۲- پیشینه تحقیق	۱۸
۳- برنامه آتی	۴۰
۴- مراجع	۴۱

۱ - مقدمه

تجزیه و تحلیل گراف در سال‌های اخیر به دلیل گستردگی شبکه‌ها در دنیای واقعی، توجه زیادی را به خود جلب کرده است. گراف‌ها (شبکه‌ها) برای بیان و نمایش اطلاعات در زمینه‌های مختلف از جمله زیست‌شناسی (شبکه‌های متقابل پروتئین-پروتئین)، علوم اجتماعی (شبکه‌های دوستی) و زبان‌شناسی (شبکه‌های کلمات هم‌رخداد)^۱ مورد استفاده قرار گرفته است. مدل‌سازی تعاملات متقابل موجودیت‌ها در قالب گراف، محققان را قادر ساخته است تا سیستم‌های شبکه‌ای مختلف را به‌صورتی سیستماتیک و نظام‌مندتر درک کنند. به عنوان مثال، از شبکه‌های اجتماعی^۲ برای اهدافی همچون دوست‌یابی، توصیه محتوا^۳ و همچنین تبلیغات استفاده شده است [۱]. از دیرباز یکی از مسائل مهم مطرح در اکثر سیستم‌های دنیای واقعی، پیش‌بینی آینده آن‌هاست، که با مدل شدن این سیستم‌های پیچیده در قالب شبکه‌های اطلاعاتی، این مسئله به مسئله پیش‌بینی یال [۲] مبدل می‌گردد.

پژوهشگران در طول سال‌های اخیر دریافته‌اند که مدل کردن برخی از شبکه‌های پیچیده به شبکه‌هایی که تنها یک نوع یال و گره دارند منجر به از دست رفتن بخش زیادی از اطلاعات و نهایتاً موجب تحلیل ناقص و همچنین پیش‌بینی اشتباه می‌گردد، آن‌ها برای حل این چالش شبکه‌های اطلاعاتی چندلایه را مطرح کردند که در آن، یال‌ها می‌توانند از انواع مختلفی باشند.

در ادامه، چشم‌اندازی کلی از انواع شبکه‌های اطلاعاتی، روش‌های تعبیه گراف و همچنین کاربردهای آن از جمله پیش‌بینی یال در آن‌ها را ارائه می‌دهیم. هدف از این بخش ترسیم یک تصویر و دید کلی، همچنین ایجاد پس‌زمینه مناسب برای مطالب موجود در بخش‌های بعدی می‌باشد.

یک شبکه اطلاعاتی^۴ یک انتزاع از اطلاعات موجود در جهان واقع است به‌گونه‌ای که تمرکز عمده این شبکه اطلاعاتی بر روی نحوه اتصال گره‌ها و تعاملات بین آن‌ها می‌باشد. این سطح از انتزاع نه تنها در ارائه و مرتب‌سازی اطلاعات جهان واقع نقش دارد بلکه یک ابزار کاربردی و اصلی برای به‌دست آوردن اطلاعات و دانش موجود در آن می‌باشد [۳]. حال در ادامه به تعریف رسمی شبکه‌های اطلاعاتی می‌پردازیم.

¹ Protein-Protein interaction networks

² Word co-occurrence networks

³ Social networks

⁴ Content recommendation

⁵ Information Networks

تعریف ۱-۱ شبکه‌های اطلاعاتی

یک گراف جهت دار به صورت $G(V, E)$ را در نظر بگیرید، همچنین یک تابع نگاشت نوع اشیا $t: V \rightarrow A$ و تابع نگاشت نوع یال $\phi: E \rightarrow R$ که در آن هر گره $v \in V$ متعلق به یک موجودیت خاص $t(v) \in A$ و هر یال $e \in E$ متعلق به یک رابطه خاص $\phi(e) \in R$ می‌باشد. حال دو یال متعلق به یک رابطه هستند، اگر نقاط ابتدایی هر دو یال از یک نوع گره و نقاط انتهایی آن‌ها نیز از یک نوع گره باشد [۳].

۱-۱- انواع شبکه‌های اطلاعاتی

به‌طور کلی می‌توان شبکه‌های اطلاعاتی را به دو گروه عمده شبکه‌های تک‌لایه^۱ و یا شبکه‌های چندلایه^۲ تقسیم‌بندی کرد. در ادامه به معرفی هر یک از این شبکه‌ها خواهیم پرداخت.

• شبکه‌های تک‌لایه

شبکه تک‌لایه شبکه‌ای است که دارای گره‌های یکسان می‌باشد که در آن گره‌ها فقط از یک طریق امکان ارتباط با یکدیگر را دارند. در برخی مقالات از این نوع شبکه‌ها به‌عنوان شبکه‌های همگن^۳ نیز یاد می‌شود. در شکل بالا مثالی از این شبکه آورده شده که گره‌ها نویسندگان مقالات و یال‌ها ارتباط بین آن‌ها (ارجاعات) را نشان می‌دهد.

• شبکه‌های چندلایه

این شبکه از اجتماع چندین شبکه تک‌لایه تشکیل شده به گونه‌ای که نوع گره‌ها در هر لایه از لایه دیگر متفاوت می‌باشد و در نتیجه نوع ارتباطات در هر لایه با لایه دیگر تفاوت خواهد داشت. یک نوع خاصی از شبکه‌های چندلایه وجود دارد که شبکه‌های چندگانه^۴ نام دارد، در این نوع شبکه گره‌ها در بین همه لایه‌ها یکسان می‌باشد اما این نوع ارتباطات بین گره‌هاست که در هر لایه با دیگری متفاوت می‌باشد [۴].

تعریف ۲-۱- شبکه‌های چندگانه

¹ Single-Layer

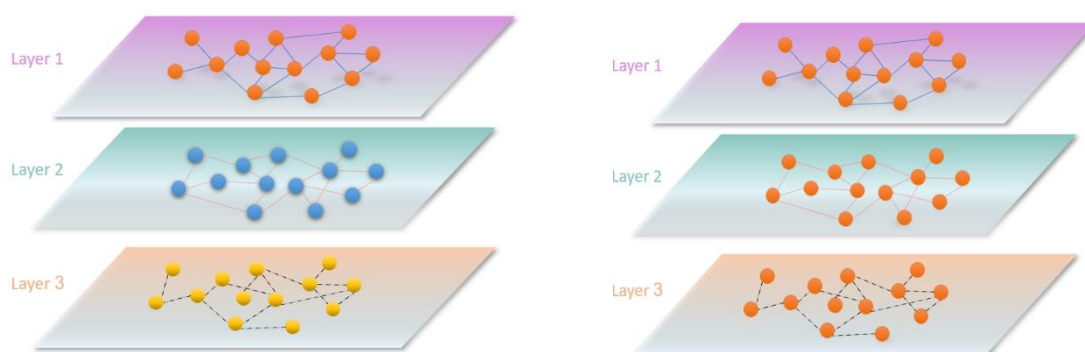
² Multi-Layer

³ Homogeneous Networks

⁴ Multiplex Network

فرض کنیم $G = (V, E_1 E_2 \dots E_N)$ یک شبکه چندگانه باشد که در آن V نشان دهنده مجموعه گره‌های هم‌نوع بوده و $E_j (j \in 1, 2, \dots, N)$ نیز مجموعه یال‌های نوع j باشد، همچنین N معادل تعداد لایه‌های این شبکه می‌باشد.

از این نوع شبکه‌ها در دنیای واقعی به وفور یافت می‌شود. برای نمونه می‌توان از شبکه مواصلاتی بین شهرها نام برد. اگر شهرها را به عنوان گره‌های این شبکه در نظر بگیریم روش‌های مواصلاتی مختلف بین این شهرها از جمله هوایی، زمینی و ریلی، لایه‌های مختلف این شبکه خواهند بود که نوع ارتباط بین گره‌ها درون لایه‌ها با



شکل ۱-۱. (سمت راست): نمایی از یک شبکه چندگانه که در آن نوع گره‌ها یکسان ولی ارتباطات در هر لایه متفاوت است. (سمت چپ): نمایی از یک شبکه چندلایه که در آن نوع گره‌ها و ارتباطات، هر دو، در هر لایه متفاوت هستند. یکدیگر متفاوت است.

باوجود تفاوت‌های یادشده بین این دو نوع شبکه، در بسیاری از موارد از کلمه شبکه‌های چندلایه به جای شبکه‌های چندگانه استفاده می‌شود، به همین جهت ما نیز در ادامه برای ساده‌سازی به جای کلمه شبکه‌های چندگانه از شبکه‌های چندلایه استفاده خواهیم کرد.

۱-۲- تعبیه گراف

جدا از بحث سیستم‌های اطلاعاتی و مزیت‌های سیستم‌های چندلایه، یکی از مسائل اصلی مطرح در این حوزه، بحث یادگیری ماشین^۱ روی این گراف‌ها و یافتن راهی برای ارائه و استخراج اطلاعات موجود در ساختار گراف به یک مدل یادگیری ماشین است. به عنوان مثال، در مسئله پیش‌بینی لینک در یک شبکه اجتماعی، ممکن است شخصی بخواهد خصوصیات بین جفت گره‌ها، از قبیل قدرت رابطه یا تعداد دوستان مشترک را مورد

^۱ Machine Learning

توجه قرار دهد. یا در مورد دسته‌بندی گره، ممکن است اطلاعاتی درباره موقعیت سراسری یک گره در گراف یا ساختار همسایگی محلی گره برای ما مهم باشد. چالشی که از دیدگاه یادگیری ماشینی وجود دارد این است که هیچ روش سراسری برای رمزگذاری^۱ و تبدیل چنین اطلاعاتی با بُعد بالا، از ساختار گراف به یک بردار ویژگی وجود ندارد [۵]. برای استخراج اطلاعات ساختاری از گراف‌ها، رویکردهای سنتی اغلب به خلاصه برداری از ویژگی‌های آماری گراف‌ها (به عنوان مثال، درجه یا ضرایب خوشه‌بندی)^۲ تکیه دارند.

با این حال، این رویکردها محدود هستند چرا که نمی‌توانند در طول یادگیری سازگار شوند. اخیراً رویکردهای زیادی پدید آمده‌اند که به دنبال یادگیری بازنمایی‌هایی^۳ هستند تا اطلاعات ساختاری یک گراف را رمزگذاری^۴ کنند. ایده این رویکردها، یادگیری نگاشتی است که گره‌ها یا کل گراف (یا زیرگراف) را به عنوان نقاطی در یک فضای برداری با بعد کم‌تر بکار برند. هدف، بهینه‌سازی این نگاشت است، به گونه‌ای که روابط بین گره‌ها در فضای تعبیه شده^۵ ساختار گراف اصلی را منعکس کند. حال پس از بهینه‌سازی فضای تعبیه، تعبیه‌های آموخته شده می‌توانند به عنوان ورودی برای الگوریتم‌های یادگیری ماشین مورد استفاده قرار گیرند. تمایز اصلی بین روش‌های یادگیری بازنمایی^۶ و کارهای قبلی نحوه برخورد آن‌ها با مسئله بازنمایی ساختار گراف است. کارهای قبلی این مشکل را به عنوان یک مرحله پیش‌پردازشی، با استفاده از آمارهای گرافی که به صورت دستی برای استخراج اطلاعات ساختاری گراف‌ها تعریف می‌شدند، حل می‌کرد. در مقابل، رویکردهای یادگیری بازنمایی این مشکل را به عنوان جزئی از وظیفه یادگیری ماشین، با استفاده از یک روش داده‌محور برای یادگیری تعبیه‌هایی که ساختار گراف را رمزگذاری می‌کنند، حل می‌کند [۶].

تعبیه گراف^۷ روشی است برای تبدیل گره‌ها، یال‌ها و ویژگی‌های آن‌ها به فضای برداری (بعد پایین‌تر) درعین حفظ حداکثری خواصی مانند ساختار گراف، اطلاعات گره‌ها و یال‌ها. گراف‌ها بسیار پیچیده هستند زیرا می‌توانند از نظر مقیاس، ویژگی و محتوا متفاوت باشند یک ساختار مولکولی می‌تواند به عنوان یک گراف کوچک، به صورت استاتیک نمایش داده شود، درحالی‌که یک شبکه اجتماعی می‌تواند توسط یک گراف بزرگ، متراکم و پویا به نمایش درآید. درنهایت، این امر پیدا کردن یک روش تعبیه که بتوان در شبکه‌های مختلف از آن استفاده کرد را دشوارتر می‌کند. روش‌های مختلف برای اهداف مختلف ارائه می‌شود که عملکرد هریک در

¹ Encode

² Clustering coefficient

³ Representation

⁴ Encoding

⁵ Embedding

⁶ Representation learning

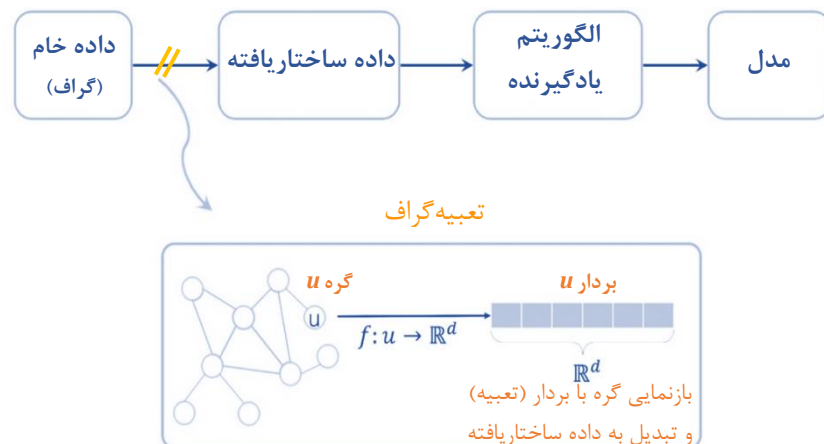
⁷ Graph embedding

مجموعه داده‌های مختلف متفاوت است [۷]. قبل از ارائه تعریف تعبیه باید اشاره کرد که اصطلاح تعبیه گراف (بازنمایی گراف) در ادبیات به دو صورت مختلف مورد استفاده قرار گرفته است:

(آ). بازنمایی کل گراف به صورت یکجا در فضای برداری

(ب). بازنمایی هر گره به صورت مجزا در فضای برداری.

ما در این طرح، از رویکرد دوم استفاده می‌کنیم و تعاریف بر مبنای رویکرد دوم می‌باشد.



شکل ۱-۲. فرآیند تبدیل داده گرافی به داده‌ی قابل استفاده برای الگوریتم‌های یادگیری ماشین و ساخت مدل

تعریف ۱-۳- تعبیه گراف: با توجه به گراف $G(V, E)$ ، تعبیه گراف یک تابع نگاشت^۱ است:

$f: v_i \rightarrow Y_i$ به گونه‌ای که $d \ll |X|$ (ابعاد فضای نگاشت شده از فضای قبلی بسیار کمتر است) و تابع f معیار شباهت تعریف شده در گراف G را حفظ کند (Y_i معادل تعبیه گره i و v_i : گره i ، d : ابعاد تعبیه، V : گره‌های گراف، E : یال‌های آن و X بیانگر بردار ویژگی گره‌ها می‌باشد). بنابراین تعبیه، هر گره را به یک بردار ویژگی^۲ با ابعاد کمتر نگاشت و سعی می‌کند میزان قدرت اتصال بین گره‌ها را حفظ کند. به عنوان مثال، فرض کنیم ماتریس S ماتریس شباهت گره‌های موجود در گراف باشد و وزن یال S_{ij} به عنوان مقیاسی برای شباهت بین گره‌های v_i و v_j در نظر گرفته شود، حال اگر $S_{ij} > S_{ik}$ باشد، v_i و v_j به گونه‌ای نگاشت می‌شوند که در فضای جدید گره j به گره i نزدیک‌تر از فاصله گره k به گره i باشد.

