

Logic meets LLMs: Attention has a logic

Miguel Moreno
Tampere University
(joint work with Veeti Ahvonen, Damian Heiman,
Antti Kuusisto, and Matias Selin)

UN Encuentro 70

18 - 22 August, 2026

Motivation

Modern computational models (e.g. neural networks) have been proven very useful for a variety of problems.

In particular, the transformers architecture in the LLMs.

However, many of these models are only understood at a heuristic level.

To address this weakness we use logic to study their expressive power.

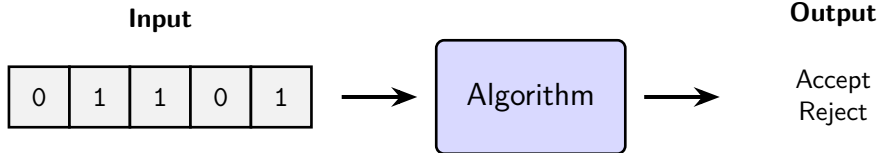
Big picture

One of the major goals in the theory of transformers is to answer a simple question:

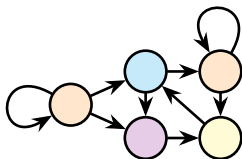
What can different transformer architectures compute?

We study the original encoder–decoder architecture, revealing a remarkably clean correspondence with temporal logic and distributed automata.

Typical sequential computing framework

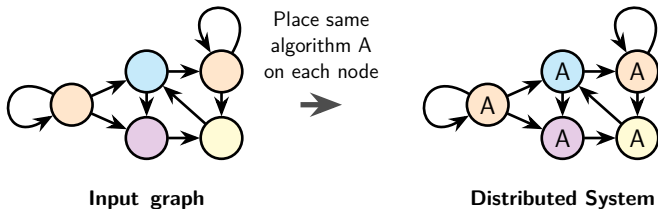


Distributed computing

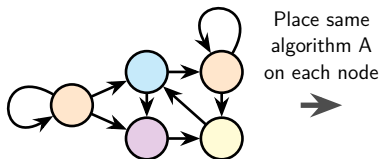


Input graph

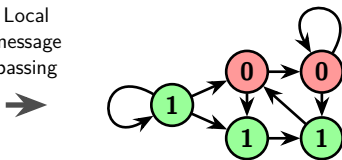
Distributed computing



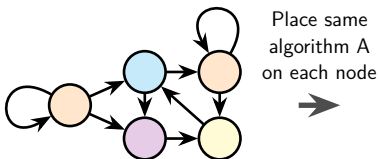
Distributed computing



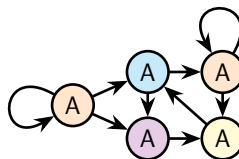
Local message passing



Distributed computing

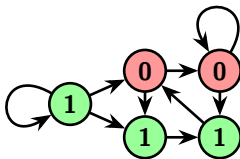


Input graph

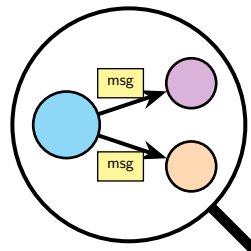


Distributed System

Local
message
passing



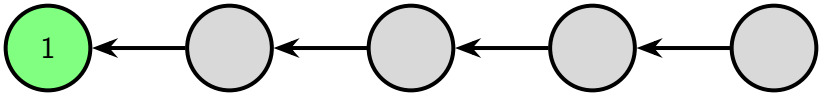
Local Outputs



Local View

Example: Reachability

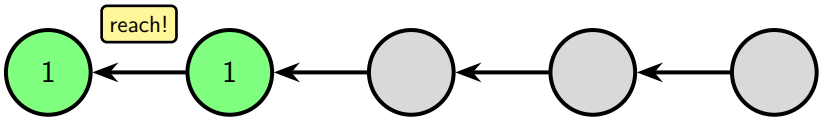
Initial State



Source (green) is reachable. Others are **unknown** (gray).

Example: Reachability

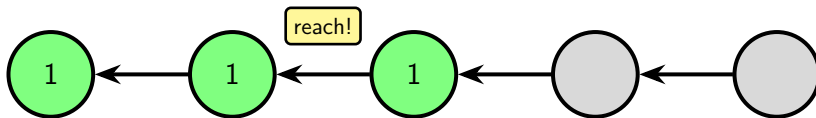
Round 1



Source sends “reach” — Node 2 receives from its out-neighbour.

