



Carnegie
Mellon
University

Graph Neural Networks

*11-485/685/785: Introduction to Deep Learning
Spring 2026*

Shravanth Srinivas



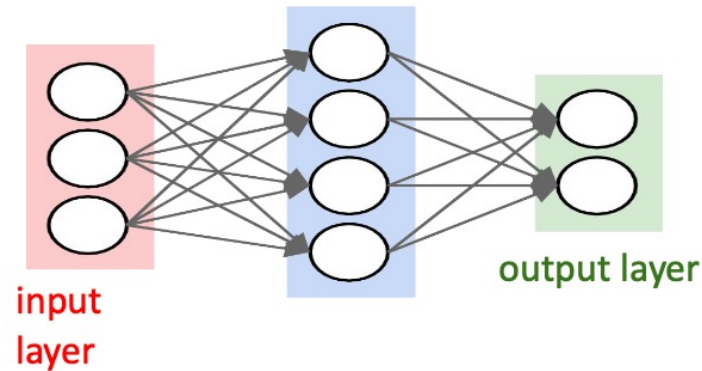
Models so far

- **MLPs are universal function approximators**
 - Boolean functions, classifiers, and regressions
- **MLPs can be trained through variations of gradient descent**
 - Gradients can be computed by backpropagation



MLP Model

Or, more generally a
vector input



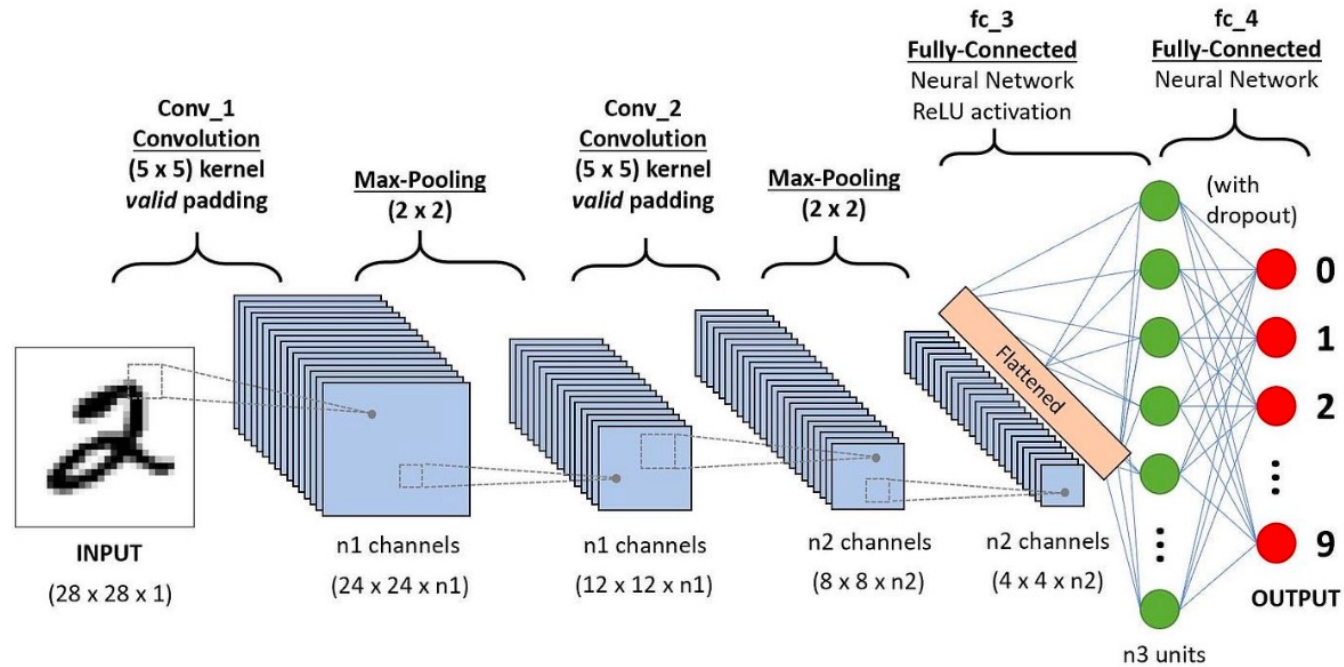
- Can recognize patterns in data
 - E.g. digits
 - Or any other vector data



Models so far

- **CNNs designed for image and spatial data**
 - Convolutional layers learn spatial patterns (e.g., edges, textures).
 - Pooling layers reduce spatial dimensions while retaining key features.
- **CNNs can be trained through variations of gradient descent**
 - Gradients can be computed by backpropagation

CNN Model





Models so far

- **Sequence-to-Sequence Models: sequential data.**
 - RNNs, LSTMS, Transformers
 - Encode input sequence and decode the encoded sequence.
- **RNNs, LSTMS, Transformers can be trained through variations of gradient descent**
 - Gradients can be computed by backpropagation

Sequence-to-Sequence Model

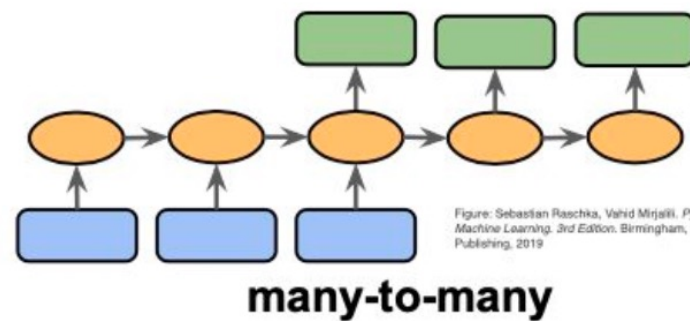
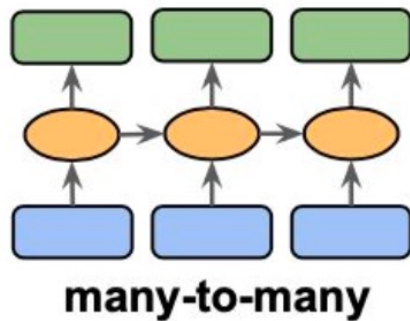
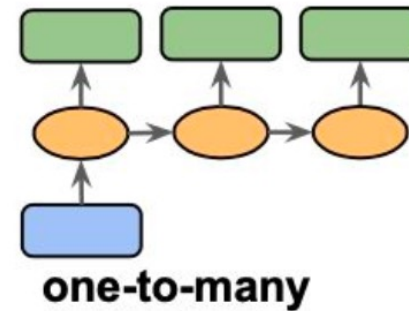
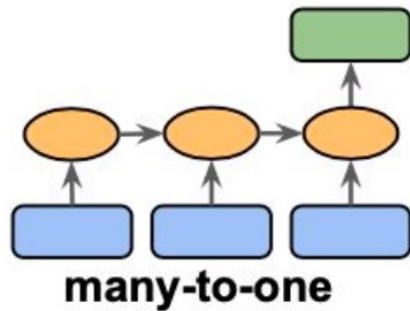
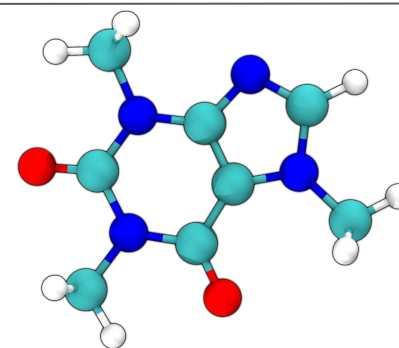


Figure: Sebastian Raschka, Vahid Mirjalili, Python Machine Learning, 2nd Edition, Birmingham, UK: Packt Publishing, 2019



How about

3d representation of the Caffeine molecule

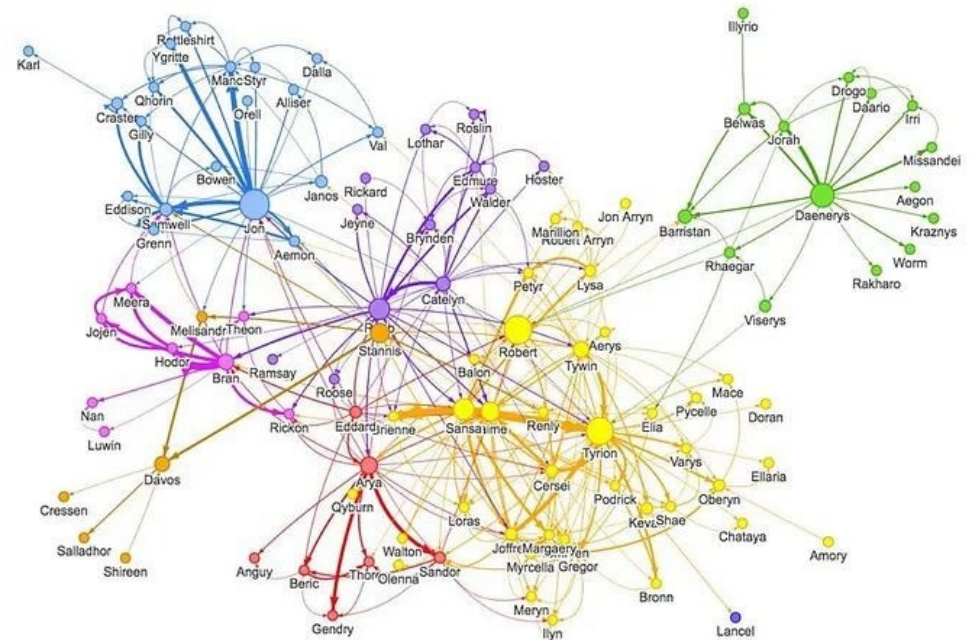


Credits: <https://distill.pub/2021/gnn-intro/>

Carnegie Mellon University

How about

Social Network



Credits: medium.com/dair-ai

Carnegie Mellon University



Why Graphs Matter

Social networks: Relationships, influence propagation

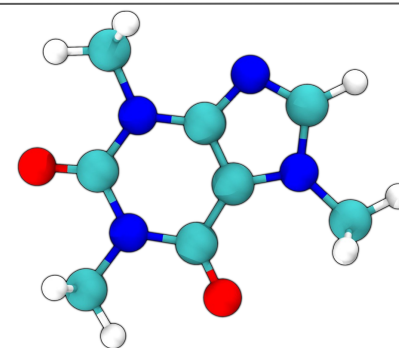
Molecules: Chemical properties depend on structure