¹ Graph Representation

² Mapping function

³ Feature vector

بنابراین هدف، رمزگذاری و نگاشت گره‌ها به گونه‌ای است که شباهت^۱ و نزدیکی گره‌ها در فضای تعبیه، تقریبی از شباهت آن‌ها در گراف اصلی باشد:

$$\text{شباهت } (u, v) \text{ در گراف اصلی} \approx Y_v Y_u^T \text{ (شباهت } (u, v) \text{ در فضای تعبیه (ضرب داخلی دو بردار))}$$

مسئله‌ای که اینجا به چشم می‌خورد، تعریف شباهت است، اینکه چه زمانی دو گره را (در شبکه اصلی) شبیه هم گوئیم یا به بیان دیگر، چه زمانی دو گره، تعبیه یکسان یا شبیه به همی را خواهند داشت؟

(الف). حالتی که به هم‌دیگر متصل باشند و یالی بینشان برقرار باشد.

(ب). حالتی که همسایه‌های بیش‌تری باهم به اشتراک بگذارند.

(پ). و یا اینکه از لحاظ ساختاری شبیه به هم باشند و نقش‌های مشابهی در شبکه ایفا کنند.

معیارهای شباهت مختلفی را می‌توان برای گره‌ها در شبکه متصور شد و هریک از این تعریف‌ها بسته به نوع کاربرد و هدفی که از تعبیه داریم می‌توانند درست باشند.

۳-۱- تعریف موضوع

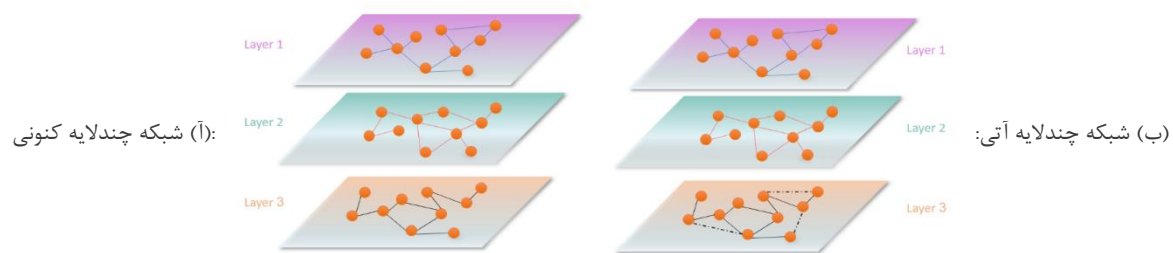
به‌طور کلی هدف تعبیه گراف در گام اول استخراج اطلاعات گراف و نگاشت آن‌ها به یک فضای برداری با بُعد کم‌تر و در گام دوم استفاده از تعبیه‌های به‌دست‌آمده در الگوریتم‌های یادگیری ماشین جهت کاربردهای مختلف از جمله دسته‌بندی گره‌ها، پیش‌بینی یال، خوشه‌بندی یا تصویرسازی می‌باشد که در ادامه به توضیح مختصر این کاربردها خواهیم پرداخت. تاکید عمده این پژوهش ارائه یک روش تعبیه بهینه و با دقت بالا برای گراف‌های دنیای واقعی همچون شبکه‌های اجتماعی و استفاده از این تعبیه در حوزه‌های کاربردی یادشده مانند پیش‌بینی یال می‌باشد. این روش تعبیه نه تنها باید قابلیت استفاده روی شبکه‌های تک‌لایه را دارا باشد بلکه با ارائه راهکاری نوین قابل استفاده بر روی شبکه‌های چندلایه نیز باشد و کاربردهایی از قبیل پیش‌بینی یال در شبکه‌های چندلایه را به‌انجام رساند.

مسئله پیش‌بینی یال در شبکه‌های مختلف پیش‌بینی یال‌های بالقوه‌ای است که در حال حاضر در شبکه موجود نیستند ولی احتمال آن می‌رود که در آینده در گراف به وجود آیند. می‌توان مسئله پیش‌بینی یال را به عنوان یک مسئله دسته‌بندی نیز در نظر گرفت. در این روش هر یال را می‌توان به عنوان یک جفت رأس در نظر گرفت لذا هر نقطه متناظر با یک جفت رای در شبکه خواهد بود و باتوجه به ویژگی‌ها و اطلاعات یال مدل در بازه آموزشی^۲ آموزش خواهد دید تا بتواند در بازه تست^۱ به پیش‌بینی یال‌هایی که ممکن است در آینده در گراف

^۱ Similarity

^۲ Training Phase

به وجود آیند پردازد [۸]. همین رویکرد را می‌توان به شبکه‌های چندلایه هم تعمیم داد با این تفاوت که در شبکه‌های چندلایه نوع یال‌ها در هر لایه متفاوت خواهد بود لذا باید به نوع یال‌ها هم توجه کرد. شکل پایین چارچوب و رویکر کلی مسئله پیش‌بینی یال در شبکه‌های سه لایه را نشان می‌دهد. در بخش (آ) حالت کنونی شبکه نشان داده شده است که مدل روی آن آموزش می‌بیند و بخش (ب) نتیجه پیش‌بینی مدل را نشان می‌دهد که یال‌های آتی را در لایه سه پیش‌بینی کرده‌است.



شکل ۱-۳. مسئله پیش‌بینی یال در شبکه‌های چندلایه (قبل و بعد از پیش‌بینی مدل)

۴-۱- اهمیت موضوع و کاربردها

به‌طور معمول، یک مدل تعریف شده برای حل مسائل مبتنی بر گراف، بر اساس ماتریس مجاورت^۱ گراف اصلی یا در یک فضای برداری مشتق‌شده از آن کار می‌کند. اخیراً روش‌های مبتنی بر نمایش شبکه‌ها در فضای برداری^۲ همچون تعبیه گراف، ضمن حفظ ویژگی‌های آن‌ها، بسیار مورد استقبال محققان قرار گرفته است [۹، ۱۰]. دستیابی به چنین تعبیه‌ای در کارهایی که درون تحلیل‌های گرافی انجام می‌پذیرد مفید است. در ادامه به معرفی این تحلیل‌ها که کاربردهای تعبیه‌گراف می‌باشند می‌پردازیم، این کاربردها به‌طور کلی به چهار دسته زیر خلاصه شود:

تصویرسازی، خوشه‌بندی، پیش‌بینی یال و دسته‌بندی گره.

¹ Test Phase

² Adjacency Matrix

³ Vector Space Representation

• تصویرسازی

کاربرد تصویرسازی گراف می‌تواند به سال ۱۷۳۶ میلادی برگردد، زمانی که اوپلر از آن برای حل مسئله پل‌های کونیگزبرگ^۱ استفاده کرد. در سال‌های اخیر، تصویرسازی گراف کاربردهای وسیعی در مهندسی نرم‌افزار، مدارهای الکترونیکی، زیست‌شناسی و جامعه‌شناسی داشته است. باتیستا و ادز طیف وسیعی از روش‌های مورد استفاده برای ترسیم گراف‌ها و تعیین معیارهای لازم برای این منظور را بررسی کرده‌اند. هرمان و همکاران این را تعمیم داده و از منظر تصویرسازی اطلاعات گراف را مورد مطالعه قرار می‌دهند [۱]. آن‌ها روش‌های سنتی مختلفی که برای ترسیم گراف‌ها از جمله طرح‌های درختی و سه بعدی استفاده می‌شوند را مطالعه و مقایسه می‌کنند.

از آنجایی که تعبیه، گراف را در یک فضای برداری نشان می‌دهد، می‌توان از تکنیک‌های کاهش بُعد مانند تجزیه و تحلیل مؤلفه اصلی (PCA) برای تصویرسازی گراف استفاده کرد. نویسندگان روش DeepWalk [۱۱] با تصویرسازی شبکه باشگاه کاراته زاکاری^۲، مزیت روش تعبیه خود را در تصویرسازی گراف نشان دادند. همچنین شبکه همکاری نویسندگان^۳ را تصویرسازی کرده و نشان دادند که این روش قادر است نویسندگان را بر اساس محتوا و موضوع مقالات خوشه‌بندی کند.

• خوشه‌بندی

خوشه‌بندی گراف می‌تواند دو نوع باشد: (الف) مبتنی بر ساختار و (ب) مبتنی بر ویژگی^۴. دسته اول را بازهم می‌توان به دو دسته تقسیم کرد، یعنی خوشه‌بندی مبتنی بر اجتماع و خوشه‌بندی مبتنی بر هم‌ساختاری. روش‌های مبتنی بر اجتماع [۱۲] هدفشان یافتن زیرگراف‌های متراکم^۵ (با تعداد زیادی یال‌های درون خوشه‌ای^۶ و تعداد کمی

¹ Königsberger Brücken Problem

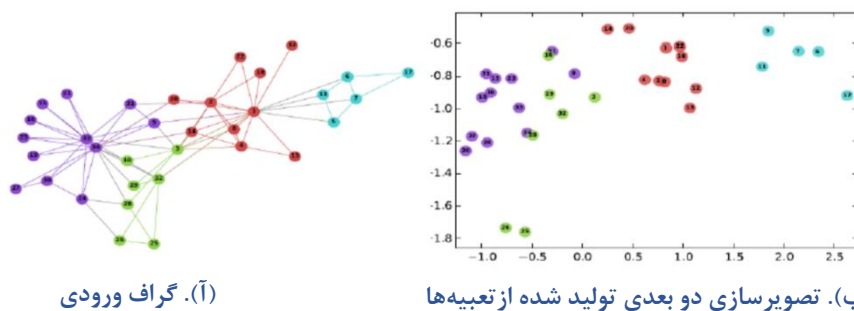
² Zachary's Karate Club

³ DBLP Co-Authorship Network

⁴ Attribute based

⁵ Dense Subgraphs

⁶ Intra-Cluster



شکل ۱-۴. (آ). شبکه اجتماعی باشگاه کاراته زاکاری (Zachary Karate Club)، که در آن گره‌های مربوط به افراد در صورت دوست‌بودن به هم متصل می‌شوند. همچنین گره‌ها با توجه به جوامع مختلفی که در شبکه وجود دارد رنگ‌آمیزی می‌شوند. (ب). تصویرسازی دو بعدی حاصل از تعبیه گره‌ها با استفاده از روش DeepWalk. فاصله بین گره‌ها در فضای تعبیه، شباهت آن‌ها را در گراف اصلی نشان می‌دهد، همچنین تعبیه گره‌ها با توجه به جوامع مختلف رنگ‌آمیزی شده و به صورت مکانی خوشه‌بندی می‌شوند [۱۱].

یال‌های بین خوشه‌ای می‌باشد. برعکس، خوشه‌بندی مبتنی بر هم‌ساختاری برای شناسایی گره‌ها با نقش‌های مشابه (مانند پل‌ها یا گره‌های دورافتاده) طراحی شده است. همچنین روش‌های مبتنی بر ویژگی، علاوه بر پیوندهای مشاهده‌شده، از برچسب گره‌ها به منظور خوشه‌بندی آن‌ها استفاده می‌کنند.

• دسته‌بندی گره‌ها

اغلب در شبکه‌ها، تنها کسری از گره‌ها برچسب گذاری می‌شوند. در شبکه‌های اجتماعی، برچسب‌ها ممکن است علائق، تمایلات، عقاید یا گروه خاصی را نشان دهند. در شبکه‌های زبانی، یک سند ممکن است بر اساس موضوعات یا کلمات کلیدی برچسب گذاری شود، در حالی که برچسب موجودیت‌ها در شبکه‌های زیستی ممکن است مبتنی بر قابلیت‌ها و کارکرد آن‌ها باشد. به دلیل عوامل مختلف، برچسب‌ها ممکن است برای بخش بزرگی از گره‌ها ناشناخته باشند.

برای مثال، در شبکه‌های اجتماعی بسیاری از کاربران به دلیل نگرانی از حریم خصوصی برخی از اطلاعات خود را ارائه نمی‌دهند. با استفاده از گره‌های برچسب زده شده و پیوندهای موجود در شبکه، برچسب‌های گم‌شده قابل استنباط و استنتاج هستند. وظیفه پیش‌بینی این برچسب‌های گم‌شده نیز به عنوان دسته‌بندی گره‌ها شناخته می‌شود. باگات و همکاران [۱۳] روش‌های مورد استفاده در ادبیات را برای این کار بررسی کرده و آن‌ها را به دو دسته طبقه‌بندی می‌کنند، دسته اول: روش‌های مبتنی بر استخراج ویژگی^۳ و دسته

¹ Inter-cluster

² Outliers

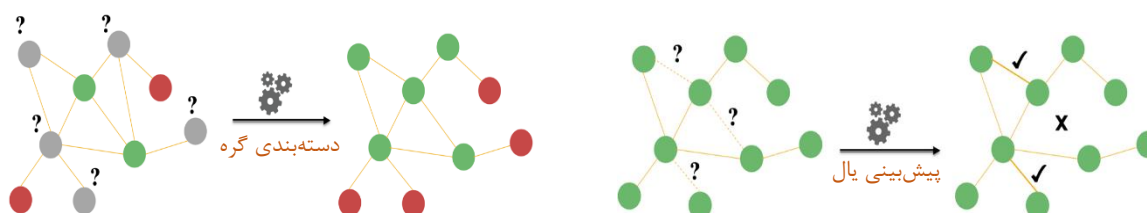
³ Feature Extraction Based Methods

دوم: قدم زدن تصادفی^۱ مدل های مبتنی بر ویژگی، ویژگی‌هایی را برای گره‌ها براساس همسایه‌ها و گره‌های محلی‌شان ایجاد می‌کنند و سپس برای پیش‌بینی برچسب‌ها از دسته‌بندی‌هایی^۲ همچون رگرسیون لجستیک^۳ استفاده می‌کنند. درمقابل، مدل های مبتنی بر قدم‌زنی تصادفی برچسب‌ها را با قدم‌زنی تصادفی و حرکت تصادفی از روی یک گره به دست می‌آورند [۱].

مقالات اخیر [۱۴] [۱۱] [۱۰] [۱] قدرت پیش‌بینی کننده تعبیه در شبکه‌های اطلاعاتی مختلف از جمله زبانی، اجتماعی، زیستی و گراف‌های همکاری را مورد ارزیابی قرار داده‌اند. آن‌ها نشان می‌دهند که تعبیه با دقت بالایی برچسب‌های گم‌شده را پیش‌بینی می‌کند.

• پیش‌بینی یال

شبکه‌ها از تعاملات مشاهده شده بین موجودیت‌ها ساخته می‌شوند که ممکن است مشاهدات ما ناقص یا نادرست باشد. این چالش در شناسایی تعاملات و پیش‌بینی اطلاعات از دست‌رفته بیش‌تر می‌شود. پیش‌بینی یال به وظیفه پیش‌بینی پیوندهای مفقود شده یا پیوندهایی که ممکن است در آینده در یک شبکه در حال گسترش



شکل ۱-۵. دسته‌بندی گره‌ها و پیش‌بینی یال از کاربردهای تعبیه

ظاهر شوند، اشاره دارد. پیش‌بینی یال در تجزیه و تحلیل شبکه‌های بیولوژیکی بسیار فراگیر و گسترده است، جایی که در آن تأیید وجود پیوند بین گره‌ها نیازمند آزمایشات پرهزینه‌ای می‌باشد. بنابراین محدود کردن آزمایش‌ها به پیوندهایی که احتمال وقوع آن‌ها بیش‌تر است، بسیار مقرون به صرفه می‌باشد. در شبکه‌های اجتماعی از پیش‌بینی یال برای پیش‌بینی دوستی‌های احتمالی آتی استفاده می‌شود که می‌تواند در سیستم‌های توصیه‌گر^۴ مورد استفاده قرار گیرد و منجر به یک تجربه کاربری رضایت بخش‌تر شود [۱۴].

