

Inherently Explainable Reinforcement Learning in Natural Language

Xiangyu Peng[†]

Mark O. Riedl[†]

Prithviraj Ammanabrolu[‡]

[†]Georgia Institute of Technology
{xpeng62, riedl}@gatech.edu,

[‡]Allen Institute for AI
raja@allenai.org

Abstract

We focus on the task of creating a reinforcement learning agent that is inherently explainable—with the ability to produce immediate local explanations by thinking out loud while performing a task and analyzing entire trajectories post-hoc to produce causal explanations. This Hierarchically Explainable Reinforcement Learning agent (HEX-RL),¹ operates in Interactive Fictions, text-based game environments in which an agent perceives and acts upon the world using textual natural language. These games are usually structured as puzzles or quests with long-term dependencies in which an agent must complete a sequence of actions to succeed—providing ideal environments in which to test an agent’s ability to explain its actions. Our agent is designed to treat explainability as a first-class citizen, using an extracted symbolic knowledge graph-based state representation coupled with a Hierarchical Graph Attention mechanism that points to the facts in the internal graph representation that most influenced the choice of actions. Experiments show that this agent provides significantly improved explanations over strong baselines, as rated by human participants generally unfamiliar with the environment, while also matching state-of-the-art task performance.

1 Introduction

Explainable AI refers to artificial intelligence methods and techniques that provide human-understandable insights into how and why an AI system chooses actions or makes predictions. Such explanations are critical for ensuring reliability and improving trustworthiness by increasing user understanding of the underlying model. In this work we specifically focus on creating deep reinforcement learning (RL) agents that can explain their actions in sequential decision making environments through natural language.

¹Code: <https://github.com/xiangyu-peng/HEX-RL>

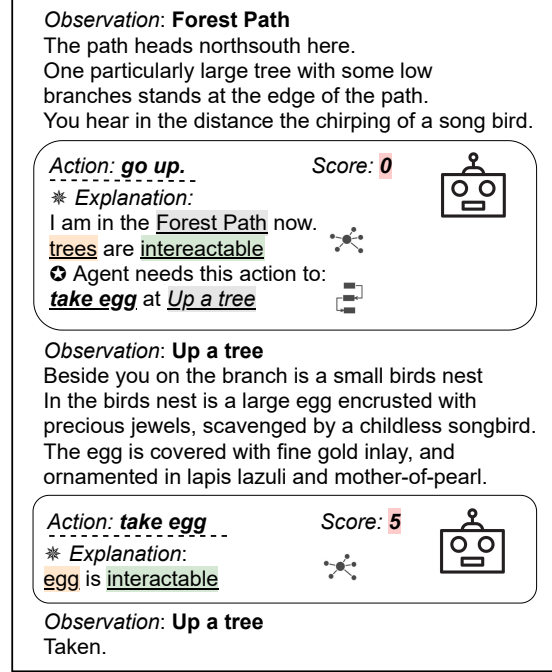


Figure 1: Excerpt from *zork1* and immediate step-by-step explanations extracted from knowledge graph represented by * and global trajectory explanations represented by E.

In contrast to the majority of contemporary work in the area which focuses on supervised machine learning problems requiring singular instance level local explanations (You et al., 2016; Xu et al., 2015; Wang et al., 2017; Wiegrefe and Marasovic, 2021), such environments—in which agents need to reason causally about actions over a long series of steps—require an agent to take into account both environmentally grounded context as well as goals when producing explanations. Agents implicitly contain beliefs regarding the downstream effects—the changes to the world—that actions taken at the current timestep will have. This requires explanations in these environments to contain an additional causal component taking the full trajectory’s context into account—complementary to the immedi-

ate step-by-step explanations.

Interactive Fiction (IF) games (Fig. 1) are partially observable environments where an agent perceives and acts upon a world using potentially incomplete textual natural language descriptions. They are structured as long puzzles and quests that require agents to reason about thousands of locations, characters, and objects over hundreds of steps, creating chains of dependencies that an agent must fulfill to complete the overall task. They provide ideal experimental test-beds for creating agents that can both reason in text and explain it.

We introduce an approach to game playing agents—**Hierarchically Explainable Reinforcement Learning (HEX-RL)**—that is designed to be *inherently* explainable, in the sense that its internal state representation—i.e. belief state about the world—takes the form of a symbolic, human-interpretable knowledge graph (KG) that is built as the agent explores the world. The graph is encoded by a Graph Attention network (GAT) (Veličković et al., 2017) extended to contain a hierarchical graph attention mechanism that focuses on different sub-graphs in the overall KG representation. Each of these sub-graphs contains different information such as attributes of objects, objects the player has, objects in the room, current location, etc. Using these encoding networks in conjunction with the underlying world KG, the agent is able to create *immediate explanations* akin to a running commentary that points to the facts within this knowledge graph that most influence its current choice of actions when attempting to achieve the tasks in the game on a step-by-step basis.

While graph attention can tell us which elements in the KG are attended to when maximizing expected reward from the current state, it cannot explain the intermediate, unrewarded dependencies that need to be satisfied to meet the long term task goals. For example, in the game *zork1*, the agent needs to pick up a lamp early on in the game—an unrewarded action—but the lamp is only used much later on to progress through a location without light. Thus, our agent also additionally analyzes an overall episode trajectory—a sequence of knowledge graph states and actions from when the agent first starts in a world to either task completion or agent death—to find the intermediate set of states that are most important for completing the overall task. This information is used to generate a *causal explanation* that condenses the *immediate*

step-by-step explanations to only the most important steps required to fulfill dependencies for the task.

Our contributions are as twofold: (1) we create an inherently explainable agent that uses an ever-updating knowledge-graph based state representation to generate step-by-step immediate explanations for executed actions as well as performing a post-hoc analysis to create causal explanations; and (2) a thorough experimental study against strong baselines that shows that our agent generates significantly improved explanations for its actions when rated by human participants unfamiliar with the domain while not losing any performance compared to the current state-of-the-art knowledge graph-based agents.

2 Background and Related Work

Interactive Fiction (IF) games are simulations featuring language-based state and action spaces. In this paper, we use IF games as our test-bed because they provides an ideal platform for collecting data, linking game states and actions to their corresponding natural language explanations. We use the definition of text-adventure games as seen in Côté et al. (2018) and Hausknecht et al. (2020).

Formally, a text game can be defined as a partially-observable Markov Decision Process (POMDP): $G = \langle S, P, A, O, \Omega, R, \gamma \rangle$, representing the set of environment states, conditional transition probabilities between states, the vocabulary or words used to compose text commands, observations, observation conditional probabilities, reward function, and discount factor, respectively. The reinforcement learning agent is trained to learned a policy $\pi_G(o) \rightarrow a$.

Knowledge Graphs for Text Games. Ammanabrolu et al. (2020) proposed Q*BERT, a reinforcement learning agent that learns a knowledge graph of the world by answering questions. Xu et al. (2020) used a stacked Hierarchical Graph Attention mechanism to construct an explicit representation of the reasoning process by exploiting the structure of the knowledge graph. Adhikari et al. (2020) present the Graph-Aided Transformer Agent (GATA) which learns to construct a knowledge graph during game play and improves zero-shot generalization on procedurally generated TextWorld games. Other works such as Murugesan et al. (2020) explore how to use KGs to endow agents with commonsense. While these works

showcase the effectiveness of KGs on task performance, they do not generate explanations.

Explainable Deep RL. Contemporary work on explaining deep reinforcement learning policies can be broadly categorized based on: (1) how the information is extracted, either via intrinsic motivation during training (Shu et al., 2017; Hein et al., 2017; Verma et al., 2018) or through post-hoc analysis (Rusu et al., 2015; Hayes and Shah, 2017; Juozapaitis et al., 2019; Madumal et al., 2020); and (2) the scope—either global (Zahavy et al., 2016; Hein et al., 2017; Verma et al., 2018; Liu et al., 2018) or local (Shu et al., 2017; Liu et al., 2018; Madumal et al., 2020; Guo et al., 2021). In our work, we create an agent that spans more than one of these categories providing immediately local explanations through extracted knowledge graph representations and post-hoc causal explanations. Inspired by Madumal et al. (2020), we learn a graphical causal model though our model focuses on using relations between steps in a puzzle instead of generating counterfactuals.

3 Hierarchically Explainable RL

Our work aims to generate (1) *immediate explanations* of an agent’s policy by capturing the importance of the current game state observation and (2) *causal explanations* that take into context an entire trajectory via a post-hoc analysis. Formally, let $\mathbf{X} = \{\mathbf{s}_t, \mathbf{a}_t\}_{t=1:T}$ be the set of game steps that compose a trajectory. Each game state s_t consists of a knowledge graph G_t representing all the information learned since the start of the game. This graph is further split into four sub-knowledge graphs $G_t^{atr}, G_t^{inv}, G_t^{obj}, G_t^{loc}$ each containing different, semantically related relationship types. This section first describes a graph attention based architecture that uses these sub-graphs to produce immediate explanations. We then describe how to filter the game states in a trajectory into a condensed set of the most important ones $\mathbf{X}' \subset \mathbf{X}$ that best capture the underlying dependencies that need to be fulfilled to complete the task—enabling us to generate global dependencies.