Citation networks: Knowledge graphs, recommendations

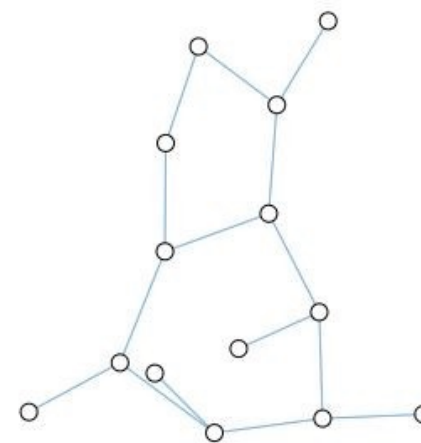
Traffic systems: Routing, congestion prediction

Why Graphs Matter

3d representation of the Caffeine molecule

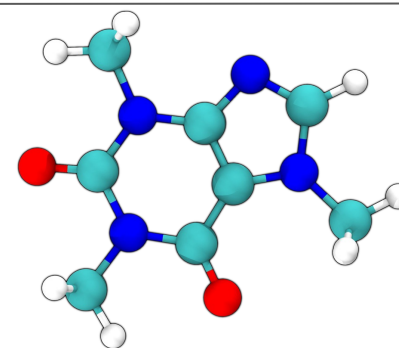


Graph representation of the Caffeine molecule

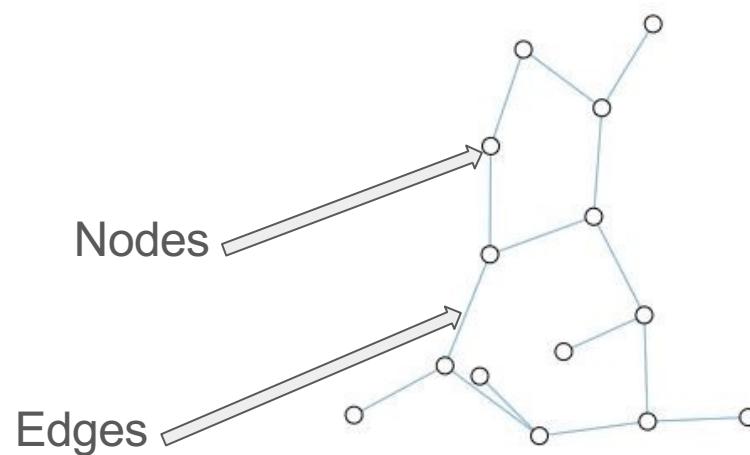


Why Graphs Matter

3d representation of the Caffeine molecule



Graph representation of the Caffeine molecule



Carnegie Mellon University



Why should we Deep Learning people care?



Why should we Deep Learning people care?

Real-World Applications

- **Drug Discovery:** Predict molecular properties, design better compounds
- **Recommendation:** User-item interactions, collaborative filtering
- **Forecasting:** Traffic, weather, energy consumption on grids



Poll 1

True or False

1. Euclidean data refers to data that lies in a space where the distance between points is calculated using the Euclidean distance formula, while non-Euclidean data involves spaces where the concept of distance may follow different rules, such as hyperbolic or graph-based distances.
2. CNNs and MLPs are specifically designed to handle non-Euclidean data, such as graphs and hyperbolic spaces, without any modifications.



Poll 1

True or False

1. Euclidean data refers to data that lies in a space where the distance between points is calculated using the Euclidean distance formula, while non-Euclidean data involves spaces where the concept of distance may follow different rules, such as hyperbolic or graph-based distances. **(True)**
2. CNNs and MLPs are specifically designed to handle non-Euclidean data, such as graphs and hyperbolic spaces, without any modifications. **(False)**



Challenges in handling Graph Data

Models like MLPs, CNNs, and RNNs are built on the assumption that data is structured in a regular, fixed grid (like pixels in an image or steps in a sequence).

Problem:

1. Graphs are irregular and unstructured. Nodes can have a variable number of neighbors, and the connections don't follow a simple, predetermined pattern.



Challenges in handling Graph Data

Models like MLPs, CNNs, and RNNs are built on the assumption that data is structured in a regular, fixed grid (like pixels in an image or steps in a sequence).

Problem:

1. Graphs are irregular and unstructured. Nodes can have a variable number of neighbors, and the connections don't follow a simple, predetermined pattern.
2. Graphs lack a natural order. Models must be permutation-invariant because re-labeling nodes cannot change the result.



Challenges in handling Graph Data

Models like MLPs, CNNs, and RNNs are built on the assumption that data is structured in a regular, fixed grid (like pixels in an image or steps in a sequence).

Problem:

1. Graphs are irregular and unstructured. Nodes can have a variable number of neighbors, and the connections don't follow a simple, predetermined pattern.
2. Graphs lack a natural order. Models must be permutation-invariant because re-labeling nodes cannot change the result.
3. Standard models cannot easily incorporate the connectivity information. They need a mechanism to aggregate features from neighbors



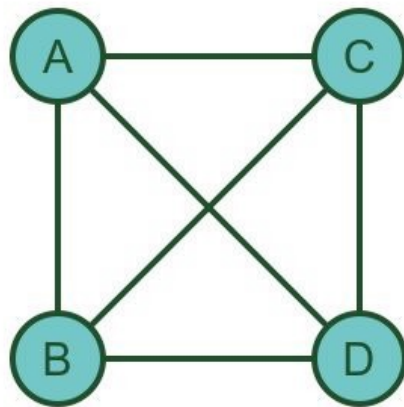
The entire research area dedicated to developing deep learning methods for Non-Euclidean data like graphs, manifolds, and more is called **Geometric Deep Learning**.

Graph Basics

Graph is an ordered pair $G = (V, E)$:

Nodes V (set of vertices)

Edges E (associated with two distinct vertices) connecting them



Complete graph

	A	B	C	D
A	0	1	1	1
B	1	0	1	1
C	1	1	0	1
D	1	1	1	0

Credits: GraphicMaths

Carnegie Mellon University

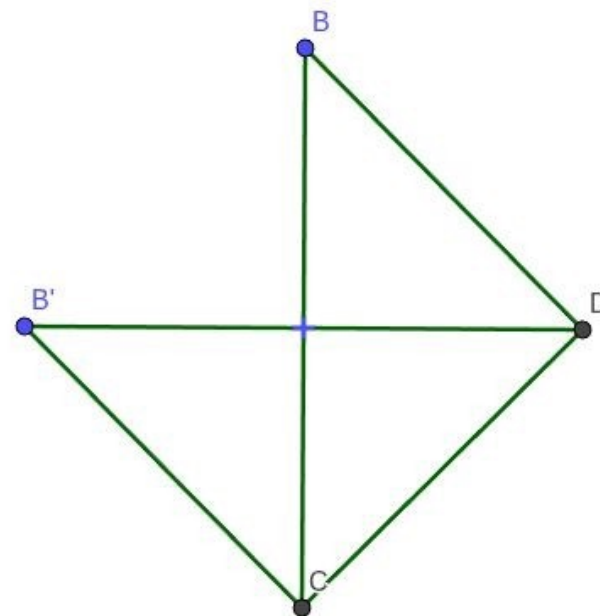
Adjacency Matrix

It is a $n \times n$ matrix in which:

- n is the number of vertices
- $A(i, j) = 1$ only if there is a link from i to j and
- $A(i, j) = 0$ if not

Adjacency Matrix:

$$\begin{pmatrix} 0 & 0 & 1 & 1 \\ 0 & 0 & 1 & 1 \\ 1 & 1 & 0 & 1 \\ 1 & 1 & 1 & 0 \end{pmatrix}$$



You can create your own Adjacency Matrix on Geogebra



Other common representation

- **Edge List**

An edge list is a collection of pairs (or tuples), where each pair (u, v) indicates that there is an edge connecting node u to node v

- **Adjacency List**

It is implemented as an array or map where the index/key represents a vertex, and the value is a list of all vertices adjacent to it (its neighbors).



Embeddings

Transform complex, high-dimensional data (like words or images) into dense, low-dimensional numerical vectors that machine learning models can easily process

- Word Embeddings: Convert words into vectors such that words with similar semantic meaning or that appear in similar contexts are positioned close together in the vector space
Example: GloVe, Word2Vec, BERT
- Image Embeddings: Compress the high-dimensional pixel data into a dense vector that encodes the image's visual features (objects, patterns, textures)
Example: CNN, Vision Transformer

These work well for unstructured or sequential data. **Fail to capture the defining characteristic of graph data: interconnectedness and complex relationships.**

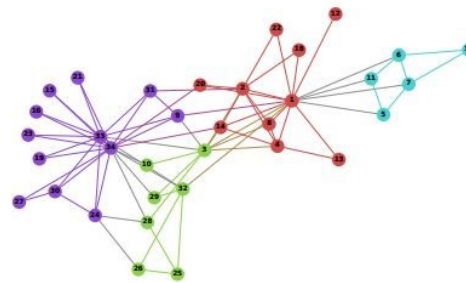
Graph Node Embeddings

To encode the topology and structural properties of the network. This includes:

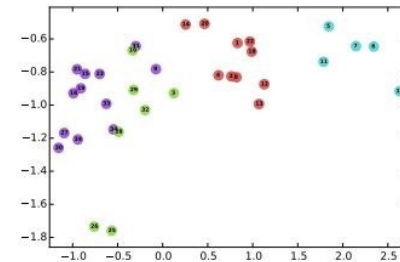
- node's own features
- Its local neighborhood (dense group of nodes within a larger graph structure)
- Its structural role

Adjacency Matrix are often huge and sparse

Graph embeddings condense this high-dimensional, sparse data into a low-dimensional, dense vector space, such that similar nodes in the graph are embedded close together.