¹ Random-Walk

² Classifier

³ Logistic regression

⁴ Recommendation Systems

• پیش‌بینی یال در شبکه‌های چندلایه

پیش‌بینی ارتباطات مابین گره‌ها در شبکه‌های چندلایه نیز اهمیت بسیار زیادی دارد، چراکه شبکه‌های دنیای واقعی بسیاری وجود دارند که چندلایه یا چندگانه هستند و پیش‌بینی ارتباط آتی بین یک جفت گره روی آن شبکه‌ها می‌تواند بسیار مفید باشد. در زیر چندین مثال از این کاربردها را مرور می‌کنیم.

- شبکه‌های اجتماعی^۱

در شبکه‌های اجتماعی نحوه ارتباط افراد با یکدیگر می‌تواند متفاوت باشد، به‌طور مثال افراد در شبکه اجتماعی توییتر^۲ می‌توانند به سه نوع مختلف پاسخ^۳، اشاره^۴ و ریتوییت^۵ با یکدیگر ارتباط برقرار کنند. به‌عنوان مثال زمانی که فردی توییتی را منتشر می‌کند افراد دیگر می‌توانند آن را ریتوییت کنند یا به آن توییت پاسخ دهند و یا فرد دیگری را مورد اشاره قرار دهند. هر یک از این ارتباطات را می‌توان به عنوان یک لایه از شبکه توییتر در نظر گرفت. در صورت پیش‌بینی صحیح یال‌ها و ارتباطات آینده می‌توان افراد مختلفی را برای یک فردی که تازه عضو آن شبکه شده پیشنهاد کرد. همچنین می‌توان از اطلاعات به‌دست آمده از این شبکه روی شبکه‌های دیگر مثل اینستاگرام استفاده کرد. در یک کاربرد دیگر می‌توان از اطلاعات به‌دست آمده از ارتباطات افراد در این شبکه‌های مجازی برای ارتباطات دنیای واقعی و بالعکس استفاده کرد [۱۵].

برای مثال شبکه توییتر به‌صورت یک شبکه چندلایه در شکل زیر به‌تصویر کشیده شده است و ارتباطات بین افراد در لایه‌های مختلف وارد شده است.

^۱ Social Networks

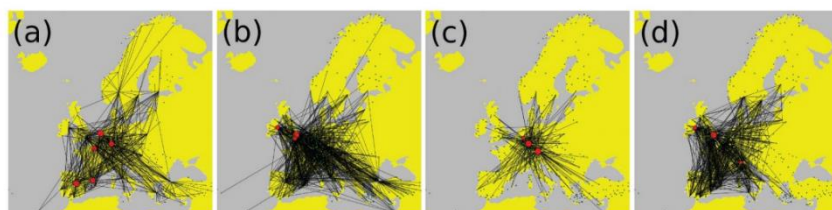
^۲ Twitter

^۳ Reply

^۴ Mention

^۵ Retweet

- شبکه حمل و نقل! در شبکه‌های حمل و نقل شهری، بین شهری یا حتی کشوری نیز می‌توان از این شبکه‌ها استفاده نمود و برای مثال ایستگاه‌های مختلف شهری را به عنوان گره‌ها همچنین نوع ارتباط بین این ایستگاه‌ها



شکل ۱-۶. تصویر شبکه چندلایه حمل و نقل اروپا [۶۴]

را به عنوان یال این شبکه در نظر گرفت و از آنجایی که این ارتباطات می‌توانند انواع مختلفی همچون تاکسی، اتوبوس تندرو، مترو و غیره داشته باشند هر کدام یک لایه جدا از این شبکه خواهد بود. هرگونه اعمال تغییر در هر یک از گره‌ها و یا یال‌های یک لایه ممکن است در لایه‌های دیگر تأثیرگذار باشد. به عنوان مثال خرابی در یکی از ایستگاه‌های مترو موجب افزایش بار ترافیکی روی لایه‌های دیگر آن ایستگاه همچون اتوبوس تندرو خواهد شد و احتمالاً در نتیجه موجب ایجاد لینک‌های جدید در لایه‌های دیگر خواهد شد. لذا مدل کردن این ارتباطات و پیش‌بینی برخی لینک‌ها می‌تواند در مدیریت بار ترافیکی شهری و جلوگیری از مشکلات احتمالی موثر واقع شود.

- پزشکی، دارویی و بیوانفورماتیک^۲

برای مثال در حوزه پزشکی می‌توان بازه‌های زمانی مختلف یک بیماری را به عنوان لایه‌های مختلف در نظر گرفت و رفتار بیماری را در لایه‌های مختلف مورد ارزیابی قرار داد و یا در حوزه دارو و کشف واکسن با تقسیم بندی نوع ارتباطات پروتئین‌ها یا نوع واکنش آن‌ها به یک نوع خاص از بیماری در لایه‌های مختلف نسبت به پیش‌بینی ارتباطات احتمالی بین پروتئین‌های مختلف اقدام کرد. همانطور که قبل‌تر نیز اشاره شد این پیش‌بینی‌ها در کاهش هزینه و زمان تولید داروها می‌تواند نقش بسیار مفیدی داشته باشد. در زیر یک مثال عملی از این کاربرد آورده شده است به گونه‌ای که در این شبکه برای درمان بیماری آسم با داشتن اطلاعاتی از سه لایه محیطی، علائم بیماری و شبکه تعاملات پروتئین-پروتئین^۳ و در نهایت با پیش‌بینی ارتباطات بین آن‌ها می‌توان داروی مختص و مورد نیاز را برای هر بیمار تجویز کرد [۱۶].

¹ Transportation Network

² Bioinformatic

³ Protein-Protein Interaction

۵-۱- چالش‌ها

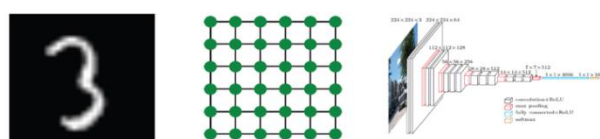
به دست آوردن بازنمایی برداری^۱ از هر گره گراف ذاتاً کاری است دشوار و چالش‌های بسیاری را ایجاد می‌کند که در زیر به برخی از آن‌ها اشاره شده است:

(الف) **انتخاب ویژگی:** یک بازنمایی برداری از گره خوب، باید ساختار گراف و اتصالات بین گره‌ها را حفظ کند. اولین چالش، انتخاب ویژگی‌هایی از گراف است که تعبیه باید آن‌ها را حفظ کند. با توجه به تعدد معیارهای مسافت^۲، خصوصیات و ویژگی‌های متعدد تعریف شده برای گراف‌ها، این انتخاب می‌تواند دشوار باشد، همچنین عملکرد و کارایی^۳ هم ممکن است به زمینه کاربرد تعبیه و با توجه به هدفی که از آن استفاده می‌کنیم بستگی داشته باشد.

(ب) **مقیاس پذیری:** بیش‌تر شبکه‌های واقعی شبکه‌هایی بزرگ هستند و حاوی هزاران یا میلیون‌ها گره و یال می‌باشند. روش‌های تعبیه باید مقیاس پذیر باشند و قادر به پردازش گراف‌های بزرگ باشند. تعریف یک مدل مقیاس پذیر می‌تواند چالش برانگیز باشد.

(پ) **ابعاد تعبیه:** یافتن ابعاد بهینه بازنمایی می‌تواند سخت باشد. به عنوان مثال، ابعاد بیش‌تر ممکن است دقت در بازسازی^۴ را افزایش دهد اما از نظر زمان و فضا از پیچیدگی بسیار بالایی برخوردار باشد. بسته به رویکرد ممکن است انتخاب ابعاد متفاوت باشد: به عنوان مثال، تعداد ابعاد کم‌تر ممکن است منجر به دقت بهتر در پیش‌بینی یال^۵ شود چرا که مدل حاصل فقط اتصالات محلی بین گره‌ها را به دست خواهد آورد [۱]. از طرف دیگر اکثر ابزارها و الگوریتم‌های طراحی شده برای یادگیری عمیق تنها روی دنباله‌ها، شبکه (grid) ها و

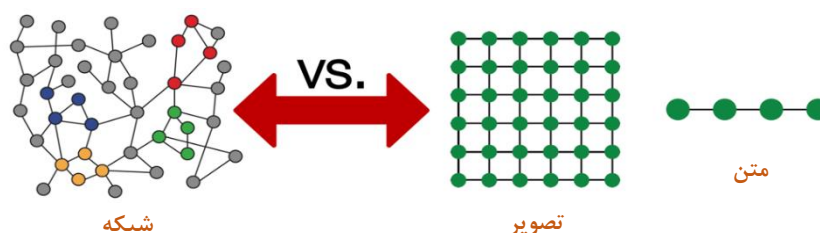
■ روش CNN برای تصاویر / (grids) هایی با سایز ثابت



■ روش RNN برای متن/صوت/ داده‌های دنباله‌ای (sequences)



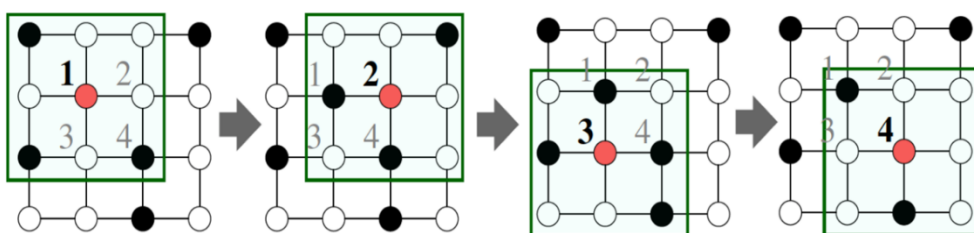
- ¹ Vector representation
- ² Distance metrics
- ³ Performance
- ⁴ Scalability
- ⁵ Dimensionality of embedding
- ⁶ Reconstruction
- ⁷ Link prediction



شکل ۷-۱. مقایسه ساختار داده‌های تصویری، متن و شبکه‌های دنیای واقعی [۶]

لتیس‌ها کاربرد دارند.

همانطور که در شکل بالا مشخص است، یک تصویر را می‌توان به عنوان یک *grid* در نظر گرفت. بنابراین روش *CNN* می‌تواند روی آن حرکت کند چرا که تنها با یک لتیس (*lattice*) بسیار مرسوم طرف هستیم. و یا در قسمت دوم شکل ۱-۹، یک داده صوتی را می‌بینیم، آنچه که در دنیای پردازش زبان طبیعی (*NLP*) شاهد هستیم، چه صوت و چه متن را می‌توانیم به عنوان داده‌های متوالی یا دنباله‌ای در نظر بگیریم. در واقع در اینجا



شکل ۱-۸. حرکت *CNN* روی داده‌های تصویر

هم با یک گراف زنجیره‌ای^۳ طرف هستیم. اما آنچه که در شبکه‌ها وجود دارد بسیار متفاوت‌تر است. شبکه‌ها در دنیای واقعی چه از لحاظ پیچیدگی و چه از لحاظ شکل و ساختار هیچ‌گونه شباهتی به این نوع داده‌ها (عکس، متن یا صوت) ندارند (قسمت سوم شکل ۱-۹).

(ت) **عدم وجود ترتیب خاص برای گره‌ها در گراف:** در گراف‌ها و شبکه‌ها ما نمی‌توانیم یک شماره گذاری ثابت روی گره‌ها داشته باشیم چرا که گره‌ها مرتب نیستند، درضمن هیچ گره‌ای را نمی‌توانیم به عنوان گره مبدا در نظر بگیریم و شماره گذاری را از آن شروع کنیم، بنابراین بی‌شمار شماره گذاری متفاوت برای گره‌های آن می‌توان متصور شد. بنابراین از الگوریتم‌هایی همچون *CNN* که بدلیل مرتب بودن گره‌های تصویر و ثابت بودن جایگیری آن‌ها مورد استفاده قرار می‌گیرد در این‌جا نمی‌توان استفاده کرد چراکه با شماره گذاری‌های متفاوت برای گره‌ها شاهد تعبیه‌های متفاوت خواهیم بود.

(ث) **پویایی شبکه‌ها و چند حالتی بودن گره‌ها:** مسئله دیگری که پیچیدگی شبکه‌ها را بیش‌تر می‌کند بحث پویایی شبکه‌ها می‌باشد، چراکه شبکه‌ها پیوسته در حال تغییرند و در طول زمان گره‌ها و یال‌های جدید افزوده و یا حذف می‌شوند، همچنین گره‌ها می‌توانند از انواع گوناگونی باشند، بطور مثال آنچه که در دنیای دادگان

^۱ lattices

^۲ Natural language processing

^۳ Chain graph

^۴ Multi-modal

نویسندگان مقالات داریم، گره‌ها می‌توانند مقالات، نویسندگان، موضوع و غیره باشد که می‌بینیم جنس گره‌ها باهم متفاوت‌اند.

(ج) پیچیدگی خاص تعبیه روی شبکه‌های چندلایه: همانگونه که اشاره شد شبکه‌های چندلایه خود به دلیل متفاوت بودن نوع گره‌ها یا یال‌ها در هر لایه دارای پیچیدگی‌های خاص می‌باشد و هرگونه تعبیه‌ای این تفاوت نوع یال‌ها یا گره‌ها را باید مورد توجه قرار دهد. از طرفی این سوال پیش می‌آید که در این شبکه‌ها لایه چه ارتباطی با یکدیگر دارند. لایه‌ها چه میزان با یکدیگر متفاوت یا شبیه هستند و اگر شباهت دارند معیار شباهت باید چه معیاری باشد. آیا تعبیه‌های هر لایه از لایه دیگر متفاوت خواهد بود یا تعبیه لایه‌ها روی همدیگر تأثیرگذار خواهند بود.

موارد ذکر شده ازجمله چالش‌هایی است که کار تعبیه گراف و به تبع آن کاربردهایی همچون پیش‌بینی یال روی شبکه‌های چندلایه را بیش از پیش با چالش مواجه خواهد کرد.

۲- پیشینه تحقیق

در یک دهه گذشته تحقیقات زیادی در زمینه تعبیه گراف با تمرکز بر طراحی الگوریتم‌های جدید تعبیه انجام شده است. محققان در تازه‌ترین تلاش‌های خود، الگوریتم‌های تعبیه مقیاس‌پذیر را به سمتی سوق داده‌اند که می‌توانند بر روی گراف‌هایی با میلیون‌ها گره و یال اعمال شوند. ما در ادامه، پیشینه‌ای از سیر پیشرفت تحقیقات در این زمینه را ارائه می‌دهیم، سپس به بررسی یک طبقه‌بندی از تکنیک‌های تعبیه گراف شامل، (الف) روش‌های تجزیه ماتریس، (ب) تکنیک‌های قدم‌زنی تصادفی و (پ) یادگیری عمیق می‌پردازیم. سپس به جهت‌اینکه در فصل‌های بعدی روش تعبیه جدیدی پیشنهاد شده و روی یکی از کاربردهای تعبیه همچون پیش‌بینی یال مورد ارزیابی قرار می‌گیرد، در ادامه به روش‌های پیش‌بینی یال و تاریخچه آن نیز پرداخته شده است.

۲-۱- تعبیه گراف، زمینه تحقیق و تکامل

در اوایل دهه ۲۰۰۰ میلادی، محققان، الگوریتم تعبیه گراف را به‌عنوان بخشی از تکنیک‌های کاهش بعد^۱ توسعه دادند. آن‌ها یک گراف شباهت را برای مجموعه‌ای از n نقطه D -بعدی براساس همسایگی ایجاد می‌کنند و سپس گره‌های گراف را در یک فضای برداری d -بعدی تعبیه می‌کنند، به گونه‌ای که $d \ll D$ باشد. ایده این تعبیه این بود که گره‌های متصل به هم را در فضای برداری (فضای تعبیه)، نزدیک به هم قرار دهند. الگوریتم‌های بردار ویژه لاپلاسیان^۲ و تعبیه خطی محلی^۳ [۱۷] نمونه‌هایی از الگوریتم‌های مبتنی بر این روش بودند. البته مقیاس‌پذیری مشکل اصلی این روش‌ها است به گونه‌ای که پیچیدگی زمانی آن‌ها $O(|V|^2)$ می‌باشد.