Example: Reachability

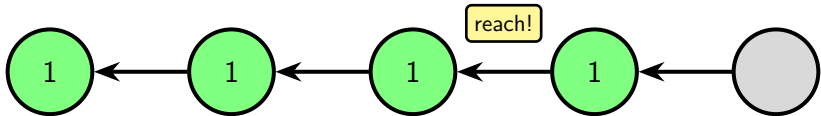
Round 2



Message propagates: each node receives from its out-neighbour.

Example: Reachability

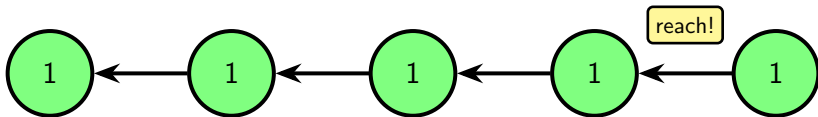
Round 3



Information continues toward the root.

Example: Reachability

Round 4



Root receives the message.

Distributed system

A distributed system is defined by a labelled directed graph $(W, R, p_1, \dots, p_k)$ together with an automaton A .

- ▶ Each node $w \in W$ contains a copy (A, w) of the automaton A .
- ▶ $R \subseteq W \times W$ is a collection of communication channels.
- ▶ Predicates $p_i \subseteq W$ encode a local input at each node.

The i -th input bit at node w is 1 iff $w \in p_i$.

Distributed automata

A **finite distributed automaton** (or FDA) is a tuple (Q, δ, π, F)

- ▶ Q is a finite non-empty set of states,
- ▶ $\delta: Q \times \mathcal{P}(Q) \rightarrow Q$ is a **transition function** that gives a new state to each node according to
 - ▶ the node's own current state
 - ▶ and the set of node's neighbours' states
- ▶ $\pi: \mathcal{P}(\{p_1, \dots, p_n\}) \rightarrow Q$ is an initialization function from node-labels to an initial state,
- ▶ $F \subseteq Q$ is a set of accepting states.

If Q is countably infinite, we simply call such an automaton a DA. A DA accepts a node of a graph if it enters an accepting state during the run.

Distributed automata counting

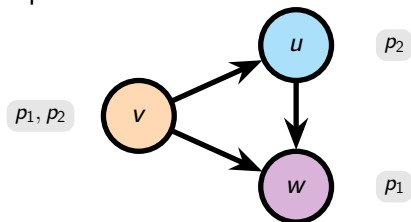
A **finite distributed automaton** (or FDA) is a tuple (Q, δ, π, F)

- ▶ Q is a finite non-empty set of states,
- ▶ $\delta: Q \times \mathcal{M}_k(Q) \times \mathcal{M}_k(Q) \rightarrow Q$ is a **transition function** that gives a new state to each node according to
 - ▶ the node's own current state
 - ▶ the k -capped multiset of node's neighbours' states
 - ▶ the k -capped multiset of all node's states
- ▶ $\pi: \mathcal{P}(\{p_1, \dots, p_n\}) \rightarrow Q$ is an initialization function from node-labels to an initial state,
- ▶ $F \subseteq Q$ is a set of accepting states.

If Q is countably infinite, we simply call such an automaton a DA.
A DA accepts a node of a graph if it enters an accepting state during the run.

Overview

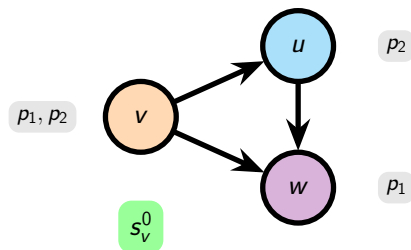
Formally a graph is a tuple $(V, E, p_1, \dots, p_n)$, where V is a set of nodes, E is a set of edges and node-labels $p_i \subseteq V$ encode a local input at each node.



Overview

Formally a graph is a tuple $(V, E, p_1, \dots, p_n)$, where V is a set of nodes, E is a set of edges and node-labels $p_i \subseteq V$ encode a local input at each node.