(a) Input: Karate Graph



(b) Output: Representation

Image credit: DeepWalk: Online Learning of Social Representations

Three Core Graph Learning Tasks

- **Node Classification**
Predict label for each node (e.g., social influence detection)
- **Edge Prediction**
Predict if edges exist or their properties (e.g., friend suggestions)
- **Graph-Level Tasks**
Classify or predict properties of entire graphs (e.g., molecular properties)

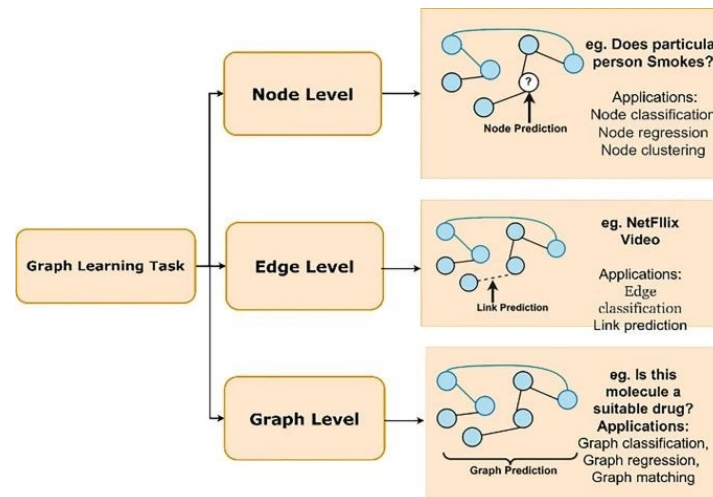


Image credit: Khemani et al. A review of graph neural networks



Graph Embeddings & MLPs??

This is a two-step process

- First, the embedding is generated
- Second, the MLP is trained on the fixed embedding
- cannot flow back to fine-tune the embedding generation process



Graph Embeddings & MLPs??

This is a two-step process

- First, the embedding is generated
- Second, the MLP is trained on the fixed embedding
- **cannot flow back to fine-tune the embedding generation process**

Other Problems:

- Re-train the entire embedding space from scratch every time you add a new node or a new edge to the graph
- They ignore or struggle to incorporate the rich node features

Solution?



Graph Embeddings & MLPs??

This is a two-step process

- First, the embedding is generated
- Second, the MLP is trained on the fixed embedding
- **cannot flow back to fine-tune the embedding generation process**

Other Problems:

- Re-train the entire embedding space from scratch every time you add a new node or a new edge to the graph
- They ignore or struggle to incorporate the rich node features

Solution?

Graph Neural Networks (GNNs)



Poll 2

True or False

Node embeddings aim to map nodes in a graph to a continuous vector space while preserving their structural and semantic properties.



Poll 2

True or False

Node embeddings aim to map nodes in a graph to a continuous vector space while preserving their structural and semantic properties. **(True)**



Graph Neural Networks

Carnegie Mellon University

Graph Neural Networks (GNNs)

Designed to handle irregular connectivity and the lack of inherent order for nodes.

The core mechanism **Message Passing**

- **Gather**
- **Aggregate**
- **Update**

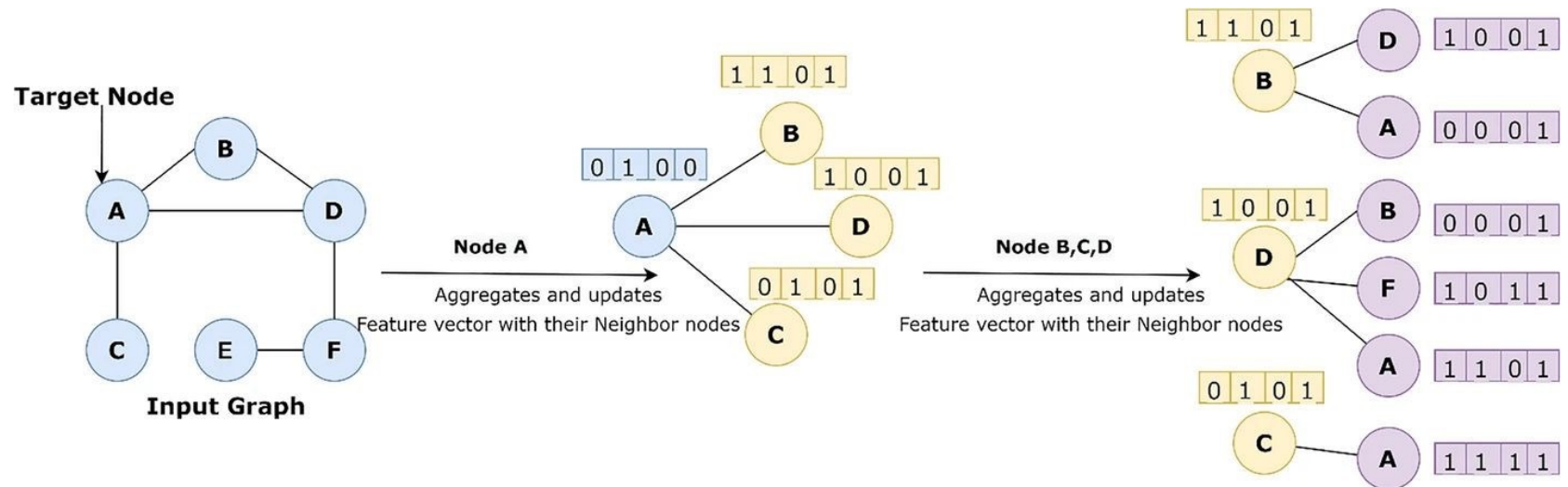


Image credit: Khemani et al. A review of graph neural networks



Graph Neural Networks (GNNs)

The Neighbourhood N_i of a node i is defined as the set of nodes j connected to i by an edge

Gather:

Each node will create a message, which will be sent to other nodes later

Example: A Linear layer $\mathbf{m}_u^{(l)} = \mathbf{W}^{(l)} \mathbf{h}_u^{(l-1)}$

Aggregate:

Node v will aggregate the messages from its neighbors $u \in N(v)$ and combines them into a single summary vector $\mathbf{h}_v^{(l)}$

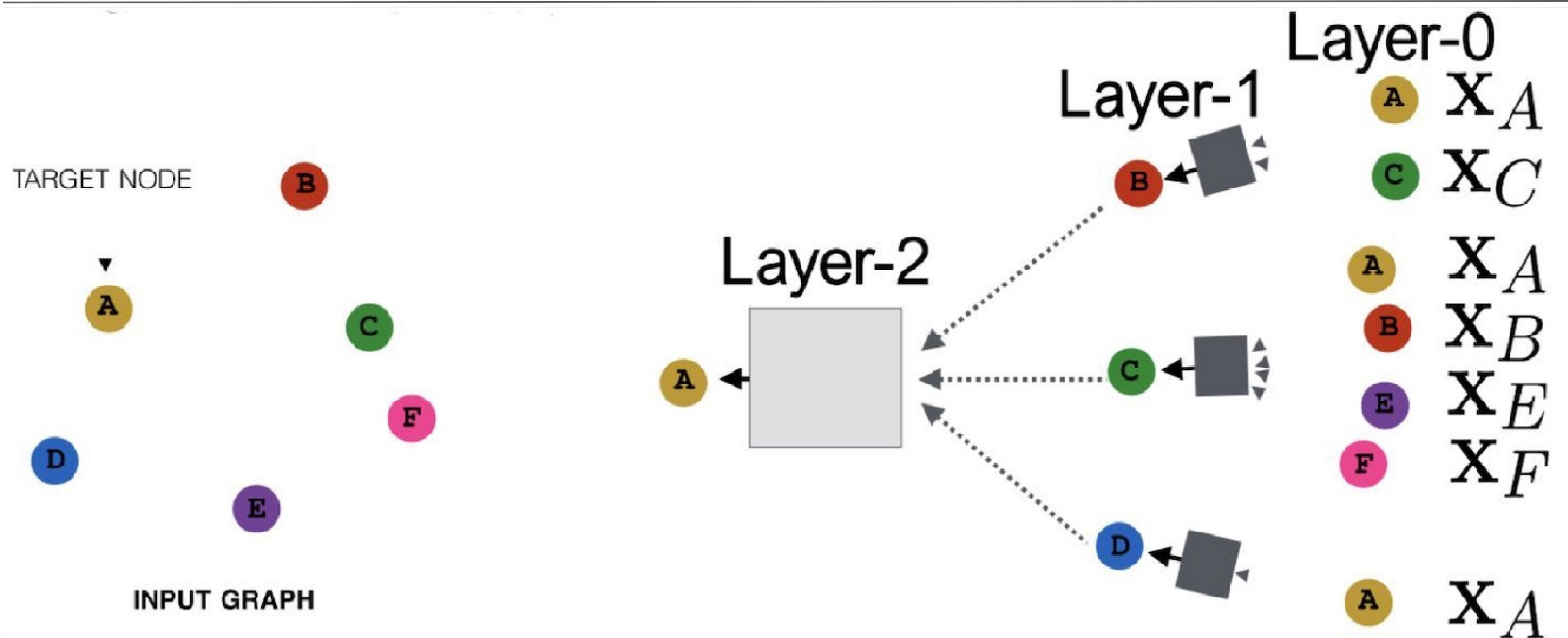
$$\mathbf{h}_v^{(l)} = \text{AGG}^{(l)} \left(\left\{ \mathbf{m}_u^{(l)}, u \in N(v) \right\} \right)$$

Node Update:

Node v combines the aggregated message $\mathbf{h}_v^{(l)}$ with its own feature vector from the previous layer $\mathbf{h}_v^{(l-1)}$. Easier to include $\mathbf{h}_v^{(l-1)}$ when computing $\mathbf{h}_v^{(l)}$.

