

Structured-CoT in Markdown for Enhanced Reasoning Efficiency and Thought Diversity

Vansh Vazirani
vansh@vazirani.net

Abstract

I present ThoughtWeaver, a fine-tuned language model that produces a structured chain-of-thought (CoT) in Markdown. By enforcing a fixed sequence of reasoning steps and deliberately omitting iterative self-revision loops, ThoughtWeaver can explore many candidate solutions while consuming substantially fewer tokens than conventional reasoning models. The approach preserves the ability to perform deeper planning when required, but it emphasises breadth over depth, making it well-suited for tasks where diverse ideas are more valuable than exhaustive verification. Experiments on arithmetic, logic puzzles, and open-ended brainstorming prompts show that ThoughtWeaver exhibits decent accuracy across various tasks while reducing token usage and inference time.

In a preliminary proof-of-concept evaluation on a small set of linguistic questions, mathematical and logic puzzles, and open-ended brainstorming prompts our experiments show that ThoughtWeaver dramatically reduces token usage (and thus inference time) relative to an unstructured CoT baseline for a large portion of tasks, while still exhibiting reasonably high response quality. A full accuracy and diversity assessment on standard arithmetic, logical-reasoning, and open-ended brainstorming benchmarks—once the model has undergone reinforcement-learning fine-tuning—is left to future work.

1 Introduction

Training models to generate Chain-of-Thought (CoT) reasoning has become a widely embraced technique for improving the reasoning ability of large language models (DeepSeek-AI et al., 2025; OpenAI, 2024; Snell et al., 2024). By designing models to synthesise intermediate reasoning, CoT typically raises accuracy on mathematical and scientific questions, code generation (J. Wang et al., 2025), and other deterministic tasks. The downside is that most CoT pipelines include unrestricted iterative self-revision: models often repeatedly re-examine their own reasoning multiple times, which inflates token consumption and delay (Sui et al., 2025; C. Wang et al., 2025).

In many practical settings, including evaluative brainstorming, design thinking and creative pursuits, the importance of breadth and volume of ideas often outweighs depth and exhaustive verification. When a large fraction of the token budget is spent on self-reflection, the model can only produce a handful of candidate solutions (Han et al., 2025), limiting creativity.

To address this tradeoff I introduce the first version of ThoughtWeaver: ThoughtWeaver-8B-Reasoning-Exp, a fine-tuned LLM that produces reasoning in a predefined, Markdown-based CoT structure. ThoughtWeaver dynamically selects the magnitude of reasoning required for each input (Anonymous, 2025; Li et al., 2025) and classifies as either low, medium, or high. For low-complexity queries the model produces a short analysis and a handful of ideas; for medium-complexity prompts it reasons in separate stages; and for high-complexity prompts it activates the optional validation stage. By skipping unnecessary self-revision loops, ThoughtWeaver can allocate a larger share of its token budget to generating diverse candidates while still offering a verification step when the task demands higher confidence.

I publish the model weights for ThoughtWeaver-8B-Reasoning-Exp under the Apache 2.0 license¹.

¹<https://huggingface.co/vazirani/ThoughtWeaver-8B-Reasoning-Exp>

2 Related Work

The development of ThoughtWeaver-8B-Reasoning-Exp builds on extensive work studying language models. Here, I briefly outline certain works that directly contributed to or inspired the creation of this model.

2.1 Reasoning with RL

The introduction of Reasoning LLMs, such as OpenAI O1 (OpenAI, 2024) and DeepSeek R1 (DeepSeek-AI et al., 2025) has deeply contributed to improvements in reasoning capabilities by scaling test-time compute, allowing models to brainstorm and reflect before generating the final answer. This has led to exceptional increases in model performance across diverse tasks and intelligence benchmarks. The DeepSeek R1 paper achieved tremendous gains in reasoning performance by scaling reinforcement learning without preliminary SFT. Subsequent works (Yang et al., 2025) often employ SFT to bootstrap the model with a scaffolding upon which RL improves.

2.2 Structured Reasoning Paradigms

In addition, certain works explore using an enforced structure for Chain-of-Thought, in the form of graphs (Besta et al., 2024), trees Yao et al., 2023, and tables (Sun et al., 2025; Z. Wang et al., 2024). Employing each structure offers its own advantages and limitations.