Knowledge Graph State Representation. Building on Ammanabrolu et al. (2020), we treat the problem of constructing the knowledge graph as a question-answering task. Knowledge graphs in these games take the form of RDF triples of $\langle \text{subject}, \text{relation}, \text{object} \rangle$ that are extracted

from text observations and update as the agent explores the world (Figure 2). The agent answers questions about the environment such as, “What am I carrying?” or “What objects are around me?”. A specially constructed dataset for question answering in text games—JerichoQA—is used to fine-tune an ALBERT (Lan et al., 2019) model to answer these questions. The answers to these questions are form a set of candidate graph vertices V_t for the current step and questions form the set of relations R_t . Both V_t and R_t are then combined with the graph at the previous step G_{t-1} to update the agent’s belief about the world state into G_t . The left side of Figure 2 showcases this.

In an attempt to enable more fine grained explanation generation and inspired by Xu et al. (2020), we divide the knowledge graph G into multiple sub-graphs $G^{atr}, G^{inv}, G^{obj}, G^{loc}$, each representing (1) attributes of objects, (2) objects the player has, (3) objects in the room, and (4) other information such as location (see right side of Figure 2) based on the corresponding relationship types extracted by the ALBERT-QA module. The union of all sub-graphs is equivalent of V_t and R_t extracted from the current game state. The full knowledge graph G_t captures the overall game state since the start of the game. The sub-graphs easily reflect different relationships of the current game state.

Template Action Space. Agents output a language string into the game to describe the actions that they want to perform. To ensure tractability, this action space can be simplified down into templates. Templates consist of interchangeable verbs phrases (VP), optionally followed by prepositional phrases ($VP PP$), e.g. ($[\text{carry/take}] __$) and ($[\text{throw/discard/put}] __ [\text{against/on/down}] __$), where the verbs and prepositions within $[_]$ are aliases. Actions are constructed from templates by filling in the template’s blanks using words in the game’s vocabulary.

3.1 Immediate Explanations

Our immediate explanations consist of finding the subset of triplets in sub-graphs $G^{atr}, G^{inv}, G^{obj}, G^{loc}$ that most influence the action decision made at the current step. We introduce a deep RL architecture capable of this.

Hierarchical Knowledge Graph Attention Architecture. At every step, a total score R_t and an observation o_t is received—consisting of $(o_{t_{desc}}, o_{t_{game}}, o_{t_{inv}}, a_{t-1})$ corresponding to the

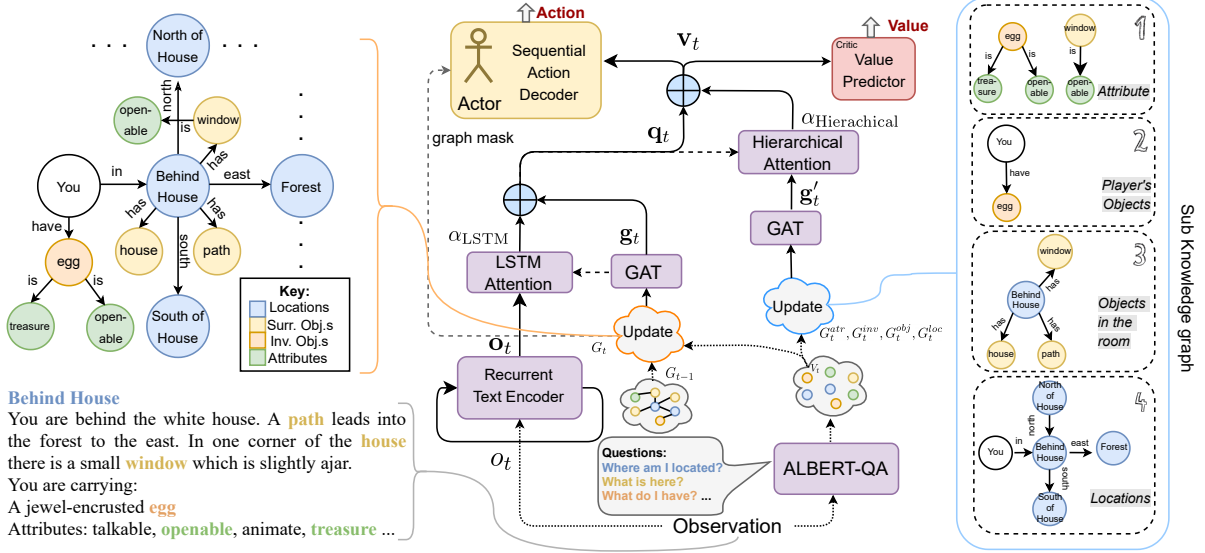


Figure 2: One-step knowledge graph extraction and overall hierarchically explainable reinforcement learning agent (HEX-RL) architecture at time step t .

room description, game feedback, inventory, and previous action. These components are processed using a GRU based encoder, utilizing the hidden state from the previous step and combined to have a single observation embedding $\mathbf{o}_t$ (bottom of Figure 2).

The full knowledge graph G_t is processed via Graph Attention Networks (GATs) (Veličković et al., 2017) followed by a linear layer to get the graph representation $\mathbf{g}_t$ (middle in Figure 2). Then we compute the *LSTM attention* between $\mathbf{o}_t$ and $\mathbf{g}_t$ by:

$$\alpha_{\text{LSTM}} = \text{softmax}(\mathbf{W}_1 \mathbf{h}_{\text{LSTM}} + \mathbf{b}_1) \quad (1)$$

$$\mathbf{h}_{\text{LSTM}} = \tanh(\mathbf{W}_g \mathbf{g}_t \oplus (\mathbf{W}_o \mathbf{o}_t + \mathbf{b}_o)) \quad (2)$$

where $\oplus$ denotes the addition of a matrix and a vector. $\mathbf{W}_1$, $\mathbf{W}_g$, $\mathbf{W}_o$ are weights and $\mathbf{b}_1$, $\mathbf{b}_o$ are biases. The overall representation vector is updated as,

$$\mathbf{q}_t = \mathbf{g}_t + \sum_i^c \alpha_{\text{LSTM},i} \odot \mathbf{o}_{t,i} \quad (3)$$

where $\odot$ denotes dot-product, c is the number of $\mathbf{o}_t$'s components.

Sub-graphs are also encoded by GATs to get the graph representation $\mathbf{g}'_t$. The *Hierarchical Graph Attention* between $\mathbf{q}_t$ and $\mathbf{g}'_t$ is calculated by:

$$\alpha_{\text{Hierarchical}} = \text{softmax}(\mathbf{W}_H \mathbf{h}_H + \mathbf{b}_H) \quad (4)$$

$$\mathbf{h}_H = \tanh(\mathbf{W}_{g'} \mathbf{g}'_t \oplus (\mathbf{W}_q \mathbf{q}_t + \mathbf{b}_q)) \quad (5)$$

Where $\mathbf{W}_H$, $\mathbf{W}_{g'}$, $\mathbf{W}_q$ are weights and $\mathbf{b}_H$, $\mathbf{b}_q$ are biases. Then we get state representation, consisting of the textual observations full knowledge graph and sub-knowledge graph.

$$\mathbf{v}_t = \mathbf{q}_t + \sum_i^s \alpha_{\text{Hierarchical},i} \odot \mathbf{g}'_{t,i} \quad (6)$$

where s is the number of sub-graphs (4 in our paper). The full architecture can be found in Figure 2.

The agent is trained via the Advantage Actor Critic (A2C) (Mnih et al., 2016) method to maximize long term expected reward in the game in a manner otherwise unchanged from Ammanabrolu et al. (2020) (See Appendix A.1). These attention values thus reflect the portions of the knowledge graphs that the agent must focus on to best achieve this goal of maximizing reward.

Hierarchical Graph Attention Explanation.