از سال ۲۰۱۰، تحقیقات در مورد تعبیه گراف به دلیل تنگ‌بودن شبکه‌های دنیای واقعی، به سمت ایجاد تکنیک‌های تعبیه گراف مقیاس‌پذیر سوق پیدا کرد. روش‌هایی همچون شبکه‌های کانولوشنی گرافی (GCN)^۴، رویکردهای مقیاس‌پذیر جدیدی هستند که پیچیدگی زمانی آن‌ها $O(|V|)$ می‌باشد [۱].

^۱ Dimensionality Reduction

^۲ Laplacian Eigenmaps (LE)

^۳ Locally Linear Embedding (LLE)

^۴ Sparse

^۵ Graph Convolutional Networks (GCN)

۲-۲- دسته‌بندی انواع روش‌های تعبیه گراف

در سال‌های اخیر شاهد تحقیقات گسترده‌ای در مورد تعبیه گره‌ها هستیم که منجر به تنوع پیچیده‌ای از نمادها و مدل‌های مختلف شده‌است. بنابراین، قبل از بحث در مورد تکنیک‌های مختلف، ابتدا یک چارچوب رمزگذاری-رمزگشایی^۱ یکپارچه را ارائه می‌کنیم. میتوان روش‌های مختلف تعبیه را حول دو تابع نگاشت کلیدی سازماندهی کرد: تابع رمزگذار، که هر گره را به یک بردار کم‌بعدتر، نگاشت یا تعبیه می‌کند و تابع رمزگشا، که اطلاعات ساختاری درباره گراف را از تعبیه‌های آموخته‌شده رمزگشایی می‌کند (شکل ۱-۲).

تعریف ۱-۲- رمزگذار (ENC): تابعی است که

$$(1,2)$$

$$ENC : V \rightarrow \mathbb{R}^d$$

گره‌ها را به تعبیه‌های برداری $Y_i \in \mathbb{R}^d$ (که در آن Y_i تعبیه مربوط به گره $v_i \in V$ است) نگاشت می‌کند.

تعریف ۲-۲- رمزگشا (DEC): تابعی است که مجموعه‌ای از تعبیه‌های گره‌ها را گرفته و آمار گرافی و اطلاعات مشخص شده توسط کاربر را از این تعبیه‌ها رمزگشایی می‌کند. به عنوان مثال، رمزگشا ممکن است وجود یال بین گره‌ها را با توجه به تعبیه‌های آن‌ها پیش‌بینی کند [۱۸]، یا ممکن است اجتماع یک گره را پیش‌بینی کند [۲۰] [۱۹]. در اصل، رمزگشاهای زیادی می‌توان استفاده کرد. با این حال، عمدتاً از یک رمزگشای زوجی استفاده می‌کنند:

$$(2,2)$$

$$DEC : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}^+$$

وقتی از رمزگشای زوجی روی یک زوج تعبیه (Y_i, Y_j) استفاده کنیم، بازسازی شباهت بین یک زوج گره (v_i, v_j) در گراف اصلی به دست می‌آید که هدف، بهینه‌سازی توابع نگاشت رمزگذار و رمزگشا برای به حداقل رساندن خطا در این بازسازی خواهد بود، به گونه ای که:

$$DEC(ENC(v_i), ENC(v_j)) = DEC(Y_i, Y_j) \approx S_G(v_i, v_j) \quad (3,2)$$

¹ Encoder-Decoder

² Reconstruction Error

که در آن S_G ماتریس شباهت بین گره‌ها در گراف G می‌باشد که توسط کاربر تعریف شده است. به عبارت دیگر، ما می‌خواهیم مدل رمزگذار-رمزگشای خود را طوری بهینه کنیم تا بتوانیم شباهت بین جفت گره‌ها را در گراف اصلی (S_G) از روی تعبیه‌های (Y_i, Y_j) رمزگشایی کنیم. به عنوان مثال، می‌توان شباهت بین گره‌ها در گراف اصلی را متناظر با ماتریس مجاورت آن (A) در نظر گرفت:

$$S_G(v_i, v_j) = A_{ij} \quad (4,2)$$

طوری‌که در صورت همسایه بودن دو گره، عدد یک و در غیر این صورت عدد صفر برای آن جفت گره در ماتریس شباهت (S_G) لحاظ کرد [۲۰]. یا ممکن است معیار تشابه بین دو گره را با توجه به احتمال هم‌رخداد بودن آن دو گره در یک قدم‌زنی تصادفی با طول ثابت تعریف کرد، روشی که در مدل‌های مبتنی بر قدم‌زنی تصادفی شاهد آن هستیم [۱۴]. در عمل، در بیش‌تر روش‌ها هدف بازسازی، به حداقل رساندن خطای $\mathcal{L}$ بر روی مجموعه‌ای از جفت‌گره‌های آموزشی (D) می‌باشد:

$$\mathcal{L} = \sum_{v_i, v_j \in D} \ell(DEC(Y_i, Y_j), S_G(v_i, v_j)) \quad (5,2)$$

که در آن ℓ یک تابع زیان است که توسط کاربر تعیین می‌شود و درواقع میزان اختلاف بین مقادیر شباهت به‌دست‌آمده از رمزگشایی تعبیه گره‌ها ($DEC(Y_i, Y_j)$) و شباهت بین گره‌ها در گراف اصلی ($S_G(v_i, v_j)$) را اندازه‌گیری می‌کند [۲۰]. پس از بهینه‌سازی سیستم رمزگذار-رمزگشا، می‌توانیم از رمزگذار آموزش‌دیده برای تولید تعبیه گره‌ها استفاده کنیم که از آن به‌عنوان ورودی برای کارهای یادگیری ماشین استفاده می‌شود. به‌عنوان مثال، می‌توان تعبیه‌های آموخته‌شده را به یک دسته‌بند رگرسیون لجستیک داد تا اجتماع متعلق به یک گره را پیش‌بینی کند، یا می‌توان از فاصله بین تعبیه‌ها برای توصیف پیوندهای دوستی در یک شبکه اجتماعی استفاده کرد [۱۴].

حال با علم به رویکرد رمزگذار-رمزگشا، در ادامه یک دسته‌بندی از رویکردهای تعبیه ارائه می‌کنیم و روش‌های تعبیه را به سه دسته کلی طبقه‌بندی می‌کنیم: (۱). مبتنی بر تجزیه ماتریس، (۲). مبتنی بر قدم‌زنی تصادفی و (۳). مبتنی بر یادگیری عمیق. سپس ویژگی‌های هر یک از این دسته‌ها را توضیح می‌دهیم. همچنین نمونه‌هایی از هر گروه در جدول ۱-۲، ارائه شده است.

¹ Co-Occurrence

² Loss Function

³ Logistic Regression Classifier

۳-۲- روش‌های مبتنی بر تجزیه ماتریس

روش‌های اولیه یادگیری تعبیه‌گره‌ها تا حد زیادی روی رویکردهای مبتنی بر تجزیه ماتریس متمرکز شده‌اند، که به‌طور مستقیم از تکنیک‌های کلاسیک برای کاهش بُعد الهام گرفته می‌شوند [۲۰]. الگوریتم‌های مبتنی بر تجزیه، اتصالات بین گره‌ها را به‌صورت ماتریس نشان می‌دهند و این ماتریس را برای به‌دست آوردن تعبیه ماتریس حاصل، تجزیه (فاکتورگیری) می‌کنند. ماتریس مورد استفاده برای نشان دادن تشابه بین گره‌ها می‌تواند شامل ماتریس مجاورت، ماتریس لاپلاسی^۱، ماتریس احتمال انتقال گره‌آ و یا ماتریس تشابه $Katz$ باشد [۱].

۳-۲-۱- بردار ویژه لاپلاسی

یکی از اولین و شناخته‌شده‌ترین نمونه‌ها، تکنیک بردار ویژه لاپلاسی (LE) [۲۱] است که می‌توانیم در چارچوب رمزگذار-رمزگشا، تابع رمزگذار و رمزگشای آن را بدین صورت تعریف کنیم:

تابع رمزگذار این روش یک ماتریس لاپلاسی و همچنین تابع رمزگشای آن بدین صورت می‌باشد:

$$(6,2)$$

$$DEC(Y_i, Y_j) = \|Y_i - Y_j\|_2^2$$

همچنین تابع زیان، جفت گره‌ها را با توجه به شباهت آن‌ها در گراف اصلی (مثل A_{ij}) وزن‌دهی می‌کند [۲۱]:

$$(7,2)$$

$$\mathcal{L} = \frac{1}{2} \sum_{(v_i, v_j)} DEC(Y_i, Y_j) \cdot S_G(v_i, v_j)$$

در این روش هدف این است که وقتی دو گره در گراف اصلی شبیه هم هستند (برفرض مثال وقتی وزن A_{ij} زیاد است)، تعبیه دو گره را نزدیک هم نگه دارد [۱]. برای این منظور تابع زیان را به حداقل می‌رسانند.

بردار ویژه لاپلاسی از یک تابع زیان^۴ درجه دو استفاده می‌کند. بنابراین تابع هدف بیش‌تر بر حفظ تفاوت بین گره‌ها تأکید می‌کند تا شباهتشان. این مسئله ممکن است موجب عدم حفظ ساختار محلی شود، چراکه اگر به

¹ Laplacian matrix

² Node transition probability matrix

³ Katz similarity matrix

⁴ Loss Function

تابع هدف توجه کنیم شباهت بین یال‌ها $(S_G(v_i, v_j))$ از درجه یک است ولی تفاوت بین یال‌ها از درجه دو می‌باشد $\langle \|Y_i - Y_j\|_2^2 \rangle$.

۲-۳-۲- روش‌های تجزیه ماتریس مبتنی بر ضرب داخلی

به دنبال روش بردار ویژه لاپلاسی، اخیراً روش‌های تعبیه متعددی تولید شده‌است که در آن تابع رمزگشا از حاصل ضرب داخلی تعبیه‌ها بدست می‌آید:

$$DEC(Y_i, Y_j) = Y_i X_{\mathcal{A}, 2}$$

به گونه‌ای که استحکام رابطه بین دو گره، متناسب با ضرب داخلی تعبیه‌های آن‌ها می‌باشد. الگوریتم‌های فاکتورگیری گراف (GF) [۲۲]، $GraRep$ [۲۳] و $HOPE$ [۲۴] همه در این کلاس قرار می‌گیرند و هر سه روش از یک رمزگشا با ضرب داخلی استفاده می‌کنند، همچنین تابع زیان آن‌ها یک تابع خطای میانگین مربعات (MSE) می‌باشد:

$$\mathcal{L} = \sum_{(v_i, v_j)} \|DEC(Y_i, Y_j) - S_G(v_i, v_j)\|_2^2 \quad (9, 2)$$

در واقع تفاوت اصلی روش‌های یاد شده $(GF, GraRep, HOPE)$ ، در معیار شباهت استفاده شده روی گراف اصلی می‌باشد $(S_G(v_i, v_j))$.

در روش فاکتورگیری گراف (GF) ، معیار شباهت را ماتریس مجاورت تعریف می‌کنند:

$$(10, 2)$$

$$S_G(v_i, v_j) = A_{ij}$$

روش $GraRep$ به منظور اخذ همسایگی‌های مرتبه بالاتر معیار شباهت دو گره در گراف اصلی را توان‌های بالاتر ماتریس مجاورت در نظر می‌گیرد، بر فرض مثال:

$$S_G(v_i, v_j) = A_{ij}^2 \quad (11, 2)$$

¹ Graph Factorization

² Mean-Squared-Error

روش HOPE برای تعیین شباهت بین دو گره از معیارهای شباهتی همچون هم‌پوشانی همسایگی جاکارد^۱ استفاده می‌کند. این روش بین «شباهت مرتبه اول»، که در آن S_G وجود یال مستقیم بین گره‌ها را در نظر می‌گیرد و «شباهت‌های مرتبه بالاتر» که S_G مفهوم کلی‌تری از همسایگی بین گره‌ها را برمی‌گرداند، تعادل برقرار می‌کند. دلیل اینکه این روش‌ها به‌عنوان روش‌های فاکتورگیری ماتریس شناخته می‌شوند این است که درواقع ماتریس شباهت (A) را به دو ماتریس Y و Y^T تجزیه یا فاکتورگیری می‌کنیم. هدف از این روش‌ها صرفاً یادگیری تعبیه برای هر گره به‌گونه‌ای است که ضرب داخلی بین بردارهای تعبیه آموخته‌شده، برخی از معیارهای شباهت بین جفت گره‌ها در گراف اصلی را تقریب بزند.

۴-۲- روش‌های مبتنی بر قدم‌زنی تصادفی

بسیاری از روش‌های موفق اخیر تعبیه گره را از طریق اطلاعات به‌دست‌آمده از قدم‌زنی تصادفی یاد می‌گیرند. نوآوری اصلی این روش‌ها بهینه‌سازی تعبیه گره‌ها به‌گونه‌ای است که گره‌هایی که در یک قدم‌زنی تصادفی کوتاه برروی گراف، هم‌رخداد می‌شوند، تعبیه مشابهی داشته باشند (شکل ۲-۲). بنابراین، به‌جای استفاده از معیارهای شباهتی که قبلاً بحث شد (منظور معیارهای شباهت در گراف اصلی است)، روش‌های قدم‌زنی تصادفی از یک معیار شباهت انعطاف‌پذیر و همچنین تصادفی استفاده می‌کنند [۶].

از قدم‌زنی تصادفی برای تقریب بسیاری از ویژگی‌های گراف از جمله مرکزیت گره^۲ و شباهت استفاده شده‌است. حالتی را تصور کنید که گراف پاره‌ای مشاهده پذیر است، به این معنا که نمی‌توان به کل گره‌ها و یال‌های گراف دسترسی داشت یا حالتی که گراف مورد مطالعه بسیار بزرگ باشد، اینجاست که کارآمدی این روش‌ها مشخص می‌شود. روش‌های مختلفی برای تعبیه با استفاده از قدم‌زنی تصادفی بر روی گراف‌ها جهت به‌دست‌آوردن بازنمایی گره ارائه شده‌است که روش‌های *DeepWalk* و *Node2Vec* دو مثال بارز از این دسته هستند [۱۱] [۱۴].

۴-۲-۱ مدل *DeepWalk*

DeepWalk یک روش دو مرحله‌ای است. این روش در مرحله اول، شبکه را با قدم‌زنی تصادفی طی می‌کند تا ساختارهای محلی را از طریق روابط همسایگی استنباط کند. در مرحله دوم، از الگوریتمی به‌نام *SkipGram* برای یادگیری و بهینه‌سازی تعبیه‌های ساختارهای به‌دست‌آمده از مرحله قبل استفاده می‌کند. مرحله ۱- کشف ساختار محلی: هدف از این مرحله، کشف همسایگی‌های هر گره در شبکه است. برای انجام این کار، تعداد ثابت (k) قدم‌زنی تصادفی را که از هر گره به‌صورت مجزا شروع شده، ایجاد می‌شود. طول هر قدم‌زنی

¹ Jaccard Neighborhood Overlaps

² Node Centrality

(l) از قبل تعیین شده است. بنابراین، هنگامی که این مرحله به پایان رسید، k دنباله گره^۱ به طول l به دست می‌آید. اصل تولید قدم‌زنی تصادفی این فرض است: گره‌های مجاور، مشابه هستند و باید تعبیه‌های مشابهی نیز داشته باشند. براساس این فرض، گره‌هایی را که هم‌زمان در یک مسیر وجود دارند^۲ به‌عنوان گره‌های مشابه تعبیر و پذیرفته می‌شوند. این بدان معنی است که میزان فراوانی هم‌زمانی دو گره در قدم‌زنی تصادفی، نشانگر شباهت آن دو گره است. درواقع این یک فرض معقول است چراکه یال‌ها، بیش‌تر نشان‌دهنده گره‌های مشابه یا گره‌های درحال تعامل در یک شبکه هستند. مرحله ۲- *SkipGram*: الگوریتم *SkipGram* یک تکنیک محبوب برای یادگیری تعبیه کلمات است. با توجه به کوله کلمات^۳ و اندازه پنجره^۴ هدف *SkipGram* این است که شباهت تعبیه کلماتی را که در همان پنجره اتفاق می‌افتند، به حداکثر برساند. این ایده بر این فرض استوار است که: کلماتی که در یک پنجره وجود دارند، معمولاً معانی نزدیکی دارند. بنابراین تعبیه آن‌ها نیز باید نزدیک به یکدیگر باشد. برای تطبیق *SkipGram* با شبکه‌ها، باید مفهوم پنجره (یا زمینه) مطابق با دنیای شبکه‌ها را تعیین کرد. اینجاست که قدم‌زنی‌های تصادفی وارد عمل می‌شوند.