Round 0

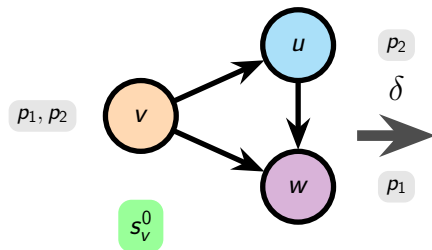


$$s_v^0 = \pi(p_1, p_2)$$

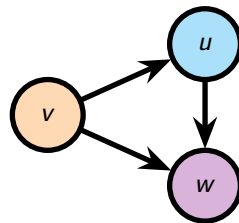
Overview

Formally a graph is a tuple $(V, E, p_1, \dots, p_n)$, where V is a set of nodes, E is a set of edges and node-labels $p_i \subseteq V$ encode a local input at each node.

Round 0



Round $t+1$

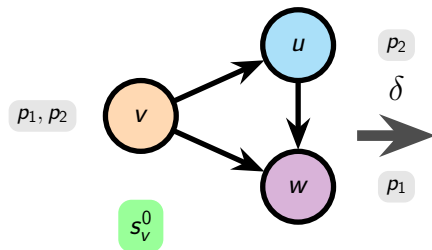


$$s_v^0 = \pi(p_1, p_2)$$

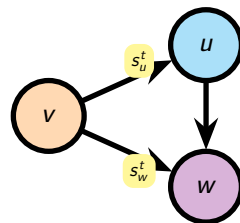
Overview

Formally a graph is a tuple $(V, E, p_1, \dots, p_n)$, where V is a set of nodes, E is a set of edges and node-labels $p_i \subseteq V$ encode a local input at each node.

Round 0



Round $t + 1$

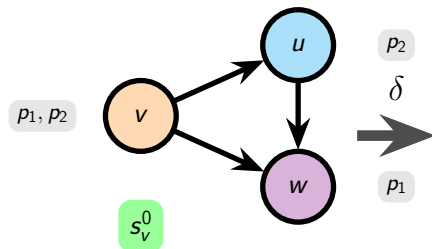


$$s_v^0 = \pi(p_1, p_2)$$

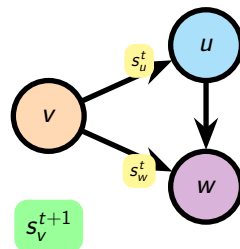
Overview

Formally a graph is a tuple $(V, E, p_1, \dots, p_n)$, where V is a set of nodes, E is a set of edges and node-labels $p_i \subseteq V$ encode a local input at each node.

Round 0



Round $t + 1$



$$s_v^0 = \pi(p_1, p_2)$$

$$s_v^{t+1} = \delta(s_v^t, \{s_u^t, s_w^t\})$$

Formally

Each communication round $t \in N$ defines a global configuration $S^t : V \rightarrow Q$ ($S_v^t = S^t(v)$).

S_v^0 := initial state at v obtained by the initialization function π .

Define Q' := the set of states realized by the neighbours of v in round t . Then S_v^{t+1} := the new state $\delta(S_v^t, Q')$ at v .

Strictly speaking, by “neighbours of v ” we here mean out-neighbours.

A node is accepted if it enters an accepting state during the run.

Related work on distributed computing

Hella et al. (PODC 2012) characterize **constant-iteration** distributed automata classes via modal logics.

Modal substitution calculus

A **program** of MSC consists of two lists of **rules**:

$$\overbrace{X_1(0) :- p \wedge q}^{\text{base rule}}$$

$$\vdots$$

$$X_n(0) :- \Diamond q \vee \neg q$$

$$\overbrace{X_1 :- X_3 \vee \Box p}^{\text{induction rule}}$$

$$\vdots$$

$$X_n :- \Diamond X_1$$

- ▶ $\varphi ::= \perp \mid p \mid \neg\varphi \mid \varphi \wedge \varphi \mid \Diamond\varphi$ (ordinary modal logic)
- ▶ $\varphi ::= \perp \mid p \mid X \mid \neg\varphi \mid \varphi \wedge \varphi \mid \Diamond\varphi$

Semantics

$$\begin{array}{ll} X_1(0) :- \psi_1 & X_1 :- \varphi_1 \\ \vdots & \vdots \\ X_n(0) :- \psi_n & X_n :- \varphi_n \end{array}$$

- ▶ $X_i^0 := \psi_i$
- ▶ X_i^{t+1} is obtained from the induction formula φ_i by substituting each X_j by X_j^t .