The graph attention $\alpha_{\text{Hierarchical}}$ (seen in the upper right in Figure 2) is used to capture the relative importance of game state observations and KG entities in influencing action choice. For each sub-graph, the graph attention, $\alpha_{\text{Hierarchical},i} \in \mathbb{R}^{n_{\text{nodes}} \times m}$ is summed over all the channels m to obtain $\alpha'_{\text{Hierarchical},i} \in \mathbb{R}^{n_{\text{nodes}} \times 1}$, showing the importance of the KG nodes in the i th sub-graph. The top- k valid entities (and corresponding edges) with highest absolute value of its attention form the set

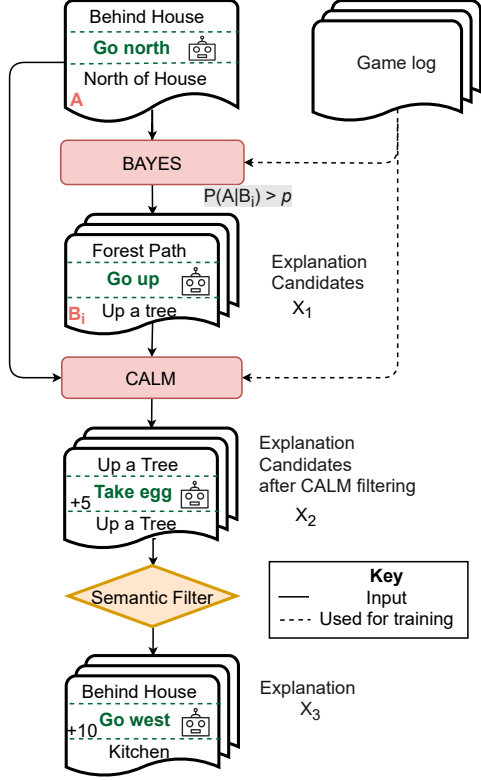


Figure 3: Pipeline for creating a causal explanation of why the agent chose the action—"go north" at "behind house".

of knowledge graph triplets that best locally explain the action a_t .

In order to make the explanation more readable for a human reader, we further transform knowledge graph triplets to natural language by template filling. For example, $\langle tree, in, forest \rangle$ is converted to "Tree is in the forest". Further details regarding template construction are found in Appendix A.2

3.2 Causal Explanations

Graph attention tells us which entities in the KG are attended to when making a decision, but is not enough alone for explaining "why" actions are the right ones in the context of fulfilling dependencies that may potentially be unrewarded by the game—especially given the fact that there are potentially multiple ways of achieving the overall task. HEX-RL thus saves trajectories for hundreds of test time rollouts of the games, performed once a policy has been trained (Table 1). The game trajectories consist of all the game state o_t , action taken, predicted critic value, game scores, the knowledge graph, and the immediate step level explanation generated as previously described. HEX-RL produces a causal

STEP: 16

Text Observation:

Up a Tree
Beside you on the branch is a small birds nest.
In the birds nest is a large egg encrusted with jewels,
apparently scavenged by a childless songbird...

Knowledge graph:

$\langle tree, in, forest \rangle$;
 $\langle egg, is, interactable \rangle$, $\langle nest, is, interactable \rangle$...

Action: take egg

Hierarchical attention explanation:

egg is interactable

Game Score is 5

Critic Value: 5.7457

Table 1: Example state saved during game play as part of a trajectory.

explanation by performing a post-hoc analysis on these game trajectories. The agent then analyzes and filters these trajectories in an attempt to find the subset of states that are most crucial to achieving the task as summarized in Figure 3—then using that subset of states to generate causal trajectory level explanations.

Bayesian State Filter. We first train a Bayesian model to predict the conditional probability $\mathbb{P}(A | B_i)$ of a game step (A) given any other possible game step (B_i) in the game trajectories. The conditional probability $\mathbb{P}(A | B_i)$ is calculated by,

$$\mathbb{P}(A | B_i) = \frac{\mathbb{P}(B_i | A)\mathbb{C}(A)}{\mathbb{C}(B_i)} \quad (7)$$

where $\mathbb{C}(A)$ and $\mathbb{C}(B_i)$ stand for the raw count of game step A and B_i in the collected trajectories. The key intuition here being that state, action pairs that appear in a certain ordering in multiple trajectories are more likely to dependant on each other. The set of game steps with the highest $\mathbb{P}(A | B_i)$ is used to explain taking the action associated with game state A . For example, "take egg" (A) is required to "open egg" (B), and $\mathbb{P}(A | B) = 1$, hence "open egg" is used as a reason why action "take egg" must be taken first. The initial set of game states $\mathbf{X}$ is filtered into $\mathbf{X}_1$ by working backwards from the final goal state by finding the set of states that form the most likely chain of causal dependencies that lead to it.

Language Model Action Filter. Following this, we apply a GPT-2 (Radford et al., 2019) language model trained to generate actions based on transcripts of text games from human play-throughs to

further filter out important states—known as the Contextual Action Language Model (CALM) (Yao et al., 2020). As this language model is trained on human transcripts, we hypothesize that it is able to further filter down the set of important states by finding the states that have corresponding actions that a human player would be more likely to perform—thus potentially leading to more natural explanations. CALM takes into observation o_t , action a_t and the following observation o_{t+1} , and predicts next valid actions a_{t+1} . In our work, we use CALM as a filter to look for the relations between a game step A and the explanation candidates $B_i \in \mathbf{X}_1$. We feed CALM with the prompt o_A, a_A, o_{B_i} to get an action candidate set. When the two game steps A and B_i are highly correlated, given o_A, a_A and o_{B_i} , CALM should successfully predict a_{B_i} with high probability. The game steps B_i , whose associated action a_{B_i} is in this generated action candidates set, are saved as the next set of filtered important candidate game states ($\mathbf{X}_2$).

Semantic State-Action Filter. To better account for the irregularities of the puzzle like environment, we adopt a *semantic filter* to obtain the final important state set $\mathbf{X}_3$. Here, given $A, B_i \in \mathbf{X}_2$, states are further filtered on the basis of whether one of these scenarios occurs:

- a_A and a_{B_i} contain the same entities. For example, “take egg” and “open egg”.
- G_A and G_{B_i} KGs for both states share the same entities. For example, “lamp” occurs in both observations.
- A and B_i are occur in the same location. For example, after taking action a_A , the player enters “kitchen” and B occurs in “kitchen”.
- The state has a non-zero reward or a high absolute critic value, indicating that it is either a state important for achieving the goals of the game or it is a state to be avoided.

This final set of important game states $\mathbf{X}_3$ is used to synthesize post-hoc causal explanations for why an action was performed in a particular state—as seen in Figure 1—taking into account the overall context of the dependencies required to be satisfied and building on the immediate step level explanations for each given state in $\mathbf{X}_3$.

4 Evaluation

Our evaluation consists of four phases: (1) We show that HEX-RL has the comparable performance to state-of-art reinforcement learning agents

on text games in Section 4.1. (2) Then in Section 4.2, we evaluate our immediate attention explanation model by comparing the explanations generated by HEX-RL and agents that do not use knowledge graphs (See Figure 2 and Section 3.1). (3) In Section 4.3 we compare immediate to causal explanations, focusing on the effects that including trajectory level context when evaluating explanations in the context of agent goals. (4) In Section 4.4 we conduct human participant ablation study evaluating the individual contributions of the filtration pipeline for generating causal explanations seen in Figure 3.

4.1 Task Performance Evaluation

We compare HEX-RL with three strong state-of-art reinforcement learning agents—focusing on contemporary agents that use knowledge graphs—on an established test set of 9 games from the Jericho benchmark (Hausknecht et al., 2020).

- **LSTM-A2C.** This is a baseline that only uses the natural language observation as state representation that is encoded with an LSTM-based policy network.
- **KG-A2C.** Instead of training a question-answering system like Q*BERT to build knowledge graph state representation, KG-A2C (Ammanabrolu and Hausknecht, 2020) extracts knowledge graph triplets from the text observations using a rules based approach built on OpenIE (Angeli et al., 2015).
- **SHA-KG.** Is adapted from Xu et al. (2020) and uses a rules-based approach to construct a knowledge graph for the agent which is then fed into a Hierarchical Graph Attention network as in HEX-RL. This agent separates the sub-graphs out using a rules-based approach and makes no use of any edge relationship information in the graph.
- **Q*BERT.** (Ammanabrolu et al., 2020) uses a similar method of creating the knowledge graph through question answering but does not use the hierarchical graph attention architecture combined with the sub-graphs.

These baselines are all trained via the Advantage Actor Critic (A2C) (Mnih et al., 2016) method—further comparisons to other contemporary agents can be found in Appendix A.7. It is also worth noting that most contemporary state of the art deep RL

Experiment	LSTM-A2C		KG-A2C		SHA-KG		Q*BERT		HEX-RL <i>Game Only</i>		HEX-RL <i>Game_and_IM</i>		Max
Metric	Eps.	Max	Eps.	Max	Eps.	Max	Eps.	Max	Eps.	Max	Eps.	Max	-
zork1	27	31.2	34	35	33.6	34.5	35	35	29.8	40	30.2	40	350
library	8.2	10	14.3	19	10.0	15.8	18	18	15.94	19	13.8	21	30
detective	141	188	207.9	214	246.1	308	274	310	276.65	330	276.93	330	360
balances	10	10	10	10	9.8	10	10	10	9.95	10	10	10	51
pentari	50.4	55	50.7	56	48.2	51.3	50	56	34.61	55	44.7	60	70
ztuu	5	5	5	5	5	25	5	5	5	5	5.08	9	100
ludicorp	14.4	18	17.8	19	17.6	17.8	18	19	14.0	18	17.6	18	150
deephme	1	1	1	1	1	1	1	1	1	1	1	1	300
temple	8	8	7.6	8	7.9	6.9	8	8	8	8	7.58	8	35
% compl.	22.6	25.9	27.3	30.8	27.2	33.1	30.8	34.9	27.2	33.9	28.2	35.8	100

Table 2: Asymptotic scores on games by different methods across 5 independent runs. *Eps.* indicates scores averaged across the final 100 episodes and *Max* indicates the maximum score seen by the agent over the same period. We present results on two training rewards, *game only* and *game_and_IM*.

agents for text games use recurrent neural policy networks as opposed to transformer networks due to their improved performance in this domain.