هدف از مدل سازی زبانی، تخمین احتمال ظاهر شدن یک توالی خاص از کلمات، در یک کوله کلمات است. به‌طور رسمی‌تر، با توجه به دنباله‌ای از کلمات:

(۱۲،۲)

$$W_1^n = (w_0, w_1, \dots, w_n)$$

که در آن w_i کلمات موجود در کوله کلمات و W_1^n دنباله‌ای از این کلمات هستند، هدف، ماکزیمم کردن احتمال زیر (احتمال مشاهده کلمه n ام با وجود مشاهده n کلمه قبلی) برای تمامی کلمات موجود کوله کلمات می‌باشد:

(۱۳،۲)

$$Pr(w_n | w_0, \dots, w_{n-1})$$

حال در این روش، یک تعمیم از مدل‌سازی زبانی، جهت اکتشاف گراف از طریق جریانی از قدم‌زنی‌های تصادفی کوتاه ارائه می‌شود. قدم‌زنی‌ها در یک شبکه را می‌توان معادل جملات و عبارات کوتاه در یک متن بلند در نظر گرفت. با این معادل‌سازی باید گفت در اینجا نیز هدف، این است که احتمال مشاهده رأس v_i را با توجه به تمام رؤس قبلی که تاکنون در مسیر قدم‌زنی تصادفی بازدید شده‌است، تخمین زد.

¹ Degree Sequence

² Co-Occurrence Nodes

³ Bag of Words (BoW)

⁴ Window Size or Context Size

(۱۴,۲)

$$Pr(v_i | (v_0, \dots, v_{i-1}))$$

اما از آنجایی که هدف اصلی در اینجا یادگیری تعبیه‌هاست، بنابراین این معادله را باید به صورت زیر بازنویسی کرد:

(۱۵,۲)

$$Pr(v_i | (Y_0, \dots, Y_{i-1}))$$

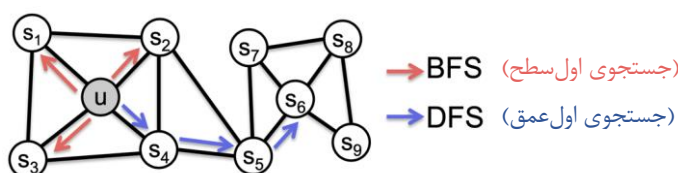
و این یعنی احتمال مشاهده رأس v_i ، به شرط مشاهده تعبیه تمام رئوس مشاهده شده قبلی در قدم‌زنی تصادفی. این هدف نیز دو اشکال اساسی دارد. اول اینکه در صورت طولانی شدن طول قدم‌زنی، قسمت شرطی معادله بالا بزرگ‌تر شده و عملاً محاسبه آن غیرممکن می‌شود. همچنین باتوجه به اینکه در فضای زبانی هم، کلمات قبل و بعد از یک کلمه خاص مهم هستند، با ایده گرفتن از آن در اینجا هم، همسایه‌های قبل و بعد از گره v_i را در محاسبات لحاظ می‌کنیم، بنابراین طول قدم‌زنی تصادفی به $2k+1$ افزایش می‌یابد. لذا برای حل این مشکل و رفع پیچیدگی زمانی یادشده، هدف جایگزین، ماکزیمم کردن تابع هدف زیر تعریف می‌شود:

$$\log Pr(v_{i-k}, \dots, v_{i-1}, v_{i+1}, \dots, v_{i+k} | Y_i) \quad (۱۶,۲)$$

برای یادگیری تعبیه‌ها از طریق *SkipGram* بردارهای d -بعدی، ابتدا به صورت تصادفی برای هر گره ایجاد می‌شود، سپس، مدل باتوجه به تابع هدف بالا، تعبیه‌ها را به گونه‌ای بهینه‌سازی می‌کند که احتمال مشاهده همسایه‌های یک گره (گره‌های هم‌رخداد در یک قدم‌زنی تصادفی) در صورت دیدار آن گره، افزایش یابد. این کار موجب می‌شود گره‌هایی که همسایگان مشابهی دارند، تعبیه‌های مشابهی نیز برای آن‌ها محاسبه شود.

۲-۴-۲- مدل Node2Vec

این مدل در اصل، ارتقاء یافته مدل *DeepWalk* با انجام برخی تغییرات در بخش قدم‌زنی تصادفی آن می‌باشد [۱۴]. در روش *DeepWalk*، فاکتور طول ثابت (k) اجازه نمی‌داد که قدم‌زن گره‌های دورتر را مشاهده و کشف کند، فرض کنید دو گره شبیه به هم وجود داشته باشند اما فاصله آن‌ها بیش‌تر از k باشد، در این صورت *DeepWalk* قادر به کشف این شباهت نخواهد بود. در واقع، تفاوت بسیار مهم *DeepWalk* و این روش در این است



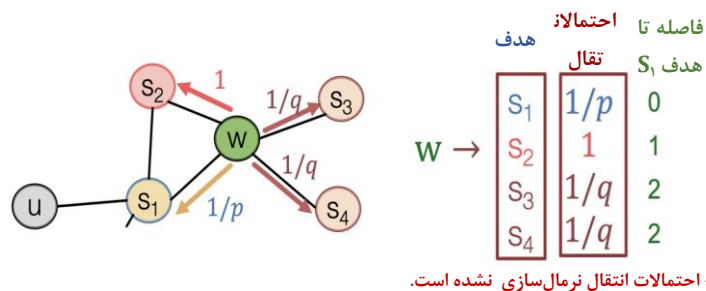
همسایه‌های کشف شده توسط جستجوی اول سطح: $\{s_1, s_2, s_3\} = N_{BFS}(u)$

همسایه‌های کشف شده توسط جستجوی اول عمق: $\{s_4, s_5, s_6\} = N_{DFS}(u)$

شکل ۲-۱. تفاوت بین قدم‌زنی‌های تصادفی براساس جستجوی اول سطح و جستجوی اول عمق [۶۳]

که این روش از قدم‌زنی تصادفی مغرضانه^۱ استفاده می‌کند، طوری‌که بین جستجوی گراف اول‌سطح (BFS) و اول‌عمق (DFS) یک تعادل و توازنی ایجاد کند، و از این رو تعبیه‌هایی با کیفیت بالاتر و حاوی اطلاعات مفیدتری نسبت به مدل سابق تولید می‌کند. داشتن رویکرد جستجوی اول‌عمق به ما این امکان را می‌دهد تا گره‌هایی با فاصله دورتر را کشف کنیم.

انتخاب تعادل مناسب بین رویکرد محلی (جستجوی اول‌سطح) و رویکرد سراسری (جستجوی اول‌عمق)، این روش را قادر می‌سازد تا ساختار اجتماع‌ها و همچنین تعادل ساختاری بین گره‌ها را بهتر حفظ کند.



شکل ۲-۲. نمایی از چگونگی ایجاد تعادل در قدم‌زنی تصادفی با استفاده از پارامترهای p و q در روش Node2Vec. با فرض اینکه قدم زن از S_1 به W منتقل شده است، برچسب یال‌ها، با احتمال انتخاب شدن آن گره به‌عنوان گره بعدی برای حرکت بعدی قدم‌زن متناسب هستند [۲۰].

به‌منظور حفظ تعادل بین جستجوی اول‌سطح و جستجوی اول‌عمق از دو پارامتر برای تولید و کشف همسایه‌ها استفاده می‌کنند: پارامتر p ، که پارامتر برگشت به گره قبلی است و پارامتر q ، نسبت جستجوی اول‌سطح به اول‌عمق می‌باشد. برای توضیح بیشتر و قابل لمس بودن این پارامترها در زیر شکلی آورده شده است تا با یک مثال ساده به توضیح این دو پارامتر بپردازیم. فرض کنیم که قدم زن از S_1 حرکت کرده و رسیده به W و الان در W قرار دارد، همانطور که می‌بینیم دیگر احتمال‌های انتقال باهم برابر نیستند. بنابراین با احتمال $1/p$ به گره قبلی (S_1) برمی‌گردد، با احتمال $1/q$ گره‌های جدید را کشف می‌کند (S_3, S_4) و با احتمال ۱ به S_2 می‌رود، احتمال رفتن به S_2 برابر ۱ شده است، چراکه با رفتن به این گره، فاصله از گره مبدأ یعنی S_1 تغییری نمی‌کند. درضمن احتمالات نوشته شده نرمال‌سازی نشده‌اند و باید تقسیم‌بر مجموع احتمالات شوند تا نرمال‌سازی انجام و مجموع احتمالات یک شود. حال پیداست که هرچه مقدار p کمتر باشد قدم‌زن تصادفی ما بیشتر به

¹ Biased-Random Walks

² Breadth First Search

³ Depth First Search

⁴ Community

⁵ Walker

⁶ Unnormalized

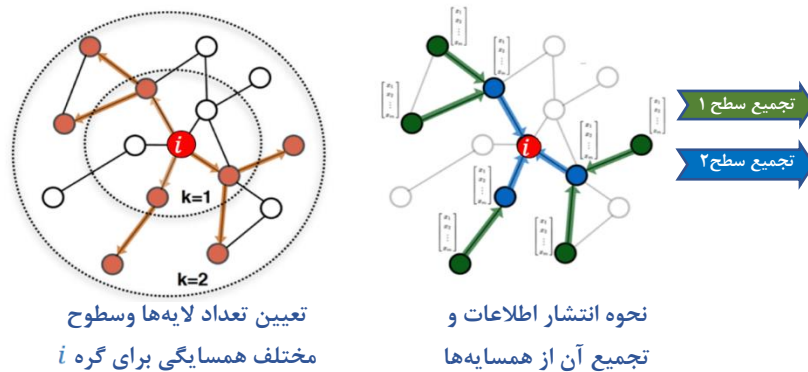
جستجوی اول سطح متمایل خواهد بود، از طرفی هرچه مقدار q کم‌تر باشد تمایل به جستجوی اول عمق بیش‌تر خواهد بود. عمده اشکال هر دو روش معرفی شده، یعنی *DeepWalk* و *Node2Vec* این است که، تعبیه گره را به‌طور تصادفی برای آموزش مدل‌ها آغاز می‌کنند. از آنجایی که تابع هدف آن‌ها غیر محدب^۱ است، این چنین شروع‌های تصادفی می‌تواند در بهینه‌های محلی^۲ گیر کند. حال برای رفع این محدودیت‌ها مدل‌های نوین‌تری ارائه شدند که در ادامه به آن می‌پردازیم.

۵-۲- روش‌های مبتنی بر یادگیری عمیق

رشد تحقیقات بر روی یادگیری عمیق باعث شده تا روش‌های مبتنی بر شبکه‌های عصبی عمیق بر گراف‌ها اعمال شوند [۹] [۱۰] [۲۵]. این روش‌ها، به‌دلیل توانایی آن‌ها در مدل‌سازی ساختار غیرخطی از داده‌ها، به منظور کاهش بُعد مورد استفاده قرار گرفته‌اند، از طرفی روش‌های مبتنی بر قدم‌زنی تصادفی این کاستی را داشتند که ویژگی‌های خود گره‌ها تأثیر داده نمی‌شد، هر گره ممکن است برای خود ویژگی‌هایی داشته باشد که باید در فرآیند تعبیه دخیل شوند. لذا این موارد باعث شدند تا محققان به تعبیه‌های قدرتمندتری روی بیاورند.

^۱ Non-Convex

^۲ Local Optima



شکل ۲-۳. نحوه تجمیع پیام‌های رسیده از همسایه‌های گره i [۲۰]

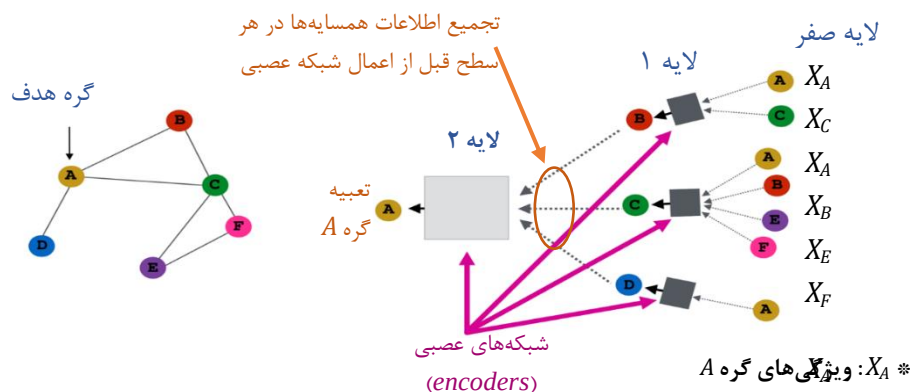
۱-۵-۲- شبکه‌های کانولوشن گرافی^۱ (GCN)

شبکه‌های کانولوشن گرافی (GCN) [۱۹] با تعیین یک عملگر حلقوی^۲ بر روی گراف، به صورت تکراری تعبیه‌های همسایه‌ها را برای هر گره تجمیع می‌کند و از یک تابع برای ترکیب تعبیه به دست آمده از همسایه‌ها و تعبیه خود آن گره که از تکرار قبلی موجود است، جهت ایجاد تعبیه جدید استفاده می‌کند. همانطور که در شکل نیز نشان داده شده است، مزیت این روش این است که هم، ساختار همسایگی گره‌ها را اخذ می‌کند و هم اینکه برای تعبیه یک گره، اطلاعات خود آن گره را نیز در نظر می‌گیرد. معمولاً در این روش، در انتخاب تعداد لایه‌ها (عمق) از ۶ لایه بالاتر نمی‌روند، چراکه هرچه عمیق‌تر شویم پیام‌های رسیده از همسایه‌ها اهمیتشان برای ما بیش‌تر می‌شود، در واقع باید ببینیم تا چه عمقی همسایه‌ها برای ما مهم هستند. از طرفی، باتوجه به اینکه شبکه‌های دنیای

¹ Graph Convolutional Networks

² Convolutional operator

³ Aggregate

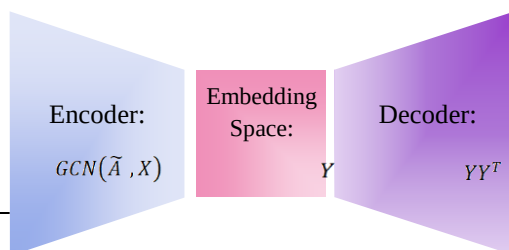


شکل ۲-۴. نمایی از نحوه رمزگذاری و تجمیع اطلاعات همسایه‌ها. جهت تولید تعبیه برای گره A ، این مدل اطلاعات همسایه‌های محلی A (یعنی B ، C و D) را جمع‌آوری می‌کند، سپس اطلاعات دریافت‌شده از این همسایگان را به همراه اطلاعات موجود در خود A تجمیع و به رمزگذار تحویل می‌دهد. در این مثال حالت دولایه از این مدل نشان داده شده است (یعنی اطلاعات همسایگی تا دو گام دورتر از گره مبدأ جمع‌آوری می‌شوند)، اما در اصل این روش می‌تواند از یک عمق دلخواه برخوردار باشند. در "عمق" یا "لایه" صفر اطلاعات اولیه همان ویژگی‌های خود گره‌ها هستند [۱۴].