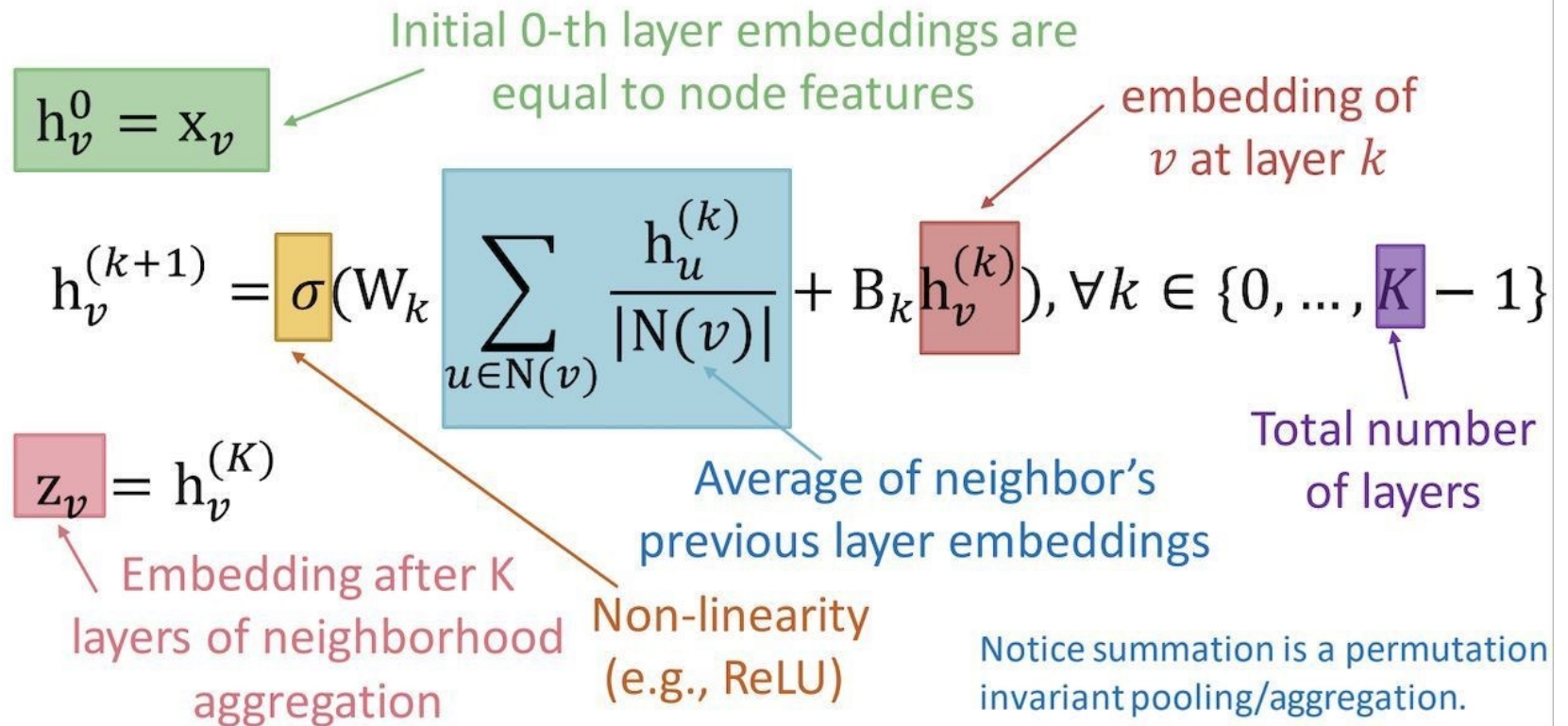
This process is repeated L times (where L is the number of GNN layers) to capture features from the L -hop neighborhood.

Graph Neural Networks (GNNs)



Jure Leskovec, Stanford CS224W: Machine Learning with Graphs, <http://cs224w.stanford.edu>

Graph Neural Networks (GNNs)



Jure Leskovec, Stanford CS224W: Machine Learning with Graphs, <http://cs224w.stanford.edu>



Aggregation Functions & Information Propagation

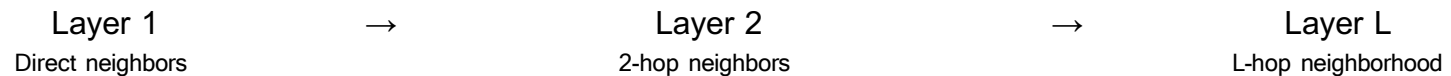
Aggregation function must be permutation-invariant. Commonly used are - Sum, Mean/Average, Max Pooling & Attention

Sum Aggregation: $h = \sum m$ (most common, permutation invariant)

Mean Aggregation: $h = (1/|N|) \sum m$ (normalized version)

Max Aggregation: $h = \max(m)$ (captures outliers)

Information Propagation



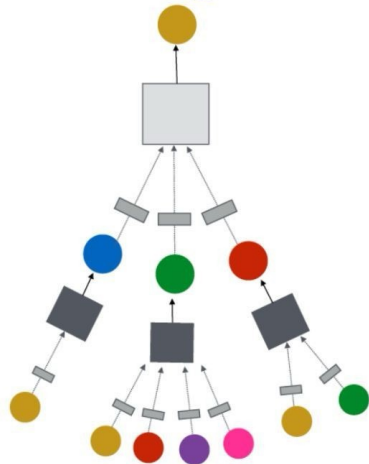
Typical: 2-4 layers (diminishing returns & over-smoothing)

Jure Leskovec, Stanford CS224W: Machine Learning with Graphs, <http://cs224w.stanford.edu>

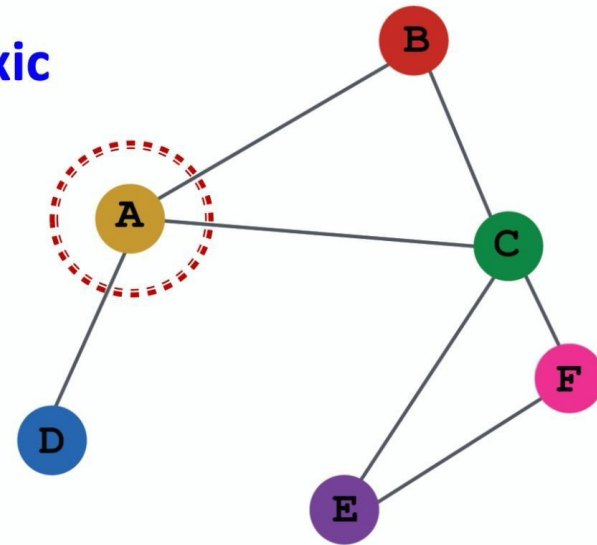
GNN Model Training

Directly train the model for a supervised task (e.g., node classification)

Safe or toxic
drug?



Safe or toxic
drug?



E.g., a drug-drug
interaction network

Jure Leskovec, Stanford CS224W: Machine Learning with Graphs, <http://cs224w.stanford.edu>

GNN Model Training

Use cross entropy loss for Node Classification

$$\mathcal{L} = - \sum_{v \in V} y_v \log(\sigma(z_v^T \theta)) + (1 - y_v) \log(1 - \sigma(z_v^T \theta))$$

Encoder output: node embedding

Classification weights

Node class label

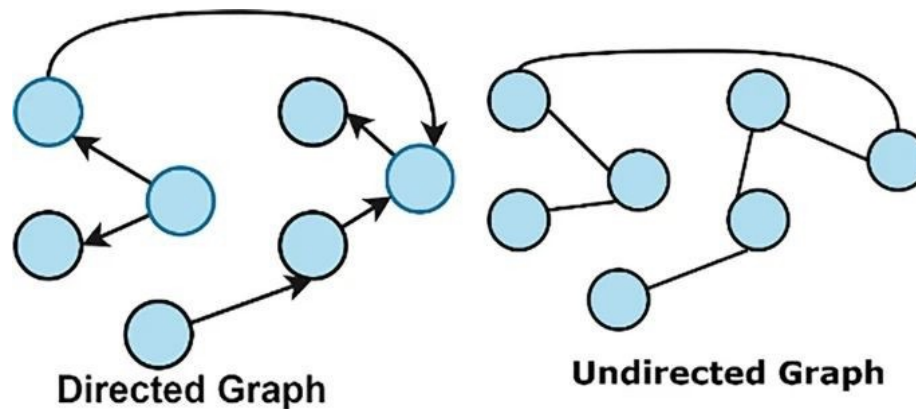
Safe or toxic drug?

What Types of Graphs can we train?

Directed/undirected graphs:

Focus is more on the flow of information. Standard Graph Convolutional Networks (GCNs), GraphSAGE is typically implemented for undirected graphs

For directed graphs a node v will only aggregate the messages from its neighbors $u \in N^-(v)$ (nodes u such that $u \rightarrow v$)



Are GNNs enough?

The first GNN paper proposed a recurrent approach that required iteration to a fixed point.

The Issue: Achieving convergence takes a significant, variable number of steps and requires backpropagation through a potentially very long sequence of updates.

This made the original GNNs computationally expensive and slow to train, especially on large graphs.

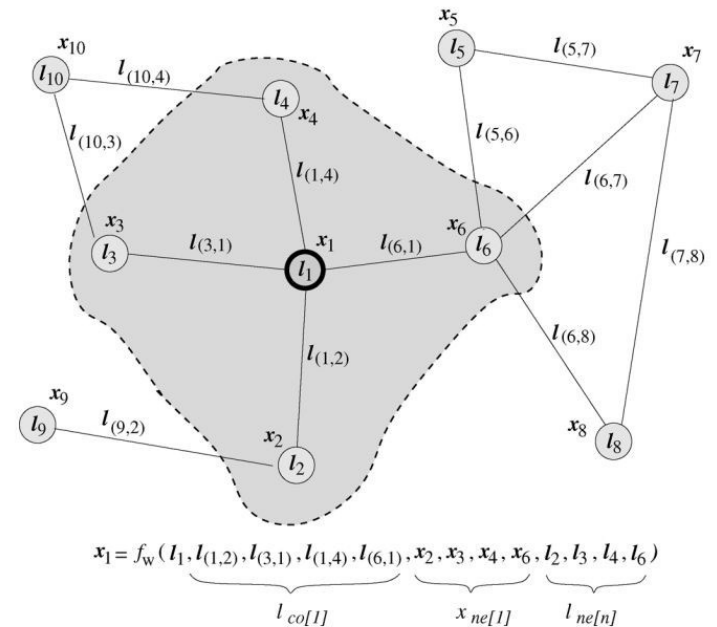


Fig. 2. Graph and the neighborhood of a node. The state x_1 of the node 1 depends on the information contained in its neighborhood.

Are GNNs enough?

The first GNN paper proposed a recurrent approach that required iteration to a fixed point.

The Issue: Achieving convergence takes a significant, variable number of steps and requires backpropagation through a potentially very long sequence of updates.

This made the original GNNs computationally expensive and slow to train, especially on large graphs.

WAIT! Our MLPs also faced similar issue!!
What was the solution??