HEX-RL Training. We trained HEX-RL on two reward types: (a) game only and (b) game with intrinsic motivation (IM). *Game only* indicates that we only use score obtained from the game as reward. *Game and IM* contains an additional intrinsic motivation reward based on knowledge graph expansion as seen in Ammanabrolu et al. (2020)—where the agent is additionally rewarded for learning more about the world by finding new facts for knowledge graph. Further details are found in Appendix A.4.

Table 2 shows the performance of HEX-RL and the other four baselines. We can see that designing the HEX-RL agent to be inherently explainable through the use of Hierarchical Graph Attention and the sub-graphs improves the overall maximum score seen during training when compared to any of the other agents. In terms of the average score seen during the final 100 episodes, HEX-RL with intrinsic motivation outperforms all baselines with the exception of Q*BERT—there HEX-RL significantly outperforms Q*BERT on one game, is outperformed on two games, and comparable on the remaining six games. HEX-RL thus performs comparably to other state-of-the-art baselines in terms of overall task performance while also boasting the additional ability to explain its actions.

4.2 Immediate Explanation Evaluation

Having established that HEX-RL’s performance while playing text games is comparable to other state-of-the-art agents, we attempt to answer the question of exactly how useful the knowledge graph based architecture is when generating im-

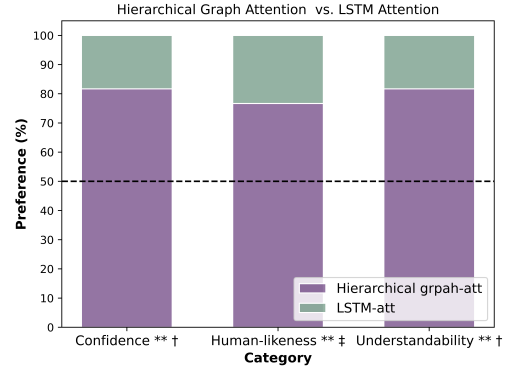


Figure 4: Human evaluation results showing the proportion of participants that prefer Hierarchical Graph Attention vs. LSTM Attention explanations, * indicates $p < 0.05$, ** indicates $p < 0.01$, † indicates $\kappa > 0.2$ or fair agreement. ‡ indicates $\kappa > 0.4$ or moderate agreement.

mediate step-by-step explanations by comparing HEX-RL to a baseline that doesn’t use knowledge graphs in a human participant study. The two models being compared are where step-by-step explanations are generated by:

- **LSTM Attention explanations.** Extracts the most important substring in the observations through LSTM attention α_{LSTM} and then uses those words to create an explanation.
- **Hierarchical Graph Attention explanations.** Extracts the knowledge graph triplets that most influenced the choice of actions by Hierarchical Attention $\alpha_{Hierarchical}$ and then transforming them into readable language explanations through templates.

We recruited 40 participants—generally unfamiliar with the environment at hand—on a crowd

sourcing platform. Each participant reads a randomly selected subset of 10 explanation pairs (drawn randomly from a pool totaling 60 explanation pairs), generated by Hierarchical Graph Attention and LSTM attention explanation on three games in the Jericho benchmark: *zork1*, *library*, and *balances*.

They answer three questions that evaluates dimensions such as confidence, human-likeness, readability and understandability. Variations of these questions have been used to evaluate other explainable AI systems (eg. Ehsan et al. (2019)). Participants are given the following metrics and asked to choose which explanation they prefer with respect to that metric:

- *Confidence*: This explanation makes you more confident that the agent made the right choice.
- *Human-likeness*: This explanation expresses more human-like thinking on the action choice.
- *Understandability*: This explanation makes you understand why the agent made the choice.

At least 5 participants give their preference for each explanation pair. Further details are found in Appendix B.1

Figure 4 shows the result of the human evaluation of attention explanations. Hierarchical graph attention explanation is preferred over LSTM attention explanation in all three dimensions. These results are statistically significant ($p < 0.05$) with fair inter-rater reliabilities. We also observe that these three dimensions are highly, positively correlated using Spearman’s Rank Order Correlation.²

A slightly higher proportion of participants preferred the LSTM Attention explanations in the human-likeness dimension compared to the other two. When this small proportions was asked to justify their choice in under 50 words, the participants preferring LSTM Attention explanation stated that they found it intuitive and found the Hierarchical Graph Attention explanations to be more robotic.

On the other hand, whenever participants did not prefer the LSTM Attention, they justified it by stating that the explanation was relatively incoherent. This appears to indicate that the LSTM Attention based explanation—constructed directly to be a substring of the human-written observation—

² $r_s = 0.70, p < 0.01$, between “confidence” and “understandability”; $r_s = 0.67, p < 0.01$, between “confidence” and “human-likeness”; $r_s = 0.89, p < 0.01$, between “human-likeness” and “understandability”

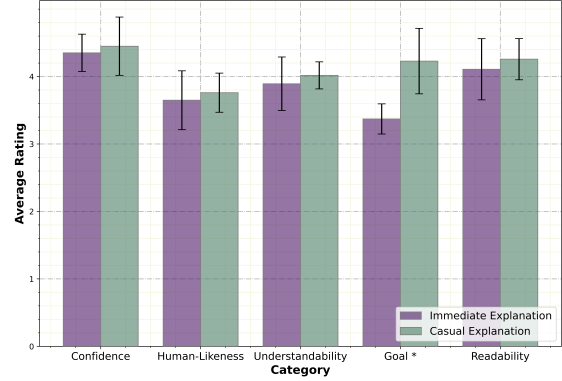


Figure 5: Human judgment⁵ results on immediate and causal explanation, * indicates $p < 0.05$. A confidence level of 95% over all the explanations is also presented.

presents a more human-like explanation for the actions than the templated Hierarchical Graph Attention explanations but only when it is coherent enough to be understood. The knowledge graph sacrifices a small amount of human-likeness in return for much greater overall understandability. Overall, we conclude that on a step-by-step level using knowledge graphs with Hierarchical Graph Attention networks gives us explanations that are more easily understood and inspire greater confidence in the agent’s decisions.

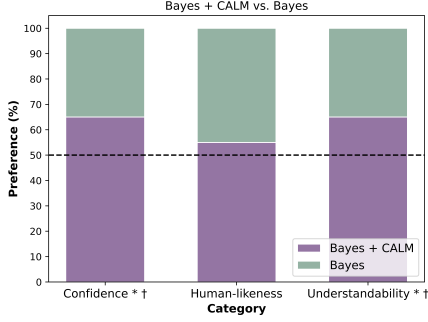
4.3 Immediate vs. Causal Explanation Evaluation

Having proved the effectiveness of the knowledge graph at the immediate step-by-step explanation level, we attempt to evaluate our method of producing causal explanations and how they compare to the immediate explanations. We assess whether the causal explanation condensing a trajectory into important steps (1) is able to maintain the coherence like the immediate explanations do; and (2) informs a human better about the agent’s actions when taken in the context of the goals of the agent.

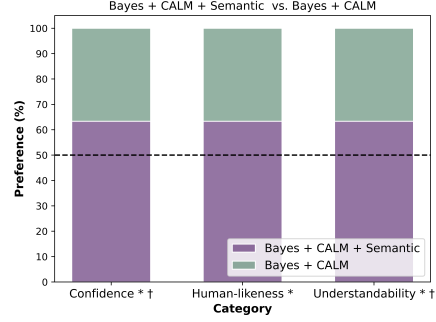
Participants first read the full trajectory of the game combined with step-by-step immediate explanations, along with summary of the game goal, and indicate how much they agree with the five statements on a Likert scale³. For this study, we take all the metrics from the previous study and add two more:

- *Goal context*: You are able to understand why the agent takes this particular sequence of ac-

³1: Strong Disagree; 2: Disagree; 3: Undecided; 4: Agree; 5: Strong Agree



(a) Bayes vs. Bayes+CALM explanation



(b) Bayes+CALM vs. Bayes+CALM+Semantic explanation

Figure 6: Human evaluation results on ablation study, * indicates $p < 0.05$, † indicates $\kappa > 0.2$ or fair agreement.

tions given what you know about the goal.