Semantics

$M, v \models \text{program}$ iff $M, v \models X_1^t$ for some $t \in \mathbb{N}$.

A program **accepts** a node v of a graph if for some $t \in \mathbb{N}$, the formula X_1^t is true at v .

Here $M = (V, R, p_1, \dots, p_k)$ is our distributed system and $v \in V$ a node.

Recall $M, v \models \Diamond \varphi$ iff $M, u \models \varphi$ for some u such that $R(v, u)$.

MSC and DAs

Theorem (Kuusisto CSL13)

MSC and finite distributed automata have the same expressive power.

Graph neural network

A **graph neural network** (or GNN) is a counting distributed automaton (δ, π, F) that is obtained by a **learning process**.

The initialization function $\pi: \mathcal{P}(\Pi) \rightarrow \mathbb{R}^d$ gives an initial feature vector s_v^0 for each node v .

The transition function δ is induced by

- ▶ a combination function $\text{COM}: \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}^d$
- ▶ an aggregation function $\text{AGG}: \text{mult}(\mathbb{R}^d) \rightarrow \mathbb{R}^d$ (typically sum)

In round $t = 1, 2, \dots$, every node v computes a new state

$$s_v^{t+1} := \text{COM}(s_v^t, \text{AGG}(\{\{s_u^t \mid (v, u) \in E\}\}))$$

Similarly to DAs, a GNN accepts a node v of a graph if in some round $t \in \mathbb{N}$, v visits an accepting feature vector.

GNN and GMSC

A GNN with **floating-point numbers** acts similarly to a GNN with real numbers.

GMSC, or **graded modal substitution calculus**, is obtained by allowing diamonds $\diamond^{\geq k}$ in MSC in addition to standard diamonds $\diamond$.

Theorem (Ahvonen, Heiman, Kuusisto, Lutz, 24)

The following have the same expressive power:

- ▶ GMSC
- ▶ floating-point GNNs

CPG Automaton

A **counting past-global distributed automaton** (CPG-automaton) is a tuple (Q, π, δ, n, b)

- ▶ $Q \subseteq \{0, 1\}^m$ is a finite non-empty set of states,
- ▶ $\delta: Q \times \mathcal{M}_k(Q) \times \mathcal{M}_k(Q) \rightarrow Q$ is a transition function, where $\mathcal{M}_k(Q)$ is the set of **k -capped multisets** over Q ,
- ▶ $\pi: \Sigma \rightarrow Q$ is an initialization function,
- ▶ $n \in \mathbb{N}$ is the output round,
- ▶ $b \in \mathbb{Z}_+$ is the output length.

GPTL⁻ Logic

The Σ -formulae of the temporal logic **GPTL⁻** are constructed by the grammar

$$\varphi ::= \perp \mid p \mid \neg\varphi \mid (\varphi \wedge \varphi) \mid \langle G \rangle_{\geq k}^{\text{pre}} \varphi \mid \langle P \rangle^{\text{suf}} \varphi,$$

Word-shaped graphs

Let Σ be a finite set of **tokens**.

A Σ -labelled word-shaped graph is a tuple (V, E, λ) :

- ▶ $V = [n]$
- ▶ $(i, j) \in E$ iff $j = i + 1$
- ▶ $\lambda : V \rightarrow \Sigma$ (labelling function)

Word-shaped graphs

A Σ -labelled **two-sorted word-shaped graph** is a $\Sigma \cup \{\text{BOS}\}$ -labelled word-shaped graph such that:

Only one vertex v has BOS assigned.

The **prefix** is the subgraph induced by $[v - 1]$.

The **suffix** is the subgraph induced by $[n] - [v - 1]$.