واقعی دارای میانگین فاصله^۱ و یا قطر شبکه^۲ کوچک‌تر از ۶ هستند^۳، بنابراین اگر تا این عمق پیش برویم تقریباً کل گره‌های موجود در شبکه را مشاهده کرده‌ایم و دیگر نیازی نیست بیش‌تر از این بر عمق بی‌افزاییم (معمولاً ۲ تا ۴ سطح بسته به نوع کاربرد برای ما کافی است). همچنین در این مدل از تابع زیان اختلاف آن‌تروپی دودویی جهت بهینه‌سازی مدل^۴ استفاده می‌شود.

۲-۵-۲ مدل $Graph\ Auto\ Encoder\ (GAE)$

$Autoencoder$ یک شبکه عصبی است که از یک رمزگذار و رمزگشا تشکیل شده است. این مدل چارچوبی را برای آموزش روی گراف‌ها به صورت بدون ناظر^۵ ارائه می‌کند. رمزگذار داده‌های ورودی را به یک چکیده تبدیل می‌کند (داده‌های ورودی را فشرده‌سازی کرده و به یک فضای با بعد کم‌تر می‌برد)، این در حالی است که رمزگشا داده‌های خروجی رمزگذار را گرفته و سعی در بازسازی آن داده‌ها دارد. درواقع $Graph\ Auto\ Encoder$



شکل ۲-۵. چارچوب کلی روش GAE

¹ Average Distance

² Network Diameter

³ Six-Degrees of Separation Phenomena

⁴ Model Optimization

⁵ Unsupervised Learning

یکی از روش‌های تعبیه گراف است و بر پایه شبکه‌های عصبی طراحی شده که داده‌های ورودی آن، اطلاعات یک گراف می‌باشد [۲۶]. این روش نیز به دلیل پیچیدگی زمانی کم و توانایی بسیار خوب در کاهش بُعد، توجه بسیاری را در حوزه تعبیه گراف به خود جلب کرده است.

تابع رمزگذار در این مدل یک شبکه عصبی است که عمدتاً از روش تعبیه GCN در آن استفاده می‌شود. این تعبیه‌کننده (GCN) برای اینکه بتواند تعبیه گره‌ها را تولید کند به دو ماتریس نیاز دارد. (الف) ماتریس همبندی (A) گراف، (ب) ماتریس ویژگی گره‌ها (X). بنابراین خواهیم داشت:

$$\begin{aligned} (17,2) \\ Y = GCN(\tilde{A}, X) = ReLU(\tilde{A} X W_0) \\ (18,2) \\ A \leftarrow A + I, \quad \tilde{A} = D^{-1/2} A D^{-1/2} \end{aligned}$$

که در آن D ماتریس قطری تشکیل یافته از درجات گره‌ها، I ماتریس همانی، Y خروجی تابع رمزگذار (GCN) و همان ماتریس گره‌های تعبیه شده در فضای جدید می‌باشد. برای تولید این تعبیه‌ها، GCN ماتریس همبندی و ماتریس ویژگی گره‌ها را به عنوان ورودی گرفته، ابتدا یک فرم نرمالایز شده از ماتریس همبندی ($\tilde{A}$) را ایجاد کرده و سپس با ضرب ماتریس حاصله در ماتریس ویژگی و همچنین ماتریس وزنی (ماتریس ضرایب لایه‌های GCN) و در آخر با اعمال یک تابع فعال ساز^۱ (ReLU)، ماتریس تعبیه گره‌ها را به دست می‌آورد. این GCN می‌تواند از چندین لایه تشکیل شود، به طور مثال در [۲۶] نویسنده از دو لایه GCN به عنوان رمزگذار استفاده کرده است. در صورت استفاده از دو لایه GCN خواهیم داشت:

$$Y^{l2} = GCN(\tilde{A}, Y^{l1}) = ReLU(\tilde{A} Y^{l1} W_1) = \tilde{A} ReLU(\tilde{A} X W_0) W_1$$

Y^{l1} تعبیه‌های تولید شده توسط لایه اول و Y^{l2} تعبیه‌های تولید شده توسط لایه دوم می‌باشد، همچنین Y^{l1} (تعبیه‌های لایه اول GCN) همان ماتریس ویژگی گره‌ها و ماتریس‌های W_1 و W_0 ماتریس ضرایبی است که مدل در هر لایه سعی در آموزش آن دارد.

¹ Activation Function

تابع رمزگشا: حال که تعبیه گره‌ها به دست آمد، مدل توسط تابع رمزگشا سعی در بازسازی^۱ این تعبیه‌ها دارد به گونه‌ای که پس از بازسازی، ماتریس حاصله به گراف اصلی هرچه نزدیک‌تر و شبیه‌تر باشد. درواقع مدل تلاش دارد خطای بازسازی^۲ را به حداقل برساند. یکی از متداول‌ترین روش‌ها برای بازسازی (تابع رمزگشا) استفاده از ضرب داخلی می‌باشد، به این ترتیب تابع رمزگشا با ضرب داخلی تعبیه گره‌های تولیدشده توسط رمزگذار، سپس با اعمال یک تابع فعال‌ساز مثل سیگموید^۳ احتمال رخداد یال بین هر جفت گره را بازسازی می‌کند. مدل با مقایسه احتمال یال‌های بدست‌آمده از تابع رمزگشا و سنجش صحیح یا ناصحیح بودن این پیش‌بینی‌ها با یال‌های واقعی موجود در گراف، تابع زیان را به روزرسانی کرده و سعی در بهبود تعبیه‌های تولیدشده و در نتیجه بهبود پیش‌بینی‌ها خواهد داشت. برای این منظور از تابع زیان اختلاف آنروپی دودویی لاجستیک^۴ استفاده می‌شود [۲۷].

$$\mathcal{L} = \text{Average} \{l_1, \dots, l_N\} \quad (20, 2)$$

$$l_n = -w_n [y_n \cdot \log \sigma(x_n) + (1 - y_n) \cdot \log(1 - \sigma(x_n))] \quad (21, 2)$$

در معادله بالا، N تعداد یال‌های موجود در فاز یادگیری است که در اختیار مدل قرار گرفته، w_n ماتریس ضرایبی است که مدل سعی در یادگیری آن دارد، y_n برچسب واقعی یال‌هاست که در صورت وجود در گراف مقدار ۱ و در غیر این صورت مقدار ۰ می‌گیرد، x_n مقدار پیش‌بینی شده توسط مدل برای آن یال و σ تابع فعال‌ساز سیگموید می‌باشد، با ضرب سیگموید، آن مقدار پیش‌بینی شده (خروجی تابع رمزگشا) تبدیل به یک احتمال می‌شود که در بازه صفر و یک قرار دارد. هرچه مقدار احتمال پیش‌بینی شده برای هر یال به ماتریس همبندی گراف اصلی نزدیک‌تر باشد، مقدار این تابع زیان کوچک‌تر شده و در نهایت منجر به پیش‌بینی‌های باکیفیت‌تری خواهد شد. همچنین نویسنده پیشنهاد می‌کند، در صورتی که ماتریس ویژگی‌گره‌ها در دسترس نباشد، با جایگزینی آن با ماتریس مشخصه^۵ (ماتریس قطری که همه درایه‌های قطر اصلی آن یک می‌باشد) وابستگی مدل به ماتریس ویژگی از بین می‌رود [۲۶].

¹ Reconstruction

² Reconstruction Error

³ Sigmoid Function (Logistic Function)

⁴ Binary Cross Entropy with Logits Loss

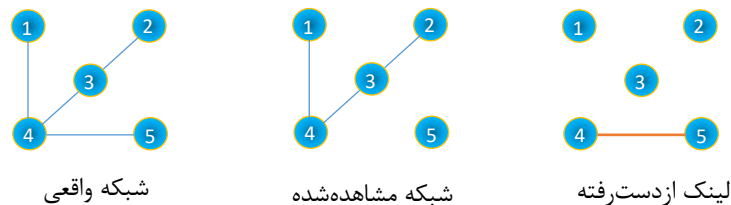
⁵ Identity Matrix

۶-۲- مسائل پیش‌بینی یال به‌عنوان یکی از کاربردهای تعبیه

حال که به بررسی برخی از مهم‌ترین و پرکاربردترین روش‌های تعبیه گراف پرداختیم، به دلیل استفاده از پیش‌بینی یال به‌عنوان یکی از کاربردهای تعبیه گراف، جهت سنجش و ارزیابی تعبیه پیشنهادی در فصول آتی، بهتر است نگاهی به مسائل و روش‌های پیش‌بینی یال غیرتعبیه‌ای هم بی‌اندازیم و برخی کارهای انجام‌شده اخیر در این حوزه را مرور کنیم.

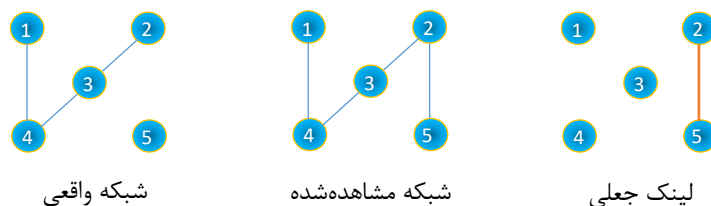
۱-۶-۲- پیش‌بینی یال‌های از دست‌رفته^۱ و جعلی^۲

معمولاً داده‌های شبکه‌های دنیای واقعی کامل نیستند و با برخی خطاها یا نواقصی همراه هستند. به‌طور مثال در بسیاری از موارد دیده می‌شود که افراد به‌دلیل برخی نگرانی‌های امنیتی یا به‌دلیل حفظ حریم خصوصی، دسترسی به صفحات خود در شبکه‌های اجتماعی را محدود می‌کنند. مسائل این‌چنینی باعث به وجود آمدن



شکل ۶-۲. نمایی از لینک‌های از دست‌رفته

برخی نواقص در دادگان شبکه‌های مورد مطالعه می‌شود و کار تحلیل و ارزیابی را با مشکل مواجه می‌کند [۲۸]. برای حل چنین مشکلی می‌توان از روش‌های پیش‌بینی یال کمک گرفت و یال‌های از دست‌رفته را بازیابی کرد. به‌طور کلی یال‌های گمشده یا از دست‌رفته به یال‌هایی گفته می‌شود که در شبکه واقعی حضور داشته اما به هر دلیلی در داده‌های مشاهده‌شده ما وجود ندارند. در شکل زیر نمونه‌ای از این مسئله نشان داده شده است.



شکل ۷-۲. نمایی از لینک‌های جعلی

^۱ Missed Links

^۲ Spurious Links

یال‌های جعلی نیز در نقطه مقابل یال‌های ازدست‌رفته قرار دارند و به آن دسته از یال‌هایی گفته می‌شوند که در گراف اصلی وجود ندارند ولی به اشتباه در گراف مشاهده شده حضور داشته‌اند. شکل زیر این موضوع را نشان می‌دهد [۲۹].

۲-۶-۲- پیش‌بینی یال در شبکه‌های مختلف

ساختار و نوع اطلاعات موجود در شبکه‌های اطلاعاتی موجب به‌وجود آمدن شبکه‌های مختلفی شده، شبکه‌هایی همچون ناهمگن^۱، چندلایه و دوبخشی^۲ نمونه‌ای از انواع شبکه‌ها می‌باشند. هرکدام از این شبکه‌ها ویژگی‌ها و مشخصه‌های خاص خود را دارند برای مثال در برخی از آن‌ها گره‌ها در همه لایه‌ها یکسان بوده و در برخی دیگر نه تنها گره‌ها بلکه نوع ارتباطات بین آن‌ها نیز متفاوت است. لذا هنگام ارائه مدل پیش‌بینی یال، حال چه با روش تعبیه و چه روش‌های مبتنی بر شباهت و یا روش‌های مبتنی بر قدم‌زنی تصادفی، لازم است که نوع شبکه را مورد توجه قرار داده و باتوجه به ویژگی‌های شبکه مورد بررسی روشی مناسب را اتخاذ کرد. در زیر به معرفی اجمالی این شبکه‌ها می‌پردازیم [۳۰].

- پیش‌بینی یال در شبکه‌های ناهمگن^۳

پیش‌بینی یال در این شبکه‌ها با توجه به عدم وجود یال و گره‌های یکسان و تعدد انواع آن‌ها در این شبکه‌ها ویژگی‌های متفاوتی را می‌طلبد چرا که انواع مختلف یال بین انواع مختلف گره به‌طور یکنواخت و با رویکرد یکسانی شکل نخواهد گرفت، برای این منظور محققان روش‌ها و راهکارهای متفاوتی را پیشنهاد کرده‌اند. برخی از آن‌ها شبکه را بر اساس یال‌ها تقسیم‌بندی کرده و شبکه‌های مختلفی از روی آن ایجاد می‌کنند سپس پیش‌بینی یال را روی هر کدام از این زیرشبکه‌های ایجاد شده انجام داده و اطلاعات به‌دست‌آمده از هر کدام را تلفیق می‌کنند [۳۱].

- پیش‌بینی یال در شبکه‌های دوبخشی^۴

نوع دیگری از شبکه‌ها وجود دارند که به‌عنوان شبکه‌های دوبخشی شناخته می‌شوند، شبکه‌های کاربر-محصول^۵ نمونه‌ای از این شبکه‌ها می‌باشند. برخی ویژگی‌ها همچون عدم شکل‌گیری مثلث‌ها^۶ و زیرگراف‌های کامل (کلیک)^۷ بر روی این شبکه‌ها موجب متمایز شدن روش‌های پیش‌بینی یال آن با شبکه‌هایی همچون تک‌بخشی^۱

^۱ Heterogeneous Networks

^۲ Bipartite

^۳ Heterogeneous Networks

^۴ Bipartite

^۵ User-Product

^۶ Triangles

^۷ Clique

می‌شود [۳۲]. برخی پژوهشگران برای حل مشکل، روش‌های کلاسیک پیش‌بینی یال همچون همسایگان مشترک^۲ را به این شبکه‌ها تعمیم داده و یا برخی دیگر این گراف‌ها را به گراف‌های تک‌بخشی تصویر کرده‌اند [۳۳] [۳۴].

- پیش‌بینی یال در شبکه‌های چندلایه

یکی از مهم‌ترین و بارزترین تفاوت‌های شبکه‌های چندلایه با شبکه‌های ناهمگن عدم وجود گره‌های متفاوت در این شبکه هاست و همانطور که پیش‌تر نیز اشاره شد، نوع گره‌ها در این شبکه‌ها یکسان می‌باشد و این تفاوت نوع یال‌ها و ارتباطات بین این گره‌هاست که باعث ایجاد لایه‌های متفاوت می‌شود. هرویتس و همکاران [۳۵] شبکه توپیتر-فوراسکوئر را بررسی کرده و تحلیل‌هایی در رابطه با همبستگی بین فعالیت‌های کاربران متناظر در هر دولایه و هم‌پوشانی آن‌ها ارائه می‌دهند. آن‌ها به این نتیجه می‌رسند که کاربرانی که در هر دو شبکه باهمدیگر ارتباط دارند (متصل هستند) نسبت به گره‌هایی که فقط در یکی از شبکه‌ها به هم متصل‌اند تعامل بیش‌تری داشته و فعال‌ترند. برخی پژوهشگران دیگر، آنتروپی را یکی از ویژگی‌های قابل تعمیم به شبکه‌های چندلایه می‌دانند و آن را با ویژگی‌های دیگر درون‌لایه‌ای (همچون تعداد همسایگان مشترک، ضریب جاکارد و آدامیک آدار) ترکیب کرده

و ترکیب حاصل را به‌عنوان یک ویژگی جدید شبکه‌های چندلایه معرفی می‌کنند [۳۶]. در پژوهش [۳۷] جلیلی و همکاران از ویژگی‌های مبتنی بر فرامسیر^۳ با فرامسیرهایی به طول حداکثر ۴ به‌عنوان ویژگی‌های بین‌لایه‌ای به پیش‌بینی یال در هر لایه می‌پردازند، همچنین آن‌ها در هر لایه دسته‌بندهای مختلف از جمله ماشین‌بردار ویژه^۴ و k -نزدیک‌ترین همسایه^۵ استفاده می‌کنند که در نتیجه با همین روش می‌توانند به بهترین نتیجه دست پیدا کنند.