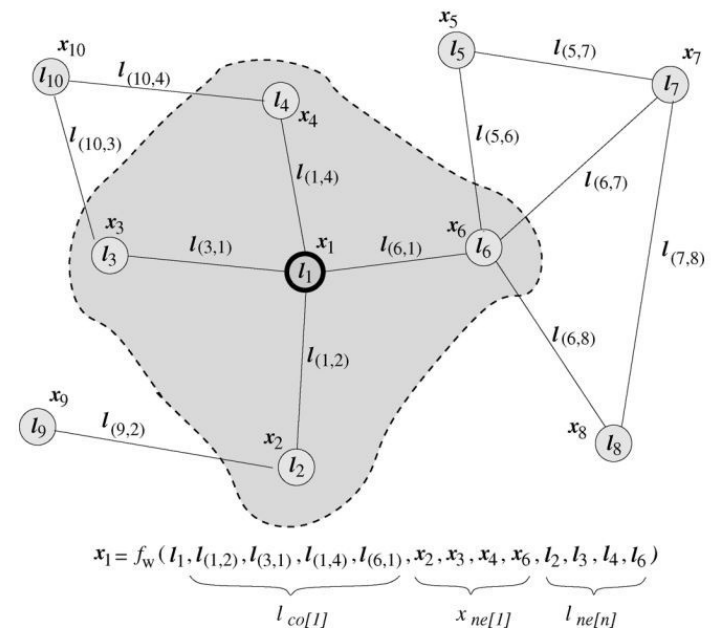


Fig. 2. Graph and the neighborhood of a node. The state x_1 of the node 1 depends on the information contained in its neighborhood.

Are GNNs enough?

The first GNN paper proposed a recurrent approach that required iteration to a fixed point.

The Issue: Achieving convergence takes a significant, variable number of steps and requires backpropagation through a potentially very long sequence of updates.

This made the original GNNs computationally expensive and slow to train, especially on large graphs.

WAIT! Our MLPs also faced similar issue!!
What was the solution??

CONVOLUTION!!

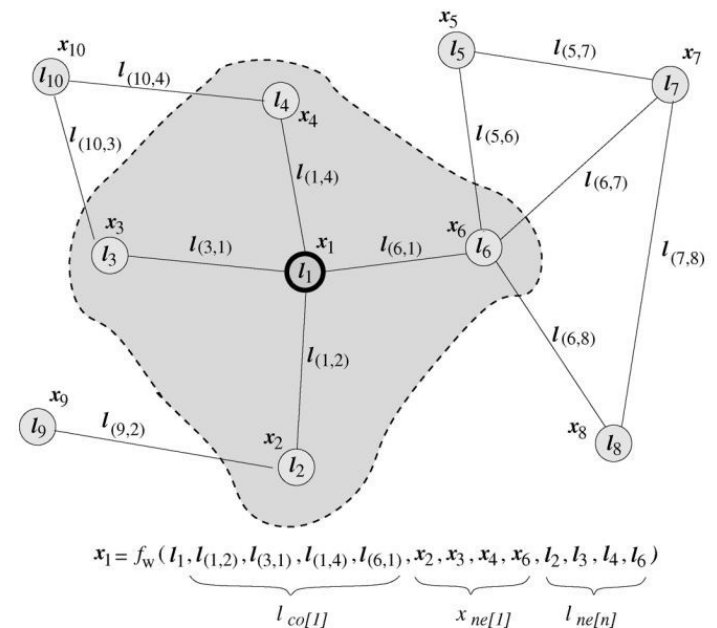


Fig. 2. Graph and the neighborhood of a node. The state x_1 of the node 1 depends on the information contained in its neighborhood.

What was the benefits of Convolution?

Extracts Local Patterns: Filters learn to recognize features by looking at a fixed, local neighborhood of pixels

Achieves Parameter Sharing: The same filter (weights) is applied across the entire image, making the model highly efficient and scalable

Is Translation Invariant: The model can recognize a feature regardless of where it appears in the image

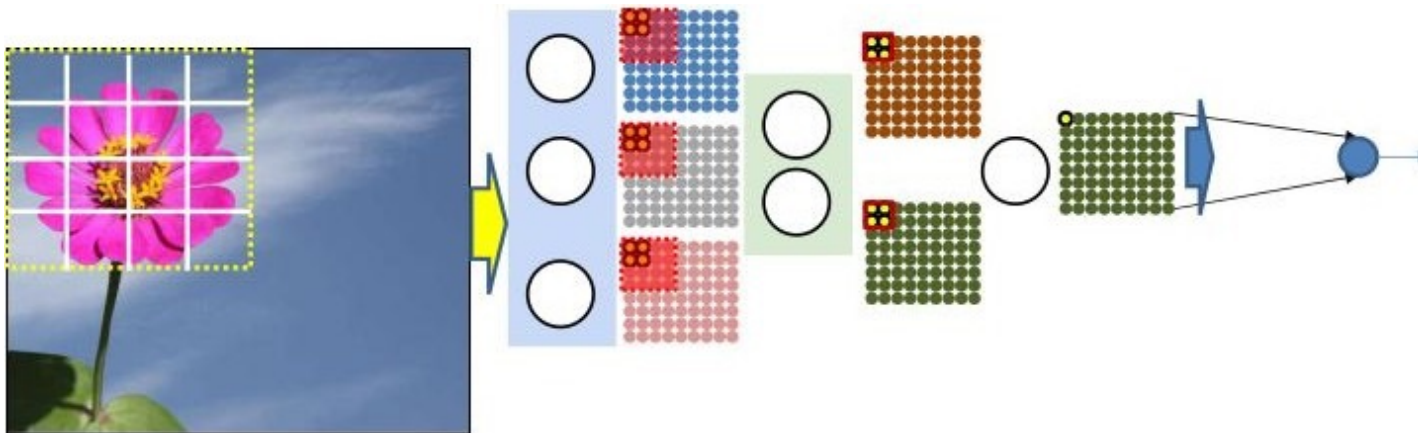


Image credit: 11-785, Prof. Bhiksha

Graph Convolutional Networks (GCN)

Foundational GNN architecture (Kipf & Welling 2017, Semi-supervised Classification with Graph Convolutional Networks)

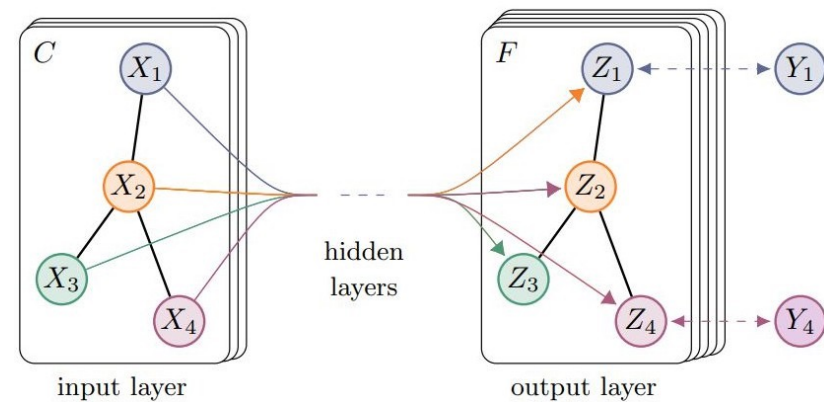
Update rule:

$$h = \sigma(W * \sum (\hat{A}[i,j] * h))$$

$\hat{A}$ = normalized adjacency matrix

W = learnable weight matrix

σ = ReLU activation



Graph Convolutional Networks (GCN)

Foundational GNN architecture (Kipf & Welling 2017, Semi-supervised Classification with Graph Convolutional Networks)

Update rule:

$$h = \sigma(W * \sum (\hat{A}[i,j] * h))$$

$\hat{A}$ = normalized adjacency matrix

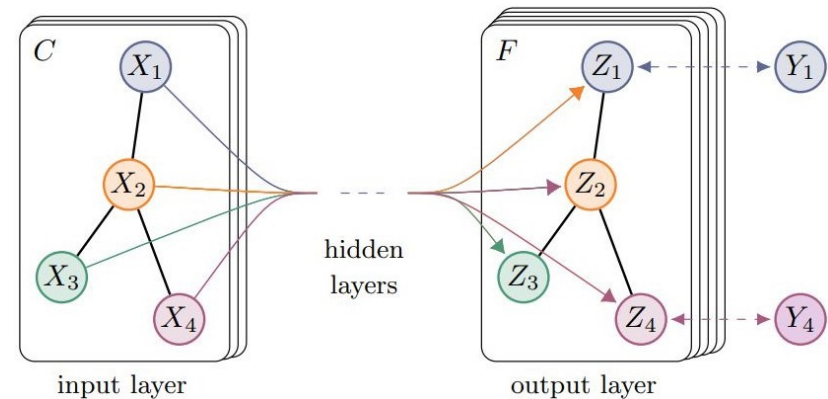
W = learnable weight matrix

σ = ReLU activation

This introduced scalable GCN model that uses the matrix multiplication formula to implement the message-passing (or convolution) step.

Solves convergence and complexity issues.

Pushed GNNs into the mainstream.