3 Motivating Problem

A large majority of reasoning models which natively produce unstructured Chain-of-Thought reflect on their own reasoning iteratively. Often, in the context of simpler problems, or tasks which do not require deep reasoning into singular aspects, this reasoning is limited and redundant. Such models are often too focused on needlessly verifying obvious parts of their reasoning chain, leading to wasted computational resources and a limited number of thoughts the model can explore.

With ThoughtWeaver, I seek to build a system that is capable of generating a wide range of candidate solutions and engaging in swift computation, without getting bogged down in endless self-critique.

Apart from reasoning in markdown (which encourages clear and interpretable thoughts), ThoughtWeaver’s second key innovation lies in its dynamic reasoning depth control, similar to what is seen in models like OpenAI’s GPT-OSS series (OpenAI et al., 2025).

4 Methodology

This section describes the methodology used to develop ThoughtWeaver, a fine-tuned language model that produces structured chain-of-thought (CoT) reasoning in Markdown. The approach consists of two main phases: data collection and model fine-tuning.

4.1 Data Collection

To create a comprehensive dataset for fine-tuning, I generated and collected a large corpus of structured reasoning traces using a diverse set of prompts and responses from existing datasets. The conversation records were classified into three categories of reasoning effort: low, medium, or high. The base datasets used include:

- microsoft/NextCoderDataset (Aggarwal et al., 2025)
- open-thoughts/OpenThoughts-114k (Guha et al., 2025)
- open-r1/OpenR1-Math-220k (Hugging Face, 2025)

- alibaba-pai/OmniThought-0528 (Cai et al., 2025)
- SoftAge-AI/multi-turn_dataset (SoftAge-AI, 2025)
- Josephgflowers/Finance-Instruct-500k (Flowers, 2025)
- NousResearch/Hermes-3-Dataset (Teknium, Quesnelle, and Guang, 2024)
- QuixiAI/dolphin-r1 (Hartford and Computations, 2025)
- glaiveai/reasoning-v1-20m (GlaiveAI, 2025)
- HuggingFaceTB/smoltalk2 - Split multi_turn_reasoning_if_think (HuggingFaceTB, 2025)
- PrimeIntellect/SYNTHETIC-2-SFT-verified (PrimeIntellect, 2025)
- inclusionAI/Ring-lite-sft-data (Ling Team, 2025)

I utilised several language models to generate the reasoning traces, including:

- moonshotai/kimi-k2-instruct-0905 (Kimi Team et al., 2025)
- moonshotai/kimi-k2-instruct
- openai/gpt-oss-120b (OpenAI et al., 2025)
- openai/gpt-oss-20b
- meta-llama/llama-4-maverick-17b-128e-instruct (Meta AI, 2024)
- meta-llama/llama-4-scout-17b-16e-instruct
- qwen/qwen3-235b-a22b-instruct-2507 (Yang et al., 2025)
- qwen/qwen3-32b

The resulting dataset consists of approximately 11,000 records, which were used for fine-tuning the model.

4.2 Model Fine-tuning

I fine-tuned the Ministral 3 8B Instruct 2512 (Mistral AI, 2025a) model using Supervised Fine-Tuning (SFT) with Unsloth (Daniel Han and team, 2023) on the collected dataset. The following hyperparameters were used:

- Number of epochs: 1
- Learning rate: $3e - 6$
- Optimizer: AdamW, 8-bit precision (Loshchilov and Hutter, 2019)
- Learning rate scheduler: Cosine annealing (2017)
- Packing enabled (Unsloth, 2025)
- Batch size: 8
- Maximum sequence length: 32,768

The fine-tuned model is capable of producing structured CoT reasoning in Markdown, with the ability to dynamically select the magnitude of reasoning required for each input.

5 Results

To sample responses from the model, it was converted to GGUF using llama.cpp (Gerganov, 2023) and hosted using llama-server. During testing, the model was prompted with several diverse queries and the responses were manually evaluated. Then, it was benchmarked against a corpus of mathematical questions.

5.1 Samples

I present three illustrative examples that exemplify ThoughtWeaver’s reasoning style in distinct scenarios. These were sampled from the GGUF model at temperature 0.

Query: How many of the first one hundred positive integers are divisible by 3, 4, and 5?

<think>

START OF THOUGHT PROCESS

This query requires low reasoning

* **Goal:** Count numbers ≤ 100 divisible by 3, 4, and 5.