۳-۶-۲- دسته‌بندی کلی روش‌های پیش‌بینی یال

روش‌های پیش‌بینی یال بسیاری تا به حال ارائه شده که می‌توان آن‌ها را به دو دسته کلی تقسیم‌بندی کرد:

(آ) رویکردهای مبتنی بر شباهت (ب) رویکردهای مبتنی بر یادگیری

¹ Unipartite (Monopartite)

² Common Neighbor

³ Meta-Path

⁴ Classifier

⁵ Support Vector Machine (SVM)

⁶ K-Nearest Neighbor (KNN)

- رویکردهای مبتنی بر شباهت

ویژگی‌های مختلف بسیاری مبتنی بر شباهت تعریف شده‌اند که این رویکردها با محاسبه این ویژگی‌ها به پیش‌بینی یال می‌پردازند. ویژگی‌های مبتنی بر شباهت خود به چند دسته مختلف از جمله ویژگی‌های مبتنی بر همسایگی گره (شامل همسایگان مشترک، ضریب جاکارد، معیار آدامیک آدار، اتصال ترجیحی)، ویژگی‌های مبتنی بر مسیر (شامل فاصله کوتاه‌ترین مسیر و معیار کاتز) همچنین ویژگی‌های مبتنی بر قدم‌زنی تصادفی (شامل زمان اصابت^۱، زمان سفر^۲ و رتبه‌بندی صفحه‌اصلی^۳) تقسیم‌بندی می‌شوند. رویکردهای مبتنی بر شباهت با محاسبه یک یا چندین معیار از معیارهای معرفی‌شده، نمراتی را به هر یال گراف تخصیص داده و سپس این نمرات را مرتب کرده و آن‌ها را رتبه‌بندی می‌کنند، جفت‌گره‌هایی که نمرات شباهت بالاتری اخذ کرده‌اند به احتمال بیش‌تری باهم در ارتباط خواهند بود و لینک بینشان برقرار است.

- رویکردهای مبتنی بر یادگیری

نقطه مقابل رویکردهای مبتنی بر شباهت، رویکردهای مبتنی بر یادگیری هستند که هرکدام از معیارهای معرفی‌شده در بخش قبلی و یا معیارهای جدید همگی به عنوان یک ویژگی در اختیار یک مدل یادگیری ماشینی قرار می‌گیرند (مانند آنچه که در تعبیه گراف‌ها شاهد آن هستیم) و سپس مدل نمره مثبت را به یال‌های واقعی (موجود در گراف) و نمره منفی را به یال‌های غایب اطلاق می‌کند و در ادامه با یادگیری و بهینه‌شدن مدل به پیش‌بینی یال‌های از دست‌رفته می‌پردازد [۳۸].

نکته قابل توجه اینکه هیچ کدام از روش‌های یاد شده را نمی‌توان به طور قطع برتر بر دیگری دانست و هرکدام از روش‌ها در کاربردهای مختلف و یا بر روی شبکه‌های مختلف عملکرد متفاوتی از خود نشان خواهند داد، اما با این حال، رویکردهای مبتنی بر یادگیری اخیراً مورد توجه بسیاری از محققان قرار گرفته و در مجموع پیش‌بینی‌های دقیق‌تر و دقت‌های نسبتاً بالاتری را به نسبت سایر رویکردها ارائه می‌کنند.

در ادامه به معرفی چندین روش پیش‌بینی یال می‌پردازیم که برخلاف روش‌های قبلی که همگی برپایه یادگیری بودند، این روش‌ها مبتنی بر شباهت هستند که روش‌های جدیدی محسوب می‌شوند و به دقت‌های نسبتاً خوبی در پیش‌بینی یال دست یافته‌اند.

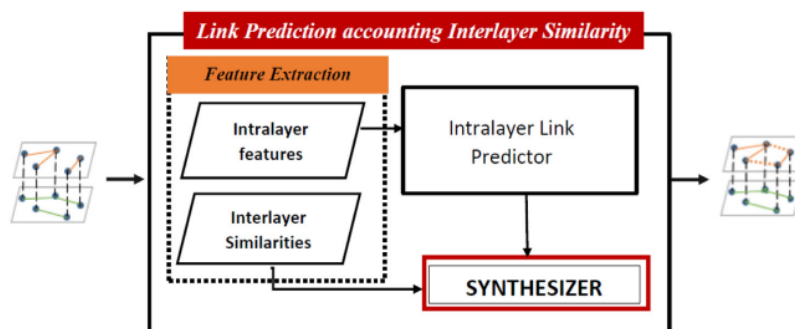
¹ Hitting Time

² Commute Time

³ Rooted PageRank

۷-۲- مدل LPIS

این مدل از جمله روش‌های مبتنی بر شباهت است که با در نظر گرفتن اطلاعات درون‌لایه‌ای و بین‌لایه‌ای اقدام به پیش‌بینی یال در شبکه‌های چندلایه می‌کند. چارچوب کلی این روش در شکل زیر آمده است.



شکل ۲-۸. چارچوب کلی مدل LPIS [۳۰]

در این روش ابتدا ویژگی‌های درون‌لایه‌ای توسط معیارهای مبتنی بر شباهت اشاره شده در بخش قبل (همچون همسایگان مشترک، ضریب جاکارد و ...) محاسبه شده و توسط دسته‌بند لاجستیک رگرسیون^۲ یک احتمال تشکیل یال برای هر جفت گره درون هر لایه اختصاص می‌یابد. سپس با استفاده از معیارهای شباهت بین‌لایه‌ای (همچون همبستگی درجه به درجه^۳ شباهت دولایه محاسبه می‌شود، حال با در دست داشتن شباهت بین لایه‌ای و همچنین

احتمال تشکیل یال درون‌لایه‌ای، از یک تابع ترکیب‌کننده برای پیش‌بینی احتمال تشکیل یال در لایه مقصد استفاده می‌شود. توجیه استفاده از این تابع این است که اگر به عنوان مثال یالی بین گره i و گره j در لایه ۱ وجود داشت و در صورت شبیه بودن لایه ۱ به لایه ۲، احتمال اینکه همان یال در لایه ۲ نیز تکرار شود وجود دارد. مشخص است که هرچه شباهت بین دولایه و احتمال تشکیل یال در لایه اول بالاتر باشد احتمال تشکیل همان یال در لایه بعدی هم بیش‌تر می‌شود [۳۰].

۸-۲- مدل EMLP

در این پژوهش، یک روش مبتنی بر شباهت جدید برای پرداختن به مسئله پیش‌بینی یال در شبکه‌های چندلایه با استفاده از نظریه شواهد^۴ ارائه شده است. نظریه شواهد، می‌تواند به طور دقیق اطلاعات نامشخص و غیرقطعی را

^۱ Link Prediction Accounting Interlayer Similarity

^۲ Logistic Regression Classifier

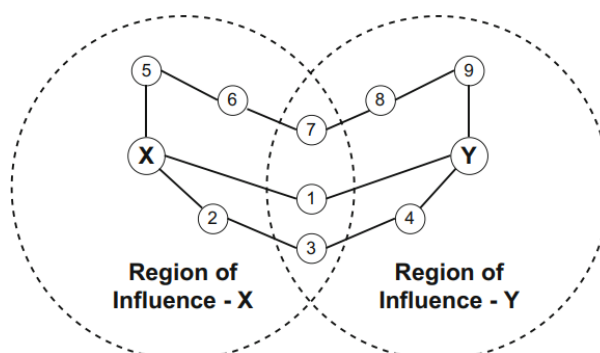
^۳ Degree-Degree Correlation (DDC)

^۴ Evidence Theory

از منابع مختلف اخذ کرده و آن‌ها را برای تصمیم‌گیری مناسب ترکیب کند. هدف از این کار استفاده از اطلاعات موجود در لایه‌های دیگر برای پیش‌بینی یال در لایه مقصد است. برای پیش‌بینی پیوندها در لایه هدف، روش پیشنهادی اطلاعات همه لایه‌ها را ترکیب می‌کند. بنابراین، هر لایه یک منبع شواهد در نظر گرفته می‌شود. برای پیوندهای بالقوه در لایه هدف، نمرات شباهت آن‌ها که از لایه‌های مختلف به دست می‌آید به‌عنوان شواهد متفاوت که از منابع متفاوت به‌دست‌آمده‌اند تلقی می‌شود. درواقع نظریه شواهد شباهت‌های به‌دست‌آمده از هر لایه را به‌عنوان یک منبع شواهد در نظر گرفته و با ترکیب این منابع (ترکیب اطلاعات به‌دست‌آمده از لایه‌های مختلف) سعی در پیش‌بینی پیوند در لایه مقصد دارد. همچنین در این مدل، برای محاسبه شباهت یک جفت‌گره درون هر لایه، یک معیار شباهت جدید به نام ECM با الهام از مدل هدایت‌الکترونیکی موجود در مدارهای الکترونیکی استفاده می‌شود. در حوزه الکترونیک، رسانایی الکتریکی، که رسانایی یک قطعه رسانا را اندازه‌گیری می‌کند، به‌عنوان معکوس مقاومت آن تعریف می‌شود. هر چه رسانایی آن قطعه بیش‌تر باشد، عبور جریان از آن آسان‌تر خواهد بود. لذا با الهام از این پدیده فیزیکی، یک شاخص برای محاسبه شباهت دو گره درون هر لایه پیشنهاد شده‌است. برای این منظور، یک گره در شبکه را به عنوان رسانا و انتقال‌دهنده جریان، همچنین درجه آن گره به‌عنوان مقاومت آن در نظر گرفته شده‌است. حال با این تعریف، معکوس درجه هر گره نشانگر رسانایی آن خواهد بود. در این سناریو می‌توان شباهت دو گره را به‌عنوان مجموع مقادیر رسانایی گره‌ها در تمام مسیرهای منتهی به آن‌ها در شبکه تعریف کرد. با این حال، از آنجایی که پیمودن تمام مسیرها بین دو گره زمان‌بر است، در این پژوهش فقط مسیریابی به طول ۲ و ۳ در نظر گرفته شده‌است. همچنین برای اندازه‌گیری شباهت بین لایه‌ها از ویژگی‌های شباهت بین‌لایه‌ای همچون ضریب خوشه‌بندی استفاده شده‌است [۳۹].

۹-۲- مدل HOPLP

هدف از این روش هم همچون روش پیشین استفاده از لایه‌های دیگر جهت پیش‌بینی یال در لایه مقصد است. در این پژوهش از یک مدل تجمیع استفاده می‌شود، این مدل تجمیع اطلاعات لایه‌های مختلف را در یک شبکه واحد وزن‌دار خلاصه می‌کند. انگیزه اصلی در پس این روش، استفاده از مسیرهای مرتبه بالاتر برای تخمین بهتر شباهت‌های همسایگی است. با توجه به ویژگی «سه درجه اثرگذاری» اشاره شده در مقاله [۴۰] می‌توان نتیجه گرفت که هر گره می‌تواند تا همسایه‌های درجه سه خود (همسایه‌هایی با طول ۳) اثربخش باشد.



¹ Electronic Conductance Measure

² Three Degrees of Influence

همانطور که در شکل بالا آمده به وضوح قابل مشاهده است که اگر ناحیه اثرگذاری گره‌ها را به فاصله یک گام^۱ از گره مبدأ فرض کنیم، تنها مسیر مرتبط بین گره‌های X و Y مسیر $X-1-Y$ خواهد بود. اگر فرض کنیم که ناحیه اثرگذاری یک گره از نقطه مبدأ، به دو گام افزایش یابد، مسیرهای مربوطه برای پیش‌بینی پیوند $X-1-Y$ و $X-2-3-4-Y$ خواهند بود. در نهایت با استفاده از پدیده ۳ درجه اثرگذاری، می‌توانیم فرض کنیم که ناحیه نفوذ یک گره تا ۳ گام دورتر از آن گسترش می‌یابد. با استفاده از این فرض، می‌توانیم ببینیم که برای محاسبه احتمال پیوند بین X و Y ، باید یک فرامسیر متشکل از تمامی مسیرهای ممکن اعم از مسیرهای کوتاه‌تر ($X-1-Y$ و $X-2-3-4-Y$) و همچنین مسیرهای طولانی‌تر (یعنی $X-5-6-7-8-9-Y$) را برای پیش‌بینی دقیق‌تر یال استفاده کنیم.

درواقع با این رویکرد به جای همسایگی‌های محلی (همسایگان اصلی هر گره) همسایگان سراسری آن نیز پوشش داده می‌شود که به عقیده نویسندگان این مقاله، رویکرد مناسب‌تری است. این روند محاسبه احتمال پیوند، بسیار شبیه به روش محاسبه شباهت مبتنی بر تخصیص منابع^۲ می‌باشد با این تفاوت که روش تخصیص منابع تنها مسیرهای با طول دو را در نظر می‌گیرد در حالی که در این روش‌های مسیرهایی با طول بیش‌تر نیز مورد توجه است.

۱۰-۲- جمع‌بندی

در جدول مروری که در ادامه آمده، مروری کلی بر پیشینه مسائل تعبیه گراف و پیش‌بینی یال داریم.

¹ Hop

² Resource Allocation (RA)

جدول ۱-۲. دسته‌بندی انواع روش‌های تعبیه گراف و پیش‌بینی‌یال

ویژگی‌ها	پیچیدگی زمانی	مدل	سال	دسته بندی
روش‌های تعبیه گراف	جزو اولین روش‌های تجزیه گراف هستند و به‌دلیل استفاده از ماتریس همبندی به‌عنوان شباهت بین گره‌ها، شباهت‌های مرتبه بالاتر بین گره‌ها حفظ نمی‌شوند (گره‌های دورتر از همسایگی اول، تأثیری در محاسبه تعبیه‌ها ندارند). تعبیه گره‌ها از طریق فاکتورگیری ماتریس همبندی گراف حاصل می‌شود، به‌گونه‌ای که گره‌های متصل به‌هم، تعبیه مشابهی داشته باشند.	<i>Laplacian Eigenmaps</i>	۲۰۰۱	تجزیه‌ماتریس <i>Matrix Factorization</i>
		<i>Graph Factorization</i>	۲۰۱۳	
		<i>GraRep</i>	۲۰۱۵	
		<i>HOPE</i>	۲۰۱۶	
	به‌جای استفاده از یک معیار ثابت جهت اندازه‌گیری شباهت بین گره‌ها (همچون ماتریس همبندی که در روش‌های تجزیه گراف استفاده می‌شد) از قدم‌زنی تصادفی برای حصول این شباهت استفاده می‌شود. این موجب حفظ شباهت‌های مرتبه بالاتر بین گره‌ها و افزایش دقت مدل در پیش‌بینی‌ها می‌شود. گره‌هایی که در یک قدم‌زنی تصادفی کوتاه، هم‌رخداد بوده‌اند، تعبیه‌های مشابهی خواهند داشت.	<i>DeepWalk</i>	۲۰۱۴	قدم‌زنی تصادفی <i>Random Walk</i>
		<i>Node2Vec</i>	۲۰۱۶	
	برخلاف روش‌های قبلی، اطلاعات و ویژگی خود گره‌ها نیز در محاسبه تعبیه آن‌ها تأثیرگذار است، پیچیدگی زمانی بسیار پایینی دارند و دقت پیش‌بینی مناسبی ارائه می‌دهند. برای شبکه‌های تک‌لایه طراحی شده‌اند و در صورت دردسترس نبودن ویژگی گره‌ها دقت پیش‌بینی آن‌ها کاهش می‌یابد، همچنین دارای طراحی استنتاجی هستند، به‌گونه‌ای که تعبیه گره‌های مشاهده‌نشده را می‌توان از روی پارامترهای آموخته‌شده بدون نیاز به بازمحاسبه تعبیه کل گره‌ها به دست آورد.	<i>GCN</i>	۲۰۱۶	یادگیری عمیق <i>Deep Learning</i>
		<i>GAE</i>	۲۰۱۷	
روش‌های پیش‌بینی یال	روش‌های نوین پیش‌بینی یال در شبکه‌های چندلایه، غیرتعبیه‌ای و غیر یادگیرنده هستند که از شباهت‌های درون‌لایه‌ای و بین‌لایه‌ای از پیش تعریف‌شده توسط کاربر، جهت اندازه‌گیری شباهت بین دو گره و در نتیجه احتمال وجود یال بین آن‌ها استفاده می‌کنند. به‌جهت اینکه برای شبکه‌های چندلایه طراحی شده‌اند، دقت‌های خوبی ارائه می‌دهند، ولی دارای مشکل مقیاس‌پذیری به‌دلیل پیچیدگی زمانی نسبتاً بالا می‌باشند.	<i>LPIS</i>	۲۰۱۹	مبتنی بر شباهت (همسایگی گره)
		<i>EMLP</i>	۲۰۲۲	
	مبتنی بر شباهت (مسیر)	<i>HOPLP</i>	۲۰۲۲	

۳- برنامه آتی

نحوه اجرا

بطور مشخص برای ادامه کار، برنامه زیر در نظر گرفته شده است:

- در قدم اول به مطالعه و بررسی روشهای ارائه شده مرتبط با موضوع طرح می پردازیم.
- پس از بررسی مدل های موجود، استخراج ایده اولیه و تحلیل نظری آن در دستور کار قرار می گیرد.
- در ادامه به آماده سازی و پیاده سازی روش پیشنهادی خواهیم پرداخت.
- پس از پیاده سازی روش پیشنهادی به ارزیابی آن و مقایسه با روش های ارائه شده دیگر خواهیم پرداخت.
- در نهایت مقاله مرتبط با روش پیشنهادی را آماده می کنیم.