- **Readability:** This explanation is easy to read. Participants then read the causal explanation of same trajectory and rate it again along the same five statements. At least 5 crowd workers rated each explanation. Further details are found in Appendix B.2

Figure 5 shows the average scores for each question for the immediate and causal explanations. The causal explanations achieve comparable performance to the immediate explanations on all metrics except for the the metric relating to goal context. On the goal context metric, the causal explanation significantly out-performs the immediate explanations. These results indicate that HEX-RL can successfully identify the most important states in a trajectory and use them to create causal explanations that are on par with immediate explanations in terms of coherence but provide significantly more context in terms of explaining an agent’s actions with respect to its task-based goals. We further note that when the participants were asked to justify these choices, a majority stated that a condensed causal explanation based on important steps made the goals of the agent easier to understand than reading through an explanation for every single step the agent took.

4.4 Causal Explanation Ablation Study

Having established the overall effectiveness of the filters in HEX-RL that create the causal explanations, we perform pair-wise ablation studies to pinpoint the relative contributions of the different filters seen in Figure 3. We first compare explanations generated using a set of important states filtered from the trajectory using the Bayes model compared to Bayes+CALM explanation. This how applying the language model action filter affects the

quality of the causal explanations. As before, we recruited 30 participants on a crowd sourcing platform. Each participant reads a randomly selected subset of explanation pairs, comprised of causal explanations filtered by Bayes and Bayes+CALM models. Figure 6a shows that after applying the CALM model to filter explanation candidates, generated explanations are significantly preferred on the “Confidence” and “Understandability” dimensions.

Similarly, we then conducted another ablation study to validate the contribution of semantic filter by comparing the Bayes+CALM filtering method to the full HEX-RL using Bayes+CALM+Semantic filters. The experiment setup is the same as the previous ablation study. Figure 6b shows that Bayes+CALM+Semantic performs significantly better than Bayes+CALM on all three dimensions.

We additionally observe that these three metrics are highly, positively correlated using Spearman’s Rank Order Correlation in both of these ablation studies⁴. When asked to justify their choices, participants indicated that the full HEX-RL system with Bayes+CALM+Semantic filters provided causal explanations that they felt was more understandable than alternatives. These results indicate that all three steps of the filtering process to identify important states are necessary for creating coherent causal explanations that effectively take into account the context of the agent’s goals.

5 Conclusions

Explaining deep RL policies for sequential decision making problems in natural language is a sparsely

⁴ $r_s = 0.86$, $p < 0.01$, between “confidence” and “understandability”; $r_s = 0.79$, $p < 0.01$, between “confidence” and “human-likeness”; $r_s = 0.90$, $p < 0.01$, between “human-likeness” and “understandability”

studied problem despite a steadily growing need. An oft given reason for this phenomenon is that deep RL methods perform better without the additional burden of being explainable. In an attempt to encourage work in this area, we create the Hierarchically Explainable Reinforcement Learning (HEX-RL) agent which treats explainability as a first-class citizen in its design by using a readily interpretable knowledge graph state representation coupled with a Hierarchical Graph Attention network. This agent is able to produce step-by-step commentary-like immediate explanations and also a condensed causal trajectory level explanation via a post-hoc analysis. We show that with careful design, it is possible to create inherently explainable RL agents that do not lose performance when compared to contemporary state-of-the-art agents and simultaneously are able to generate significantly higher quality explanations of their actions.

References

- Ashutosh Adhikari, Xingdi Yuan, Marc-Alexandre Côté, Mikuláš Zelinka, Marc-Antoine Rondeau, Romain Laroché, Pascal Poupart, Jian Tang, Adam Trischler, and William L. Hamilton. 2020. [Learning dynamic belief graphs to generalize on text-based games](#). *CoRR*, abs/2002.09127.
- Prithviraj Ammanabrolu and Matthew Hausknecht. 2020. [Graph constrained reinforcement learning for natural language action spaces](#). In *International Conference on Learning Representations*.
- Prithviraj Ammanabrolu, Ethan Tien, Matthew Hausknecht, and Mark O Riedl. 2020. How to avoid being eaten by a grue: Structured exploration strategies for textual worlds. *arXiv preprint arXiv:2006.07409*.
- Gabor Angeli, Melvin Jose Johnson Premkumar, and Christopher D Manning. 2015. Leveraging linguistic structure for open domain information extraction. In *Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing (Volume 1: Long Papers)*, pages 344–354.
- Marc-Alexandre Côté, Akos Kádár, Xingdi Yuan, Ben Kybartas, Tavian Barnes, Emery Fine, James Moore, Matthew Hausknecht, Layla El Asri, Mahmoud Adada, et al. 2018. Textworld: A learning environment for text-based games. In *Workshop on Computer Games*, pages 41–75. Springer.
- Upol Ehsan, Pradyumna Tambwekar, Larry Chan, Brent Harrison, and Mark O Riedl. 2019. Automated rationale generation: a technique for explainable ai and its effects on human perceptions. In *Proceedings of the 24th International Conference on Intelligent User Interfaces*, pages 263–274.
- Wenbo Guo, Xian Wu, Usman Khan, and Xinyu Xing. 2021. Edge: Explaining deep reinforcement learning policies. In *Thirty-Fifth Conference on Neural Information Processing Systems*.
- Matthew Hausknecht, Prithviraj Ammanabrolu, Marc-Alexandre Côté, and Xingdi Yuan. 2020. Interactive fiction games: A colossal adventure. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 34, pages 7903–7910.
- Bradley Hayes and Julie A Shah. 2017. Improving robot controller transparency through autonomous policy explanation. In *2017 12th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*, pages 303–312. IEEE.
- Daniel Hein, Alexander Hentschel, Thomas Runkler, and Steffen Udfluft. 2017. Particle swarm optimization for generating interpretable fuzzy reinforcement learning policies. *Engineering Applications of Artificial Intelligence*, 65:87–98.
- Zoe Juozapaitis, Anurag Koul, Alan Fern, Martin Erwig, and Finale Doshi-Velez. 2019. Explainable reinforcement learning via reward decomposition. In *IJCAI/ECAI Workshop on Explainable Artificial Intelligence*.
- Zhenzhong Lan, Mingda Chen, Sebastian Goodman, Kevin Gimpel, Piyush Sharma, and Radu Soricut. 2019. Albert: A lite bert for self-supervised learning of language representations. *arXiv preprint arXiv:1909.11942*.
- Guiliang Liu, Oliver Schulte, Wang Zhu, and Qingcan Li. 2018. Toward interpretable deep reinforcement learning with linear model u-trees. In *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, pages 414–429. Springer.
- Prashan Madumal, Tim Miller, Liz Sonenberg, and Frank Vetere. 2020. Explainable reinforcement learning through a causal lens. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 34, pages 2493–2500.
- Volodymyr Mnih, Adria Puigdomenech Badia, Mehdi Mirza, Alex Graves, Timothy Lillicrap, Tim Harley, David Silver, and Koray Kavukcuoglu. 2016. Asynchronous methods for deep reinforcement learning. In *International conference on machine learning*, pages 1928–1937. PMLR.
- Keerthiram Murugesan, Mattia Atzeni, Pushkar Shukla, Mrinmaya Sachan, Pavan Kapanipathi, and Kartik Talamadupula. 2020. [Enhancing text-based reinforcement learning agents with commonsense knowledge](#). *CoRR*, abs/2005.00811.