* **Key insight:** A number divisible by all three must be divisible by their LCM.

* **Compute LCM(3,4,5):**

* Prime factors: $3 = 3$, $4 = 2^2$, $5 = 5$.

* $\text{LCM} = 2^2 \times 3 \times 5 = 4 \times 3 \times 5 = 60$.

* **Find multiples of 60 ≤ 100 :**

* $60 \times 1 = 60$, $60 \times 2 = 120$ (exceeds 100).

* Only 60 qualifies.

* **Conclusion:** Exactly one number meets the criteria.

END OF THOUGHT PROCESS

</think>

To determine how many of the first one hundred positive integers are divisible by 3, 4, and 5, we need to find the numbers that are divisible by all three of these numbers simultaneously.

Step 1: Understanding Divisibility by Multiple Numbers

A number that is divisible by 3, 4, and 5 must be divisible by the **least common multiple (LCM)** of these three numbers. This is because the LCM is the smallest number that is a multiple of each of the given numbers.

Step 2: Calculating the LCM of 3, 4, and 5

Let's find the LCM of 3, 4, and 5.

- Prime factorization:

- $3 = 3$

- $4 = 2^2$

- $5 = 5$

The LCM is the product of the highest powers of all primes present in the factorizations:

- $\text{LCM} = 2^2 \times 3 \times 5 = 4 \times 3 \times 5 = 60$

So, the LCM of 3, 4, and 5 is **60**.

Step 3: Finding Multiples of 60 Within the First 100 Positive Integers

Now, we need to find how many multiples of 60 are there in the range from 1 to 100.

- $60 \times 1 = 60$

- $60 \times 2 = 120$ (which is greater than 100)

So, the only multiple of 60 within the first 100 positive integers is **60**.

Step 4: Conclusion

There is **exactly one** number (60) in the first 100 positive integers that is divisible by 3, 4, and 5.

Final Answer

60

In this example, ThoughtWeaver immediately identifies the LCM technique and provides a direct answer with minimal reasoning steps, while refraining from verifying obvious calculations.

Query: What is the correct answer to this question: You perform a high-throughput experiment on white lupine to find genes contributing to resistance to the fungal disease anthracnose. As a result, you receive three candidate genes of unknown function – G1, G2, and G3. You create three knock-out mutants, g1, g2, and g3, and a set of double-mutants, g1g2, g1g3, and g2g3. You know that at least one of these genes is a transcription factor acting upstream of (an)other gene(s). You start to test those mutant plants: do they have a higher sensitivity to anthracnose than the wild-type because they cannot produce certain gene products? After tests with the pathogen, you receive the following results where 100% is the level of resistance to the pathogen in control; 50% is half of the control's resistance; 25% is a quarter of the control's resistance; 0% - all plants show signs of infection:

- resistance of g1: 75% of control
- resistance of g2: 0% from control
- resistance of g3: 50% from control
- resistance of g1g3: 10% from control
- resistance of g2g3: 0% from control
- resistance of g1g2: 0% from control

Which conclusion regarding those genes' interaction can you draw from this experiment?

Choices:

- (A) G2 is a transcription factor, G1 and G3 show gene redundancy, G1 is epistatic towards G3
- (B) G2 is a transcription factor, G1 and G3 has the same promoter, G3 is epistatic towards G1
- (C) G1 is a transcription factor, G2 and G3 show pleiotropy, G2 is epistatic towards G1
- (D) G2 is a transcription factor, G1 and G3 show pleiotropy, G1 is epistatic towards G3