Masking and softmax

The **strict future mask** is the function on square matrices given by

$$\text{mask}(S)_{ij} := \begin{cases} -\infty & \text{when } j \geq i, \\ S_{ij} & \text{otherwise.} \end{cases}$$

Softmax is the function on vectors $\mathbf{x}$ given by

$$\text{softmax}(\mathbf{x})_i := \frac{e^{x_i - b_{\mathbf{x}}}}{\sum_{j=1}^n e^{x_j - b_{\mathbf{x}}}}.$$

Where $b_{\mathbf{x}} := \max\{x_1, \dots, x_{|\mathbf{x}|}\}$.

Attention

An **attention head** is a function defined by:

$$H_{\text{mask}}(X, Y) := \text{softmax} \left(\text{mask} \left(\frac{(Y W^Q)(X W^K)^{\top}}{\sqrt{d_k}} \right) \right) (X W^V).$$

Where W^Q , W^K , W^V are parameter matrices.

Attention

A **masked multi-head self-attention layer** is a function defined by:

$$\text{MHSA}_{\text{mask}}(X, X) := \text{concat}(H_{\text{mask}}^{(1)}(X, X), \dots, H_{\text{mask}}^{(h)}(X, X))W^O.$$

Where W^O is parameter matrix.

Multi-head cross-attention layers, **MHCA**, are defined similarly.

Multi-layer perceptrons

A **multi-layer perceptron** is a function defined by:

$$M(\mathbf{x}) := \text{ReLU}(\mathbf{x}W_1 + b_1)W_2 + b_2.$$

Where W_1 and W_2 are parameter matrices, and b_1 and b_2 are bias vectors.

ReLU is applied element-wise. M is applied to a matrix separately to each row.

Encoder

An **encoder layer** is a function given by:

$$\text{EL}(X) = M(X') + X'$$

where $X' = \text{MHSA}(X) + X$.

An **encoder of length** L is a composition of encoder layers:

$$E(X) = (\text{EL}_L \circ \dots \circ \text{EL}_1)(X).$$

Decoder

A **decoder layer** is a function given by:

$$DL(X, Y) = M(Y'') + Y''$$

where $Y'' = MHCA(X, Y') + Y'$ and $Y' = MHSA_{\text{mask}}(Y) + Y$.

A **decoder of length** L is a composition of decoder layers:

$$D(X, Y) = DL_L(X, DL_{L-1}(X, \dots DL_2(X, DL_1(X, Y)) \dots)).$$

Output

An **output head** is a function given by:

$$\text{Out}(Y) = Y \cdot W^{\text{out}} + b^{\text{out}},$$

where W^{out} and b^{out} are learned parameters.

Encoder–decoder transformer

An **encoder–decoder transformer** maps Σ -labelled two-sorted graph $\mathcal{G}$ to matrices by:

$$T(\mathcal{G}) = \text{Out}(\text{D}(\text{E}(\text{em}(\mathcal{G}_p)), \text{em}(\mathcal{G}_s))),$$

where $\mathcal{G}_p$ is the prefix and $\mathcal{G}_s$ the suffix of $\mathcal{G}$.

Equivalence

Theorem (Ahvonen, Heiman, Kuusisto, M., Selin, 26)

The following have the same expressive power:

- ▶ *Encoder–decoder transformers without the final softmax*
- ▶ *GPTL⁻ Logic*
- ▶ *CPG Automata*

References

- ▶ V. Ahvonen, D. Heiman, A. Kuusisto, and C. Lutz, *Logical Characterizations of Recurrent Graph Neural Networks with Reals and Floats.*, In A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang, editors, *Advances in Neural Information Processing Systems* **37**, 104205–104249 (2024).
- ▶ L. Hella, M. Järvisalo, A. Kuusisto, J. Laurinharju, T. Lempiäinen, K. Luosto, J. Suomela, and J. Virtema, *Weak Models of Distributed Computing, with Connections to Modal Logic*, *Distrib. Comput.* **28**, 31–53 (2015).

References

- ▶ A. Kuusisto, *Modal Logic and Distributed Message Passing Automata*, In Simona Ronchi Della Rocca, editor, *Computer Science Logic 2013 (CSL 2013)* of Leibniz International Proceedings in Informatics (LIPIcs), **23**, 452–468 (1913).
- ▶ A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, and I. Polosukhin, *Attention Is All You Need*, In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, Curran Associates **30**, (2017).