- Alec Radford, Jeff Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. 2019. Language models are unsupervised multitask learners.
- Pranav Rajpurkar, Robin Jia, and Percy Liang. 2018. [Know what you don’t know: Unanswerable questions for SQuAD](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 784–789, Melbourne, Australia. Association for Computational Linguistics.
- Andrei A Rusu, Sergio Gomez Colmenarejo, Caglar Gulcehre, Guillaume Desjardins, James Kirkpatrick, Razvan Pascanu, Volodymyr Mnih, Koray Kavukcuoglu, and Raia Hadsell. 2015. Policy distillation. *arXiv preprint arXiv:1511.06295*.
- Tianmin Shu, Caiming Xiong, and Richard Socher. 2017. Hierarchical and interpretable skill acquisition in multi-task reinforcement learning. *arXiv preprint arXiv:1712.07294*.
- Petar Veličković, Guillem Cucurull, Arantxa Casanova, Adriana Romero, Pietro Lio, and Yoshua Bengio. 2017. Graph attention networks. *arXiv preprint arXiv:1710.10903*.
- Abhinav Verma, Vijayaraghavan Murali, Rishabh Singh, Pushmeet Kohli, and Swarat Chaudhuri. 2018. Programmatically interpretable reinforcement learning. In *International Conference on Machine Learning*, pages 5045–5054. PMLR.
- Fei Wang, Mengqing Jiang, Chen Qian, Shuo Yang, Cheng Li, Honggang Zhang, Xiaogang Wang, and Xiaoou Tang. 2017. Residual attention network for image classification. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 3156–3164.
- Sarah Wiegreffe and Ana Marasovic. 2021. [Teach me to explain: A review of datasets for explainable natural language processing](#). In *Thirty-fifth Conference on Neural Information Processing Systems Datasets and Benchmarks Track (Round 1)*.
- Kelvin Xu, Jimmy Ba, Ryan Kiros, Kyunghyun Cho, Aaron Courville, Ruslan Salakhudinov, Rich Zemel, and Yoshua Bengio. 2015. Show, attend and tell: Neural image caption generation with visual attention. In *International conference on machine learning*, pages 2048–2057. PMLR.
- Yunqiu Xu, Meng Fang, Ling Chen, Yali Du, Joey Tianyi Zhou, and Chengqi Zhang. 2020. Deep reinforcement learning with stacked hierarchical attention for text-based games. *Advances in Neural Information Processing Systems*, 33.
- Shunyu Yao, Rohan Rao, Matthew Hausknecht, and Karthik Narasimhan. 2020. Keep calm and explore: Language models for action generation in text-based games. *arXiv preprint arXiv:2010.02903*.
- Quanzeng You, Hailin Jin, Zhaowen Wang, Chen Fang, and Jiebo Luo. 2016. Image captioning with semantic attention. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 4651–4659.
- Tom Zahavy, Nir Ben-Zrihem, and Shie Mannor. 2016. Graying the black box: Understanding dqns. In *International Conference on Machine Learning*, pages 1899–1908. PMLR.

A Implementation Details

A.1 A2C Architecture

Further details of what is found in Figure 2. The sequential action decoder consists two GRUs that are linked together as seen in Ammanabrolu and Hausknecht (2020). The first GRU decodes an action template and the second decodes objects that can be filled into the template. These objects are constrained by a *graph mask*, i.e. the decoder is only allowed to select entities that are already present in the knowledge graph.

A.2 Templates of Immediate Explanation

We consider four types of sub-graphs G^{atr} , G^{inv} , G^{obj} , G^{loc} , each representing (1) attributes of objects, (2) objects the player has, (3) objects in the room, and (4) other information such as location (see right side of Figure 2). Hence, we create one template for each sub-graph,

- $\langle object, is, attribute \rangle$ is converted to “*Object is attribute*”.
- $\langle player, has, object \rangle$ is converted to “*I have object*”.
- $\langle object, in, location \rangle$ is converted to “*Object is in location*”.
- $\langle location_1, direction, location_2 \rangle$ is converted to “*location_1 is in the direction of location_2*”.

A.3 Raw scores across Jericho supported games

Exp.	TDQN	DRRN	HEX-RL		Max
			<i>Game_and_IM</i>		
Metric	Eps.	Eps.	Eps.	Max	-
zork1	9.9	24.6	30.2	40	350
library	6.3	17	13.8	21	30
detective	169	197.8	276.93	330	360
balances	4.8	10	10	10	51
pentari	17.4	27.2	44.7	60	70
ztuu	4.9	21.6	5.08	9	100
ludicorp	6	13.8	17.6	18	150
deephome	1	1	1	1	300
temple	7.9	7.4	7.58	8	35
% compl.	15.2	25.5	28.2	35.8	100

Table 3: Raw scores across Jericho supported games. *Eps.* indicates scores averaged across the final 100 episodes and *Max* indicates the maximum score seen by the agent over the same period. We present results on *game and IM* reward.

A.4 Reward types

To alleviate the issue that rewards are sparse and often delayed, Ammanabrolu et al. defined an *in-*

trinsic motivation for the agent that leverages the knowledge graph being built during exploration. The motivation is for the agent to learn more information regarding the world and expand the size of its knowledge graph. They formally define *game_and_IM* reward in terms of new information learned.

$$r_{IM_t} = \Delta(\mathcal{KG}_{\text{global}} - \mathcal{KG}_t) \quad (8)$$

where $\mathcal{KG}_{\text{global}} = \bigcup_{i=1}^{t-1} \mathcal{KG}_i$. Here $\mathcal{KG}_{\text{global}}$ is the set of all edges that the agent has ever had in its knowledge graph and the subtraction operator is a set difference.

A.5 Knowledge Graph Representation QA Model

The question answering network based on ALBERT (Lan et al., 2019) has the following hyperparameters, taken from the original paper and known to work well on the SQuAD 2.0 (Rajpurkar et al., 2018) dataset. No further hyperparameter tuning was conducted.

Parameters	Value
batch size	8
learning rate	3e-5
max seq len	512
doc stride	128
warmup steps	814
max steps	8144
gradient accumulation steps	24

A.6 HEX-RL

The additional hyperparameters used for training HEX-RL are detailed below, same with Ammanabrolu et al. (2020). *graph dropout* and *mask dropout* are used for encouraging graph network to actually learn a sparse representation.

Parameters	Value
buffer size	40
batch size	16
graph dropout	0.2
mask dropout	0.1

A.7 Task Performance

We plot the training reward curve for 9 games in Figure 7 and Figure 8.

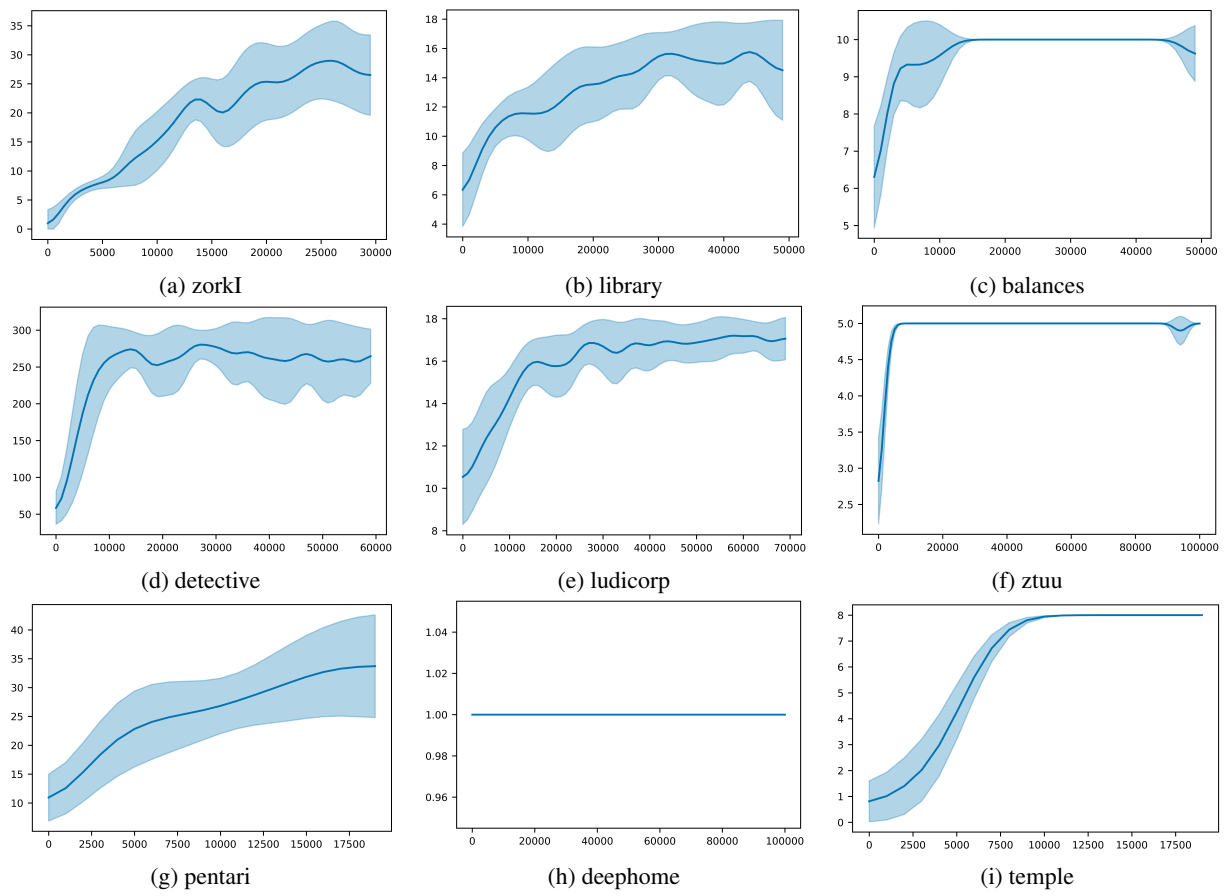


Figure 7: Eps. initial reward curves for the exploration strategies—*Game only* Reward