GraphSAGE

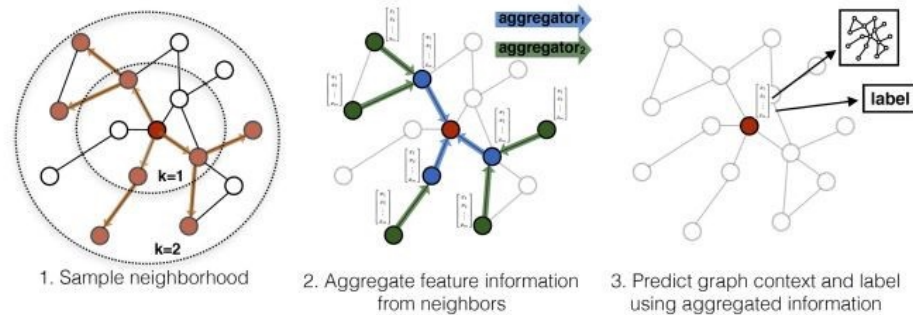
Inductive learning through sampling (Hamilton et al. 2017)

Sample k random neighbors (not all)

Aggregate their embeddings

Learn inductive functions for new nodes

Key advantage: Works on unseen graphs, scalable





GraphSAGE

Algorithm 1: GraphSAGE embedding generation (i.e., forward propagation) algorithm

Input : Graph $\mathcal{G}(\mathcal{V}, \mathcal{E})$; input features $\{\mathbf{x}_v, \forall v \in \mathcal{V}\}$; depth K ; weight matrices $\mathbf{W}^k, \forall k \in \{1, \dots, K\}$; non-linearity σ ; differentiable aggregator functions $\text{AGGREGATE}_k, \forall k \in \{1, \dots, K\}$; neighborhood function $\mathcal{N} : v \rightarrow 2^{\mathcal{V}}$

Output : Vector representations $\mathbf{z}_v$ for all $v \in \mathcal{V}$

```
1  $\mathbf{h}_v^0 \leftarrow \mathbf{x}_v, \forall v \in \mathcal{V}$  ;
2 for  $k = 1 \dots K$  do
3   for  $v \in \mathcal{V}$  do
4      $\mathbf{h}_{\mathcal{N}(v)}^k \leftarrow \text{AGGREGATE}_k(\{\mathbf{h}_u^{k-1}, \forall u \in \mathcal{N}(v)\})$ ;
5      $\mathbf{h}_v^k \leftarrow \sigma \left( \mathbf{W}^k \cdot \text{CONCAT}(\mathbf{h}_v^{k-1}, \mathbf{h}_{\mathcal{N}(v)}^k) \right)$ 
6   end
7    $\mathbf{h}_v^k \leftarrow \mathbf{h}_v^k / \|\mathbf{h}_v^k\|_2, \forall v \in \mathcal{V}$ 
8 end
9  $\mathbf{z}_v \leftarrow \mathbf{h}_v^K, \forall v \in \mathcal{V}$ 
```

Image credit: Hamilton et al., Inductive Representation Learning on Large Graphs



Poll 3

True or False

A Graph Neural Network (GNN) using graph convolution can still be trained for edge-level prediction even if there is no information in the nodes



Poll 3

True or False

A Graph Neural Network (GNN) using graph convolution can still be trained for edge-level prediction even if there is no information in the nodes **(True)**



Issues with GCNs

The standard GCN uses a predefined, isotropic (uniform) aggregation function. This means that when a node receives messages from its neighbors, it treats all neighbors equally

Issues with GCNs

The standard GCN uses a predefined, isotropic (uniform) aggregation function. This means that when a node receives messages from its neighbors, it treats all neighbors equally **irrespective of their importance, feature relevance, or the nature of their connection.**

GCN update rule

$$\mathbf{H}^{(l+1)} = \sigma \left(\tilde{\mathbf{D}}^{-\frac{1}{2}} \tilde{\mathbf{A}} \tilde{\mathbf{D}}^{-\frac{1}{2}} \mathbf{H}^{(l)} \mathbf{W}^{(l)} \right)$$

Diagram annotations: A red box highlights the normalization term $\tilde{\mathbf{D}}^{-\frac{1}{2}} \tilde{\mathbf{A}} \tilde{\mathbf{D}}^{-\frac{1}{2}}$ with an arrow pointing to the label "Normalization Term". A blue box highlights the feature transformation term $\mathbf{H}^{(l)} \mathbf{W}^{(l)}$ with an arrow pointing to the label "Feature Transformation".

Before aggregation, the current node features are transformed by a learnable weight matrix (shared by all nodes)

The normalized adjacency matrix ($\hat{\mathbf{A}}$) is the "convolution" part

SO WE NEED "Attention"

Graph Attention Networks (GAT)

Learn adaptive neighbor weights (Veličković et al. 2018)

Attention mechanism:

$$\alpha = \text{softmax}(\text{LeakyReLU}(a[Wh||Wh]))$$

Each node learns which neighbors are important

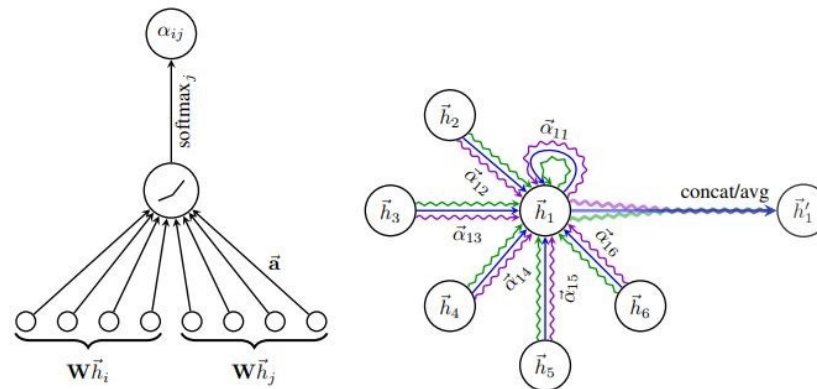


Image credit: Veličković et al., GRAPH ATTENTION NETWORKS



Training & Optimization

Layer depth: 2-4 layers optimal (over-smoothing at depth)

Pooling: Global sum/mean for graph-level tasks

Batching: Mini-batch graphs or node sampling

Regularization: Dropout, weight decay, batch norm

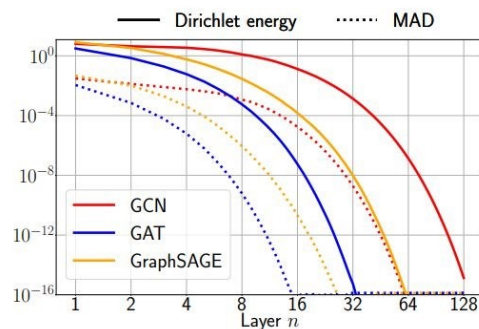
The Over-Smoothing Problem

As layers increase, node embeddings converge to similar values

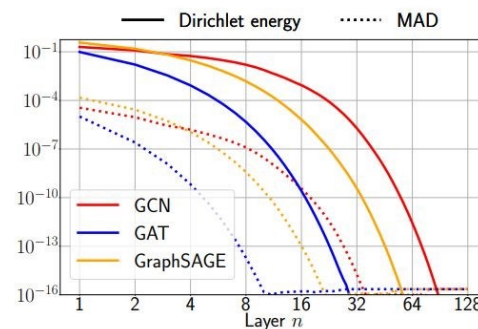
Problem:

Deep GNNs lose expressiveness. Nodes become indistinguishable.

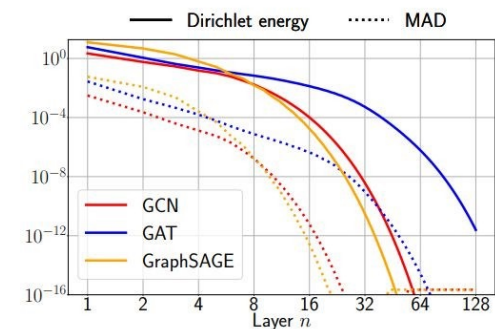
This is why 2-4 layers is standard practice



small-scale Texas graph



medium-scale Cora citation network



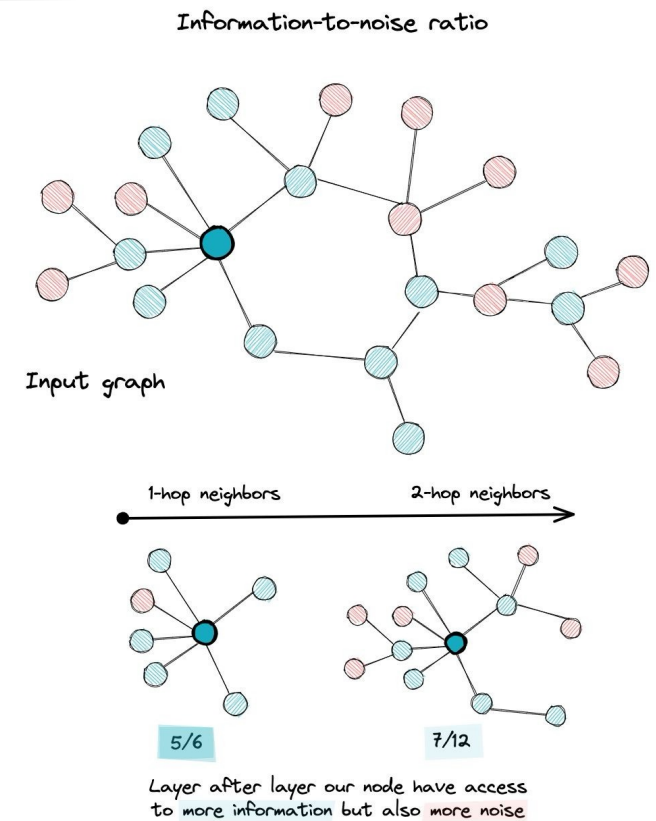
large-scale Facebook network
(Cornell5)

Image credit: Rusch et al., A SURVEY ON OVERSMOOTHING IN GRAPH NEURAL NETWORKS

The Over-Smoothing Problem

This generalization was built on the main hypothesis that while nodes are interacting, they have **either access to important information** from nodes from the same class **or noise** by interacting with nodes from other classes.

While node access to more parts of the graph we may access to noisy nodes that affect the final embedding.



Credit: Aomar Anas, Over smoothing issue in graph neural network, towards datascience

Carnegie Mellon University



GAT's Limitation

- Like the GCN, the GAT still operates within the **confines of the Message Passing paradigm** i.e aggregates information from a node's immediate neighbors (1-hop)
- Calculates **weights based on node features and local connectivity**, it doesn't explicitly encode the full positional or structural relationship between two nodes that are far apart
- Attention is calculated only on a **subset of features** (neighbors), limiting the complexity of the feature interactions

To capture global patterns or long-range dependencies, you must stack many layers (L).

Leads to the over-smoothing problem



GAT's Limitation

- Like the GCN, the GAT still operates within the **confines of the Message Passing paradigm** i.e aggregates information from a node's immediate neighbors (1-hop)
- Calculates **weights based on node features and local connectivity**, it doesn't explicitly encode the full positional or structural relationship between two nodes that are far apart
- Attention is calculated only on a **subset of features** (neighbors), limiting the complexity of the feature interactions

To capture global patterns or long-range dependencies, you must stack many layers (L).