زمانبندی طرح

برنامه زمانی اجرای طرح در جدول زیر آورده شده است.

جدول ۱-۳ برنامه زمانی اجراء طرح

زمان (ماه)	فعالیت	۱-۲	۲-۴	۴-۵	۵-۷	۷-۹	۹-۱۰	۱۱-۱۲
	تکمیل مطالعه مقالات مرتبط با روش های تعبیه گراف							
	تکمیل مطالعه روش های پیش بینی یال در شبکه های چندلایه							
	ارائه ایده اولیه حل مسئله							
	انتخاب، جمع آوری و آماده سازی مجموعه داده							
	پیاده سازی ابتدایی روش پیشنهادی							
	استخراج و ارزیابی نتایج اولیه							
	اصلاح و بهبود روش پیشنهادی باتوجه به بازخوردها							
	استخراج و ارزیابی نتایج نهایی و مقایسه با روش های پیشین							
	نگارش مقاله							

همکاران طرح

(۱) سعید پیروزفر، دانشجوی کارشناسی ارشد، دانشگاه تهران
آدرس: تهران، خ کارگر شمالی، جنب دانشکده فنی دانشگاه تهران، دانشکده علوم و فنون نوین

(۲) شقایق نجاری، دانشجوی دکتری، دانشگاه تهران
آدرس: تهران، خ کارگر شمالی، جنب دانشکده فنی دانشگاه تهران، دانشکده علوم و فنون نوین

مراجع

- [1] P. Goyal and E. Ferrara, "Graph embedding techniques, applications, and performance: A survey," *Knowledge-Based Systems*, p. 19, 2018.
- [2] A.-L. Barabási, "Network science," *Cambridge university press*, 2016.
- [3] Y. Sun, J. Han, X. Yan, P. S. Yu and T. Wu, "Pathsim: Meta path-based top-k similarity search in heterogeneous information networks," *Proceedings of the VLDB Endowment*, vol. 4, pp. 992-1003, 2011.
- [4] M. Kivelä, A. Arenas, M. Barthélemy, J. P. Gleeson, Y. Moreno and M. A., "Multilayer networks," *Journal of complex networks*, vol. 2, no. 3, pp. 203-271, 2014.
- [5] X. Wang, D. Bo, C. Shi, S. Fan, Y. Ye and P. S. Yu, "A Survey on Heterogeneous Graph Embedding: Methods, Techniques, Applications and Sources," *IEEE Transactions on Big Data*, 2022.
- [6] W. L. Hamilton, R. Ying and J. Leskovec, "Representation Learning on Graphs: Methods and Applications," *arXiv:1709.05584v3 [cs.SI]*, 2018.
- [7] M. Xu, "Understanding graph embedding methods and their applications," *SIAM Review*, vol. 63.4, pp. 825-853, 2021.
- [8] P. Goyal, S. R. Chhetri and A. Canedo, "dyngraph2vec: Capturing network dynamics using dynamic graph representation learning," *Knowledge-Based Systems*, vol. 187, no. 104816, 2020.
- [9] J. Tang, M. Qu, M. Wang, M. Zhang, J. Yan and Q. Mei, "Line: Large- scale information network embedding," *Proceedings 24th International Conference on World Wide Web*, p. 1067-1077, 2015.
- [10] D. Wang, P. Cui and W. Zhu, "Structural deep network embedding," *Proceedings of the 22nd International Conference on Knowledge Discovery and Data Mining, ACM*, p. 1225-1234, 2016.
- [11] B. Perozzi, R. Al-Rfou and S. Skiena, "Deepwalk: Online learning of social representations," *Proceedings 20th international conference on Knowledge discovery and data mining*, p. 701-710., 2014.
- [12] C. H. Ding, X. He, H. Zha, M. Gu and H. D. Simon, "A min-max cut algorithm for graph partitioning and data clustering," *International Conference on Data Mining, IEEE*, p. 107-114., 2001.
- [13] S. Bhagat, G. Cormode and S. Muthukrishnan, "Node classification in social networks," *Social network data analytics, Springer*, p. 115- 148., 2011.

- [14] A. Grover and J. Leskovec, "node2vec: Scalable feature learning for networks," *Proceedings of the 22nd International Conference on Knowledge Discovery and Data Mining, ACM*, p. 855–864, 2016.
- [15] E. Omodei, M. D. D. Domenico and A. Arenas, "'Characterizing interactions in online social networks during exceptional events," *Frontiers in Physics*, p. 59, 2015.
- [16] M. Gustafsson, C. E. Nestor, H. Zhang, A.-L. Barabási, S. Baranzini and S. Brunak, "Modules, networks and systems medicine for understanding disease and aiding diagnosis," *Genome medicine*, vol. 6, no. 10, p. 82, 2014.
- [17] S. T. Roweis and L. K. Saul, "Nonlinear dimensionality reduction by locally linear embedding," *Science* 290 (5500), p. 2323–2326, 2000.
- [18] T. Kipf and M. Welling, "Variational graph auto-encoders," *NIPS Workshop on Bayesian Deep Learning*, 2016.
- [19] T. N. Kipf and M. Welling, "Semi-supervised classification with graph convolutional networks," *arXiv preprint arXiv:1609.02907*, 2016.
- [20] W. Hamilton, R. Ying and J. Leskovec, "Inductive representation learning on large graphs," *arXiv preprint, arXiv:1603.04467*, 2017.
- [21] M. Belkin and P. Niyogi, "Laplacian eigenmaps and spectral techniques for embedding and clustering," *NIPS*, 2002.
- [22] A. Ahmed, N. Shervashidze, S. Narayanamurthy, V. Josifovski and A. J. Smola, "Distributed large-scale natural graph factorization," *WWW*, 2013.
- [23] S. Cao, W. Lu and Q. Xu, "Grarep: Learning graph representations with global structural information," *KDD*, 2015.
- [24] M. Ou, P. Cui, J. Pei, Z. Zhang and W. Zhu, "Asymmetric transitivity preserving graph embedding," *KDD*, 2016.
- [25] M. Niepert, M. Ahmed and K. Kutzkov, "Learning convolutional neural networks for graphs," *Proceedings of the 33rd annual international conference on machine learning. ACM*, 2016.
- [26] T. N. Kipf and M. Welling, "Variational graph auto-encoders," *arXiv preprint arXiv*, vol. 1611.07308, 2016.
- [27] Y. Wang, B. Xu, M. Kwak and X. Zeng, "A simple training strategy for graph autoencoder," *Proceedings of the 2020 12th International Conference on Machine Learning and Computing*, pp. 341–345, 2020.
- [28] L. Pan, T. Zhou, L. Lü and C.-K. Hu, "'Predicting missing links and identifying spurious links via likelihood analysis," *Scientific reports*, vol. 6, p. 22955, 2016.
- [29] G.-P. G., R. Aliakbarisani, A. Ghasemi and M. Á. Serrano, "Precision as a measure of predictability of missing links in real networks," *Physical Review E*, vol. 101, p. 5, 2020.
- [30] S. Najari, M. Salehi, V. Ranjbar and M. Jalili, "Link prediction in multiplex networks based on interlayer similarity," *Physica A: Statistical Mechanics and its Applications*, vol. 536, no. 120978, 2019.
- [31] J. Zhang, X. Kong and P. S. Yu, "Transferring heterogeneous links across location-based social networks," *Proceedings of the 7th ACM international conference on Web search and data mining*, pp. 303–312, 2014.
- [32] P. Kumar and D. Sharma, "A novel similarity measure for the link prediction in unipartite and bipartite networks," *Social Network Analysis and Mining*, vol. 11, no. 1, pp. 1–14, 2021.

- [33] O. Allali, C. Magnien and M. Latapy, "'Link prediction in bipartite graphs using internal links and weighted projection," *Computer Communications Workshops (INFOCOM WKSHPS)*, pp. 936-941, 2011.
- [34] S. Xia, B. Dai, E.-P. Lim, Y. Zhang and C. Xing, "Link prediction for bipartite social networks: the role of structural holes," *Advances in Social Networks Analysis and Mining (ASONAM)*, pp. 143-157, 2012.
- [35] D. Hristova, A. Noulas, C. Brown, M. Musolesi and C. Mascolo, "A multilayer approach to multiplexity and link prediction in online geo-social networks," *EPJ Data Science*, vol. 5, no. 2, p. 24, 2016.
- [36] Kanawati, M. Pujari and R., "Link prediction in multiplex networks," *NHM*, vol. 10, no. 1, pp. 17-35, 2015.
- [37] M. Jalili, Y. Orouskhani, M. Asgari, N. Alipourfard and M. Perc, "Link prediction in multiplex online social networks," *Royal Society open science*, vol. 14, no. 2, p. 160863, 2017.
- [38] J. Janssen, M. Hurshman and N. Kalyaniwalla, "'Model selection for social networks using graphlets," *Internet Mathematics*, vol. 8, no. 4, pp. 338-363, 2012.
- [39] H. Luo, L. Li, H. Dong and X. Chen, "Link prediction in multiplex networks: An evidence theory method," *Knowledge-Based Systems*, vol. 257, no. 109932, 2022.
- [40] N. A. Christakis and J. H. Fowler, "Social contagion theory: examining dynamic social networks and human behavior," *Statistics in medicine*, vol. 32, no. 4, pp. 556-577, 2013.
- [41] M. E. Newman, "Scientific collaboration networks. ii. shortest paths, weighted networks, and centrality," *Physical review E*, vol. 64, no. 1, p. 016132, 2001.
- [42] A.-L. Barabási, *Network Science*, Cambridge University Press, 2016.
- [43] F. Battiston, V. Nicosia and V. Latora, "Structural measures for multiplex networks," *PHYSICAL REVIEW E*, vol. 89, p. 032804, 2014.
- [44] D. Zhao, L. Li, H. Peng, Q. Luo and Y. Yang, "Multiple routes transmitted epidemics on multiplex networks," *Physics Letters A*, vol. 378, no. 10, pp. 770-776, 2014.
- [45] Y. Yao, R. Zhang, F. Yang, Y. Yuan, Q. Sun, Y. Qiu and R. Hu, "Link prediction via layer relevance of multiplex networks," vol. 28, no. 08, 2017.
- [46] V. Gemmetto and D. Garlaschelli, "Multiplexity versus correlation: the role of local constraints in real multiplexes," *Scientific reports*, vol. 5, no. 1, pp. 1-7, 2015.
- [47] R. J. Mondragón, "Estimating degree-degree correlation and network cores from the connectivity of high-degree nodes in complex networks," *Scientific reports*, vol. 10, no. 1, pp. 1-24, 2020.
- [48] S. Zhou and R. Mondragón, "Accurately modeling the Internet topology," *Phys. Rev. E*, vol. 70, no. 066108, 2004.
- [49] V. Colizza, A. Flammini, M. A. Serrano and A. Vespignani, "Detecting rich-club ordering in complex networks," *Nature*, vol. 2, pp. 110-115, 2006.
- [50] A. E. Brouwer and W. H. Haemers, *Spectra of graphs*, Springer Science & Business Media, 2011.
- [51] Philip, J. Zhang and S. Y., "Link prediction across heterogeneous social networks: A survey," *USA, University of Illinois at Chicago*, 2014.
- [52] M. D. Domenico, C. Granell, M. A. Porter and A. Arenas, "The physics of spreading processes in multilayer networks," *Nature Physics*, vol. 12, no. 10, p. 901, 2016.
- [53] M. D. Domenico, A. Solé-Ribalta, E. Cozzo, M. Kivelä, Y. Moreno, M. A. Porter, S. Gómez and A. Arenas, "Mathematical formulation of multilayer networks," *Physical Review X*, vol. 3, no. 4,

p. 041022, 2013.

- [54] M. D. Domenico, C. Granell, M. A. Porter and A. Arenas, "The physics of spreading processes in multilayer networks," *Nature Physics*, vol. 12, no. 10, p. 901, 2016.
- [55] S. Mishra, S. S. Singh, A. Kumar and B. Biswas, "HOPLP– MUL: link prediction in multiplex networks based on higher order paths and layer fusion," *Applied Intelligence*, pp. 1-29, 2022.
- [56] A. C. Belkina, C. O. Ciccolella, R. Anno, R. Halpert, J. Spidlen and J. E. Snyder-Cappione, "Automated optimized parameters for T-distributed stochastic neighbor embedding improve visualization and analysis of large datasets," *Nature communications*, vol. 10, no. 1, pp. 1-12, 2019.
- [57] H. Abdi and L. J. Williams, "Principal component analysis," *Wiley interdisciplinary reviews: computational statistics*, vol. 4, no. 2, 2010.
- [58] McInnes, Leland, J. Healy and J. Melville, "Umap: Uniform manifold approximation and projection for dimension reduction," *arXiv preprint arXiv:1802.03426*, 2018.
- [59] D. Liben-Nowell and J. Kleinberg, "The link-prediction problem for social networks," *journal of the Association for Information Science and Technology* 58 (7), p. 1019–1031, 2007.
- [60] A. M. D. D. S. G. a. A. A. Solé-Ribalta, "Random walk centrality in interconnected multilayer networks," *Physica D: Nonlinear Phenomena*, vol. 323, pp. 73-79, 2016.
- [61] L.v.d.Maaten and G.Hinton, "Visualizing data using t-sne," *Journal of Machine Learning Research* 9, p. 2579–2605, 2008.
- [62] A. Solé-Ribalta, M. D. Domenico, S. Gómez and A. Arenas, "Random walk centrality in interconnected multilayer networks," *Physica D: Nonlinear Phenomena*, vol. 323, pp. 73-79, 2016.
- [63] A. Grover and J. Leskovec, "Scalable Feature Learning for Networks," *arXiv:1607.00653v1 [cs.SI]*, 2016.
- [64] A. Cardillo, Gómez-Gardenes, Z. J., R. M., P. M., P. D., F. D. and S. Boccaletti, "Emergence of network features from multiplexity," *Scientific reports*, vol. 3, pp. 1-6, 2013.
- [65] Cao, Shaosheng, W. Lu and Q. Xu, "Deep neural networks for learning graph representations," *Proceedings of the Thirtieth AAAI Conference on Artificial Intelligence*, AAAI Press, p. 1145–1152, 2016.