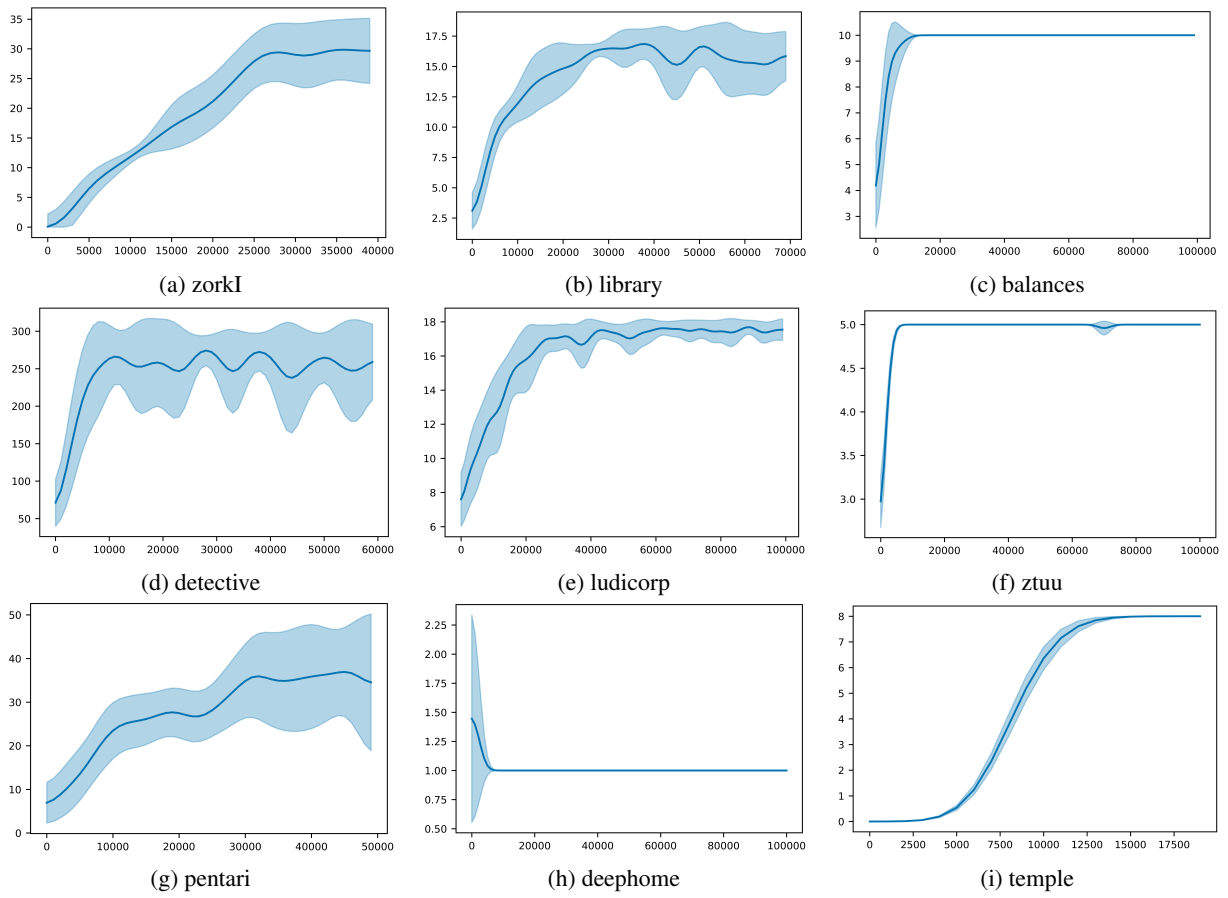
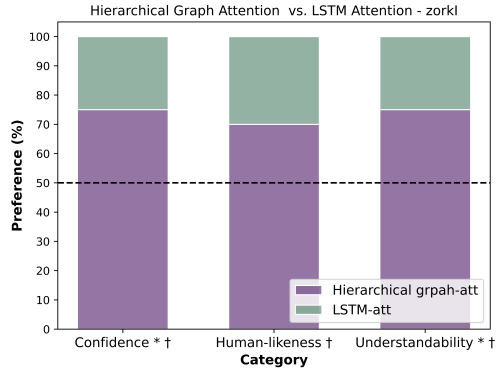


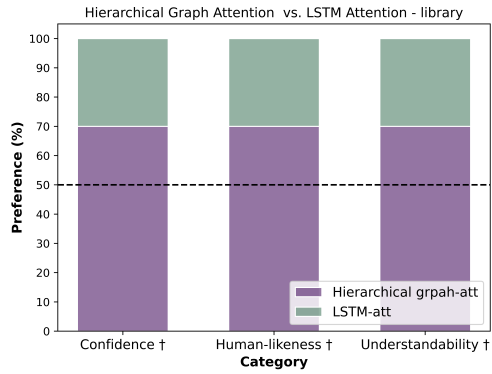
Figure 8: Eps. initial reward curves for the exploration strategies—*Game* and *IM* Reward

A.8 Immediate Explanation Evaluation

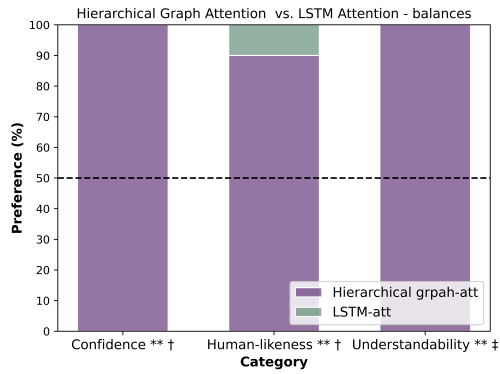
We plot the immediate explanation evaluation result per game in Figure 9.



(a) ZorkI



(b) library

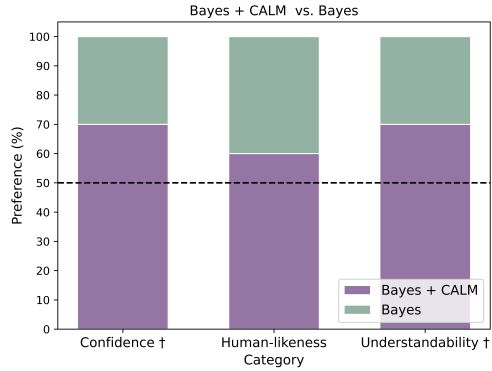


(c) balances

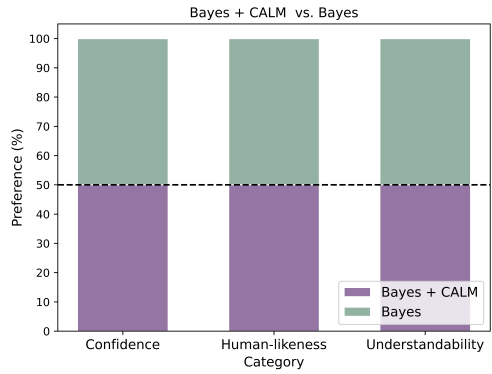
Figure 9: Human evaluation results comparing Hierarchical Graph Attention vs. LSTM Attention, * indicates $p < 0.05$, ** indicates $p < 0.01$, † indicates $\kappa > 0.2$ or fair agreement. ‡ indicates $\kappa > 0.4$ or moderate agreement.

A.9 Causal Explanation Ablation Study

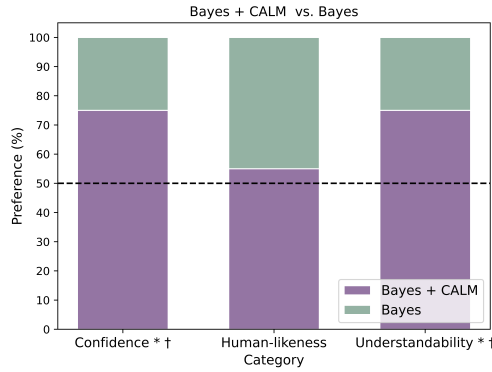
We plot the causal explanation ablation study result per game in Figure 10 and Figure 11.