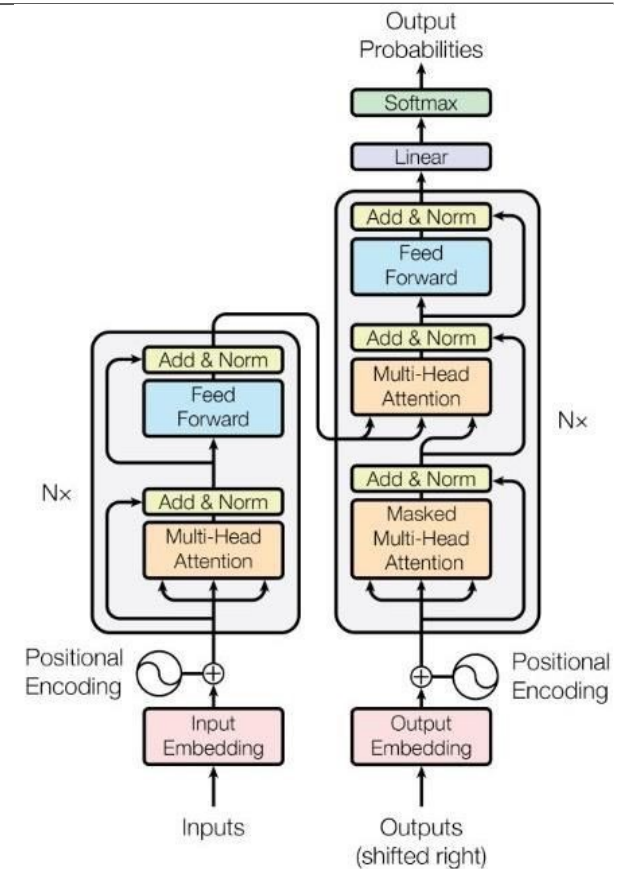
Leads to the over-smoothing problem

SO WE NEED GLOBAL ATTENTION!

Transformer??

Recap Transformers

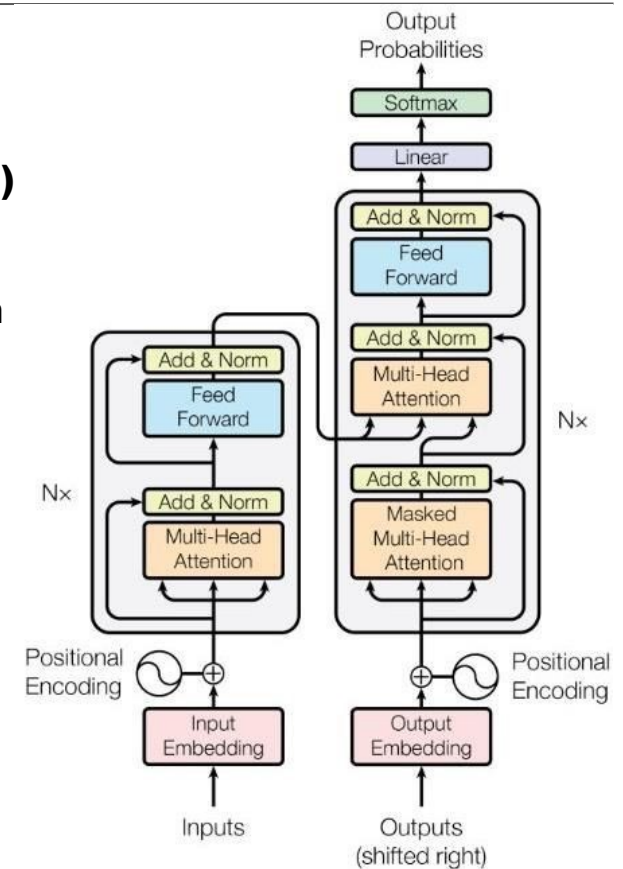
- Self-Attention is used to weigh the importance of all other words in the sequence when processing one word.
- Compute attention scores using a function of Q(Query), K (Key), and V (Value) matrices, typically via the Scaled Dot-Product Attention
- Use a Feed-Forward Network block after the attention mechanism to enhance expressiveness.



Recap Transformers

- **Self-Attention** is used to weigh the importance of all other words in the sequence when processing one word.
- Compute attention scores using a function of **Q(Query)**, **K(Key)**, and **V(Value)** matrices, typically via the Scaled Dot-Product Attention
- Use a **Feed-Forward Network** block after the attention mechanism to enhance expressiveness.

The key mechanisms are the same for Graph Transformer

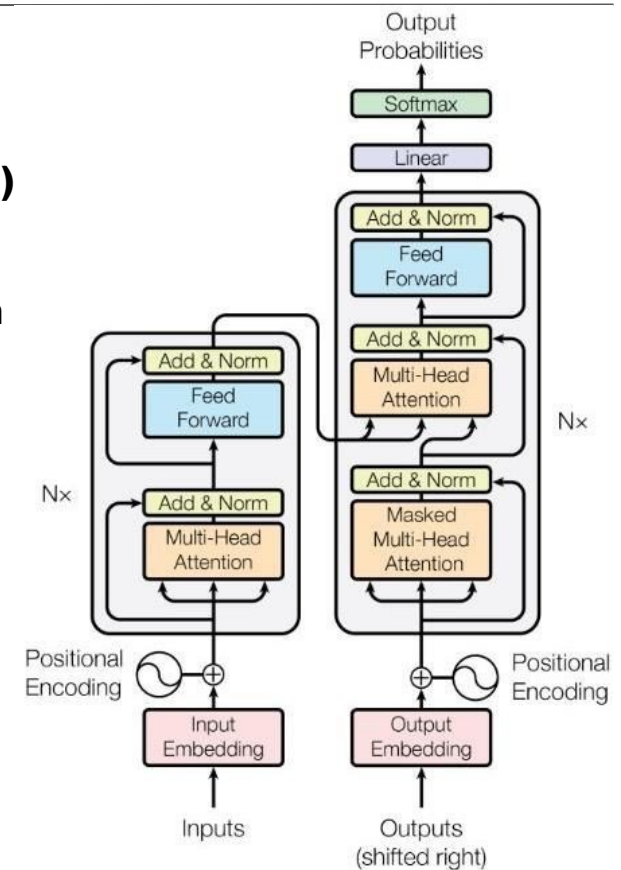


Recap Transformers

- **Self-Attention** is used to weigh the importance of all other words in the sequence when processing one word.
- Compute attention scores using a function of **Q(Query)**, **K(Key)**, and **V(Value)** matrices, typically via the Scaled Dot-Product Attention
- Use a **Feed-Forward Network** block after the attention mechanism to enhance expressiveness.

The key mechanisms are the same for Graph Transformer

Any Differences??





Transformer(NLP) vs Graph Transformer(GNN) Differences

The differences arise from the need to adapt the sequence-based operations of the original Transformer to the irregular, non-sequential structure of a graph

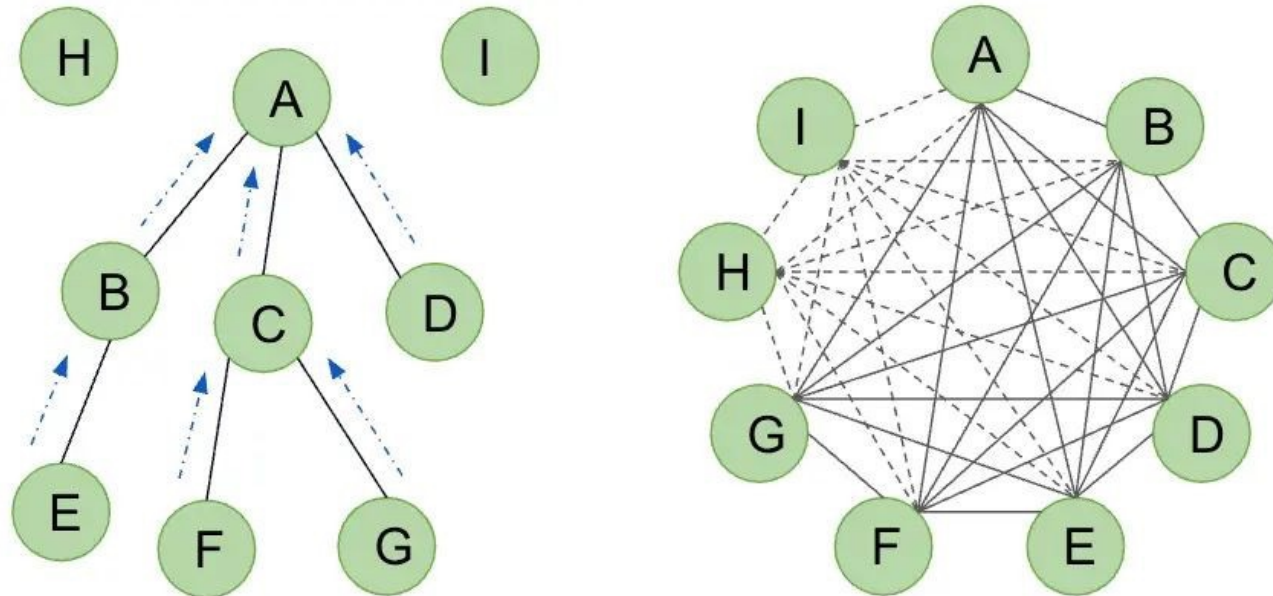
Attention Scope: Customizable for Local Attention (similar to GAT) or Global Attention

Positional Encoding: Cannot use sinusoidal function or learnable vector. No sequential order.
Must incorporate Graph Structure Encoding

Structural Encoding: Explicitly models graph properties

Masking: Cannot use Causal Masking. Adjacency Masking is an option

Graph Transformer(GNN) Information Flow



Left: GNN with local message passing. Disconnected nodes like “H” and “I” do not influence “A.”

Right: Graph Transformer with global attention. “A” can attend to all nodes, capturing long-range dependencies.

Credit: Lopez,Fey,Leskovec. "An Introduction to Graph Transformers" - kumo.ai

Graph Transformers

Full attention over all nodes (no just neighbors)

Spectral Transformers: Learnable PE (positional encoding)

Graph-scaled attention: Efficiency via subgraph sampling

- SOTA on many benchmarks

Blurs lines between local and global information

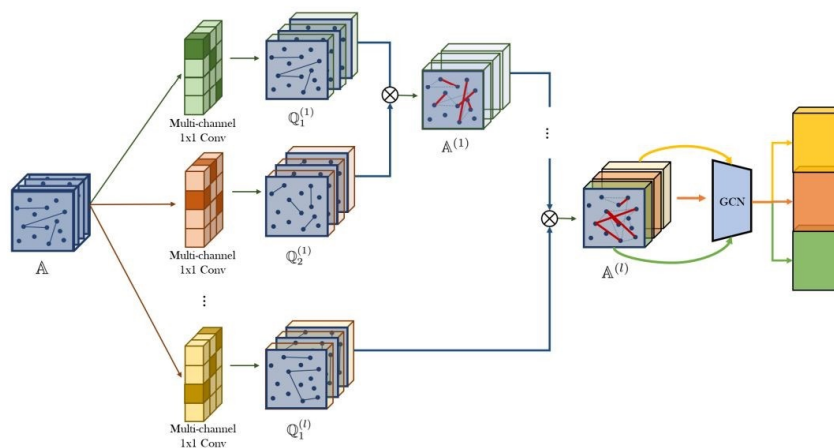


Image credit: Yun et al., Graph Transformer Networks



Confused??

GNNs vs GCNs vs GAT vs GT

Confused??

GNNs vs GCNs vs GAT vs GT

Forgot what they stand for?



Carnegie Mellon University

Confused??

**Graph Neural Network(GNNs) vs
Graph Convolutional Networks(GCNs) vs
Graph Attention Networks(GAT) vs
Graph Transformer(GT)**



Carnegie Mellon University



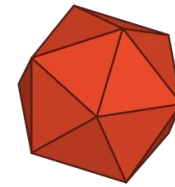
Summary

Architecture	Pros	Cons	When to Use
GNN	Message-passing paradigm.	Not a model	Parent of all Graph model!
GCN	Computationally fast	Treats all neighbors equally	For a quick, efficient solution on medium-sized, simple graphs
GAT	Assigns different, learned importance scores to neighbors using attention	High Compute Cost. Still limited to local aggregation. risk over smoothing in deep models	High-Relational Data (social networks or knowledge graphs)
Graph Transformer	Long-range dependencies, structural encoding No over-smoothing	Highest Compute Cost. Attention $O(N^2)$ for N nodes. Complicated	Long-range dependencies are critical. Prioritize expressiveness over speed. Eg (protein folding)

Ecosystems for GNNs

PyTorch Geometric

Industry standard, rich models, production-ready



PyTorch
geometric

DGL

Framework-agnostic, optimized for large graphs



NetworkX

Data preprocessing, analysis, visualization



NetworkX
Network Analysis in Python

Image credit: pytorch-geometric, dgl.ai, networkx.org

Carnegie Mellon University

NetworkX Example

A temporal, directed knowledge graph illustrating a person's health triage information. Arrows show progression from symptoms to diagnostic tests and subsequent treatment cycles.

Temporal Knowledge Graph Example

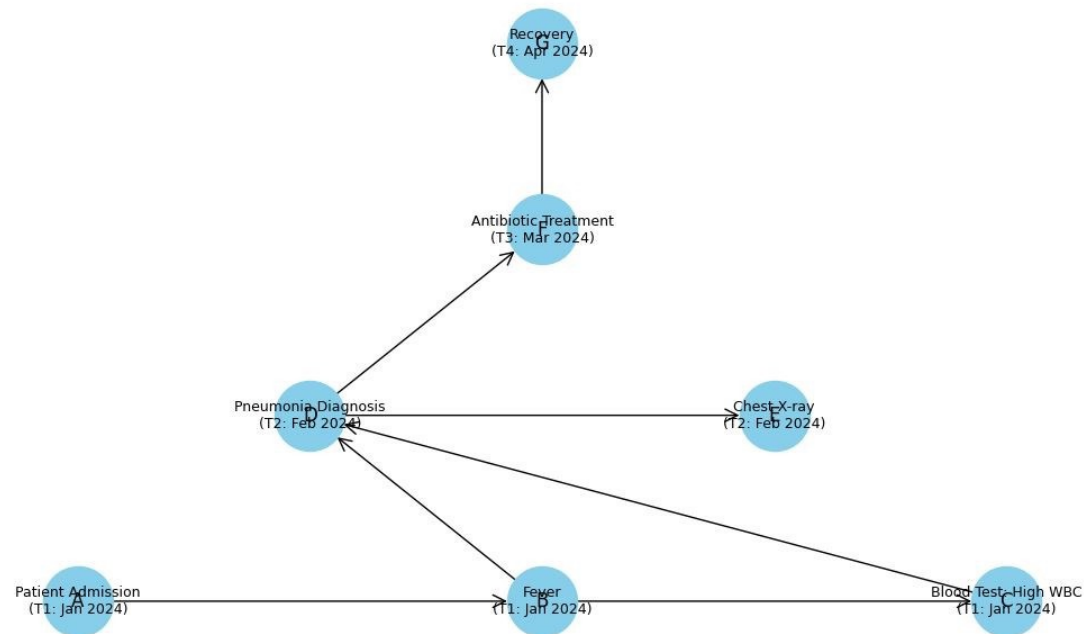


Image credit: Rutvik Joshi, 11785 Fall 25 TA



Going Deeper

Advanced Topics to Explore

Carnegie Mellon University

Expressivity: Weisfeiler-Lehman Test

Not all GNNs can distinguish all graph structures

WL-test: Iterative node coloring algorithm

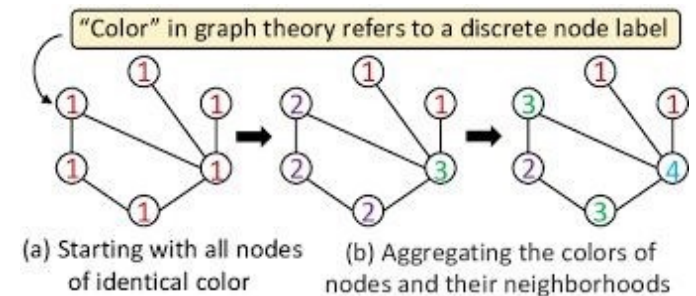
GIN (Graph Isomorphism Networks): Matches WL power

Higher-order GNNs: Use subgraph patterns

Current research: GNNs beyond message passing

Two iterations of the Weisfeiler-Lehman test on an example graph, starting with nodes of identical color/feature 1. The test performs graph coloring by first aggregating the colors of nodes and their neighborhoods and then generating unique new colors. The colors embed the structural roles of vertices in the graph. For simplicity, we assign the same initial color for all the nodes. However, if each node is assigned a unique feature vector that captures other non-structural properties, such information will also be accounted for when labeling the nodes.

Credit: Alrahis et al., Graph Neural Networks for Reverse Engineering of Gate-Level Netlists



Carnegie Mellon University

Equivariance Groups

- S_n , Permutation Group: Reordering elements (like node indices) without changing the underlying structure or connections.
- $T(3)$, Translation Group: Shifting all points by a constant vector in space without rotating or distorting the shape.
- $SO(3)$, Special Orthogonal Group: Rotating objects around a fixed point (origin) while preserving their orientation and shape.
- $SE(3)$, Special Euclidean Group: Rigid body motions that combine both rotation and translation, preserving relative distances and orientation (chirality).
- $E(3)$, Euclidean Group: The full set of spatial transformations including rotation, translation, and reflection (mirroring), preserving distances but not necessarily orientation.

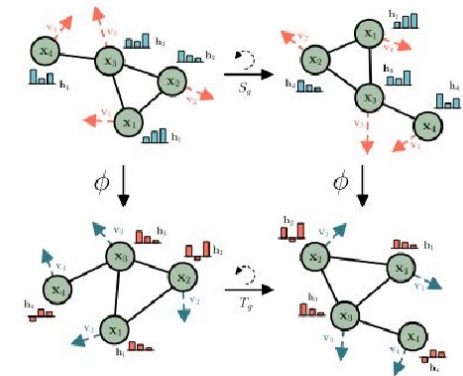


Image credit: Satorras et al., $E(n)$ Equivariant Graph Neural Networks



Equivariant & Geometric GNNs

For molecules and 3D structures, symmetries matter

SE(3)-equivariant networks: Rotation invariant

SchNet, NequIP: Molecular geometries

Enables: Drug discovery, material design

Ready for Some more Concepts??



Carnegie Mellon University

Hierarchical Graph Pooling

Coarsen graphs like CNNs pool spatial dimension

DiffPool: Differentiable pooling with assignment matrix

TopK pooling: Select important nodes

Benefit: Multi-scale graph representations

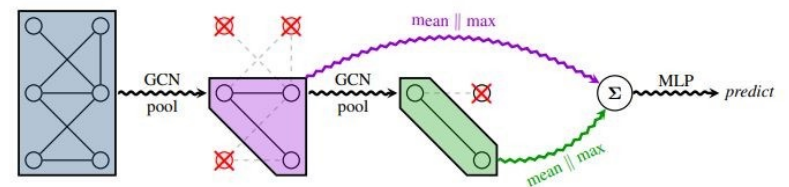
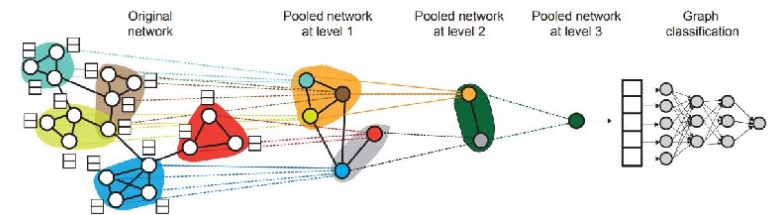


Image credit: Ying et al., Hierarchical Graph Representation Learning with Differentiable Pooling
Cangea et al., Towards Sparse Hierarchical Graph Classifiers



Spectral Graph Convolutions

- Signals on graphs: Laplacian eigenbasis approach
- Graph Fourier Transform: Eigendecomposition
- Chebyshev approximation: Efficient spatial filtering
- Connection: Spatial methods approximate spectral
- Theory-grounded signal processing perspective

Generative Models for Graphs

- Generate new graphs by learning structure and features
- Graph autoencoders: Encode $\rightarrow$ latent $\rightarrow$ decode
- Variational approaches: VAE for graphs
- Applications: Molecule generation, design



Key Takeaways

Message passing: Universal mechanism for graph learning

Architectures: GCN, GraphSAGE, GAT → easy to build variants

Practical: 2-4 layers, watch for over-smoothing

Frontier: Expressivity, geometry, generation still open



Thank you

Carnegie Mellon University



Open Research Directions

Scalability: Billion-node graphs, distributed training

Theory: Expressivity guarantees, universality results

Robustness: Adversarial attacks, certification, fairness