(a) ZorkI

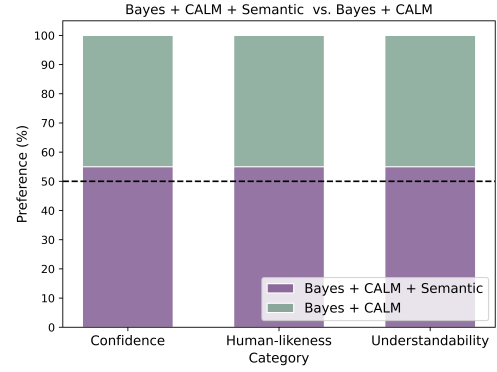


(b) library

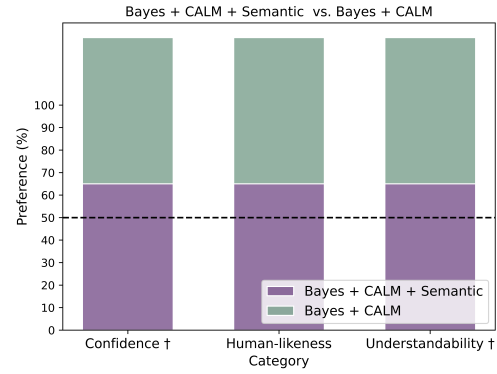


(c) balances

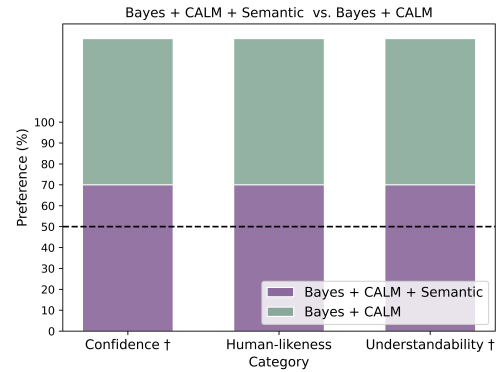
Figure 10: Human evaluation results on ablation study, * indicates $p < 0.05$, † indicates $\kappa > 0.2$ or fair agreement.



(a) ZorkI



(b) library



(c) balances

Figure 11: Human evaluation results on ablation study, * indicates $p < 0.05$, † indicates $\kappa > 0.2$ or fair agreement.

B Human Evaluation Details

B.1 Immediate Explanation Evaluation

We firstly ask participants to read an interactive game description and then ask them to answer a set of questions about this game to make sure they are qualified. They will also play a demo of an interactive text game and answer a question based on the game they played. The details can be found in Figure 12 and Figure 13. These questions are designed to improve the quality of human evaluation.

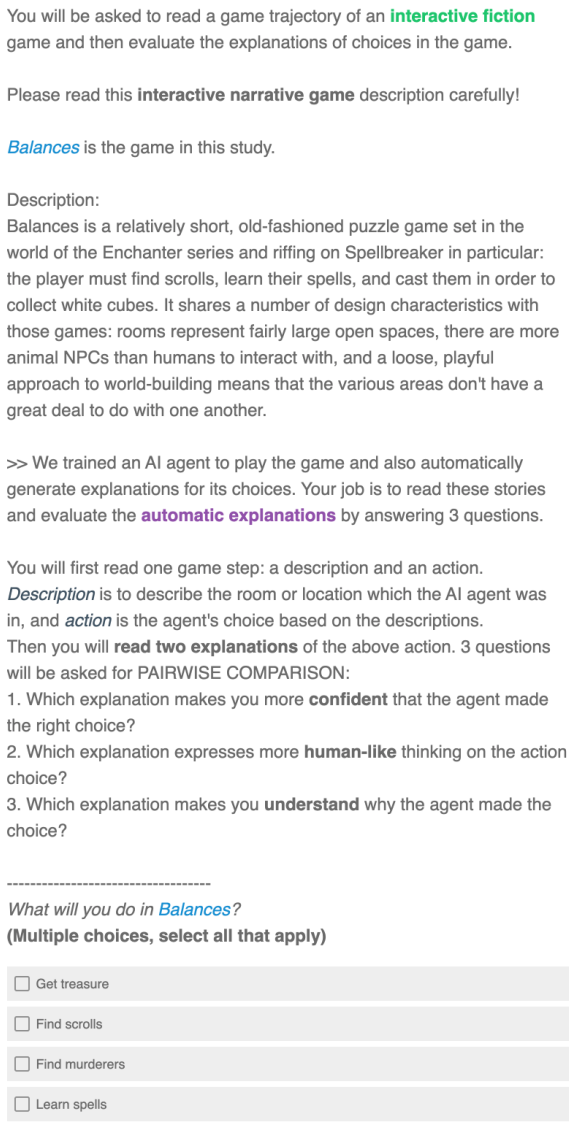


Figure 12: Screenshot of the human study instruction—game description.

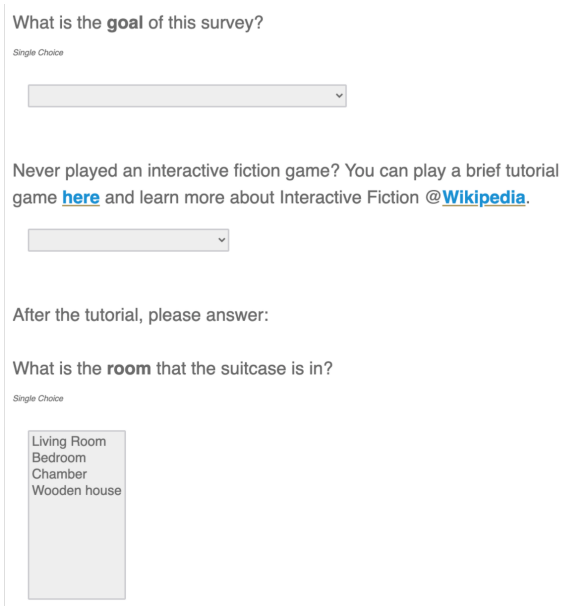


Figure 13: Screenshot of the human study instruction—task description.

Each participant reads a randomly selected subset of 10 explanation pairs (drawn randomly from a pool totaling 60 explanation pairs), generated by Hierarchical Graph Attention and LSTM attention explanation on three games in the Jericho benchmark, *zork1*, *library*, and *balances*. The following three questions are asked,

- Which explanation makes you more confident that the agent made the right choice?
- Which explanation expresses more human-like thinking on the action choice?
- Which explanation makes you understand why the agent made the choice?

B.2 Immediate vs. Causal Explanation Evaluation

Participants first read the full trajectory of the game (Figure 15) combined with step-by-step immediate explanations, along with summary of the game goal, and indicate how much they agree with the five statements on a Likert scale (Figure 16). The following five statements are used in human study.

- I am confident that I can get the same score as the agent when following this explanation.
- This explanation look like it was made by human.
- This explanation is easy to understand.
- I am able to understand why the agent takes this particular sequence of actions given what I know about the goal.
- This explanation is easy to read.

For each question, please rate which explanation best fits.

Description before the action:

Ramshackle Hut Until quite recently, someone lived here, you feel sure. Now the furniture is matchwood and the windows are glassless. Outside, it is a warm, sunny day, and grasslands extend to the low hills on the horizon.

You are carrying a spell book a silver coin a magic burin

My Spell Book gnusto spell copy a scroll into your spell book. frotz spell cause an object to give off light. yomin spell mind probe. rezrov spell open even locked or enchanted objects.

Two possible explanations of why the agent chose the action: **examine furniture** are as follows:

	This is the automatically generated explanation by the agent for why it performed this action. furniture is matchwood silver coin frotz spell	This is the automatically generated explanation by the agent for why it performed this action. I am in the Ramshackle Hut now. I have furniture furniture is interactable
1. Which explanation makes you more confident that the agent made the right choice?	☐	☐
2. Which explanation expresses more human-like thinking on the action choice?	☐	☐
3. Which explanation makes you understand why the agent made the choice?	☐	☐

Figure 14: Screenshot of the human study instruction.

Please read this transcript of a player playing *Library*:

Goal: Take the book

Description:

Second Floor Stacks

This cavernous room is filled with shelves as far as the eye can see. A doorway to the east is labelled "Computer Room", and the stairwell lies to the north. The door is unlocked but shut.

Action: *undo door*

Description:

You open the rare books door.

Action: *south*

Description after taking the above action:

Rare Books Room The shelves are nearly bare, although there is a complete set of the "New ork Times", a box labeled "Avalon", and several biographies of various computer game authors. The door out is to the north You can see a biography of Graham Nelson here.

Description:

Rare Books Room

The shelves are nearly bare, although there is a complete set of the "New ork Times", a box labeled "Avalon", and several biographies of various computer game authors. The door out is to the north You can see a biography of Graham Nelson here.

Action: *take all*

Description after taking the above action:

biography of Graham Nelson Taken. Your score has just gone up by five points.

Can you **summarize** what is happening in this transcript in less than 100 words?

Please answer the following questions about the game transcript given what you know about the goals of the game.

	Strongly Agree	Agree	Undecided	Disagree	Strongly Disagree
1. I am confident that I can get the same score as the agent when following this explanation.	☐	☐	☐	☐	☐
2. This explanation look like it was made by human.	☐	☐	☐	☐	☐
3. This explanation is easy to understand.	☐	☐	☐	☐	☐
4. I am able to understand why the agent takes this particular sequence of actions given what I know about the goal.	☐	☐	☐	☐	☐
5. This explanation is easy to read.	☐	☐	☐	☐	☐

Figure 16: Screenshot of Immediate vs. Causal Explanation Evaluation—Likert Scale.

Figure 15: Screenshot of Immediate vs. Causal Explanation Evaluation —Text Summary.