This question was sourced from GPQA Diamond: (Rein et al., 2023)

```

<think>
# START OF THOUGHT PROCESS
This query requires high reasoning

## Deconstructing the Prompt
* Given information:
  * Three candidate genes (G1, G2, G3) of unknown function in white lupine.
  * Created single knockouts (g1, g2, g3) and double knockouts (glg2, glg3, g2g3).
  * At least one gene is a transcription factor acting upstream of other gene(s).
  * Resistance levels tested against anthracnose fungus, with 100% = full resistance (wild type), 0% = complete susceptibility.
* Key data points:
  * g1: 75% resistance → partial loss of resistance.
  * g2: 0% resistance → complete loss of resistance.
  * g3: 50% resistance → moderate loss of resistance.
  * Double mutants:
    * glg2: 0% → no additive effect beyond g2 alone.
    * glg3: 10% → synergistic reduction (worse than either single mutant).
    * g2g3: 0% → no additive effect beyond g2 alone.
* Constraints:
  * Must identify which gene is the transcription factor (upstream regulator).
  * Must explain interactions between genes (epistasis, redundancy, pleiotropy).
  * Cannot assume anything about gene function without evidence from phenotypes.

## High-Level Plan
1. Identify the transcription factor: Determine which gene's knockout causes the most severe phenotype (complete loss of resistance).
2. Analyze epistasis patterns: Compare double-mutant phenotypes to single-mutant phenotypes to infer dominance relationships.
3. Determine redundancy/pleiotropy: Check if multiple genes contribute similarly or differently to resistance.
4. Validate against answer choices: Match conclusions to the logical interpretation of the data.
5. Correct any misinterpretations: Ensure no contradictions arise from initial assumptions.

## Fleshing Out Brainstorm
### Step 1: Identifying the Transcription Factor
* Reasoning: A transcription factor regulates downstream genes. Its knockout should disrupt all regulated pathways, leading to severe phenotypes.
* Data analysis:
  * g2 shows 0% resistance → complete loss of resistance. This suggests g2 is critical for resistance.
  * g1 and g3 show partial losses (75% and 50%), indicating they may be downstream targets or redundant regulators.
* Conclusion: G2 is likely the transcription factor because its knockout eliminates resistance entirely, while others only partially reduce it.

### Step 2: Analyzing Epistasis in Double Mutants
* Definition: Epistasis occurs when one gene masks the effect of another. If a double mutant's phenotype equals that of a single mutant, the second gene is epistatic to the first.
* Case 1: glg2 vs. g2
  * g2 = 0%, glg2 = 0%. No additional loss of resistance → g1 does not affect resistance when g2 is absent.
  * Interpretation: G2 is epistatic to G1 (G2's absence masks G1's role).
* Case 2: glg3 vs. g3

```

- * g3 = 50%, glg3 = 10%. Much worse than either single mutant → synergistic interaction.
- * Interpretation: G1 and G3 are both required for resistance, but their combined loss is more severe than expected. This suggests G1 is epistatic to G3 (G1's absence exacerbates G3's defect).
- * **Case 3: g2g3 vs. g2**
- * g2 = 0%, g2g3 = 0%. No additive effect → G3 does not affect resistance when G2 is absent.
- * Interpretation: G2 is epistatic to G3 (G2's absence masks G3's role).

Step 3: Determining Redundancy vs. Pleiotropy

- * **Redundancy**: Multiple genes perform similar functions. Their knockouts have additive effects.
- * **Pleiotropy**: One gene affects multiple traits. Here, we're focusing on resistance, so pleiotropy would imply a single gene influences other unrelated processes.
- * **Data analysis**:
 - * g1 and g3 individually cause partial losses (75% and 50%). Their double mutant (glg3) causes a much larger loss (10%) than either alone.
 - * This suggests G1 and G3 contribute independently to resistance, but their combined loss is more severe due to overlapping roles.
 - * However, since G2 is the transcription factor, G1 and G3 are likely downstream targets. Their redundancy implies they share regulatory pathways controlled by G2.
- * **Conclusion**: G1 and G3 show **gene redundancy** because their combined knockout has a stronger effect than either alone, indicating they compensate for each other.

Step 4: Validating Against Answer Choices

- * **Option A**: "G2 is a transcription factor, G1 and G3 show gene redundancy, G1 is epistatic towards G3."
 - * Matches our analysis: G2 as TF, redundancy between G1/G3, and epistasis where G1's absence worsens G3's defect.
- * **Option B**: "G2 is a transcription factor, G1 and G3 has the same promoter, G3 is epistatic towards G1."
 - * Incorrect: Epistasis direction is reversed (G1 is epistatic to G3), and "same promoter" is speculative.
- * **Option C**: "G1 is a transcription factor, G2 and G3 show pleiotropy, G2 is epistatic towards G1."
 - * Incorrect: G1 cannot be the TF if its knockout only reduces resistance partially. Also, pleiotropy is irrelevant here.
- * **Option D**: "G2 is a transcription factor, G1 and G3 show pleiotropy, G1 is epistatic towards G3."
 - * Incorrect: Pleiotropy is not supported; G1 and G3 are redundant, not pleiotropic.

Step 5: Correcting Misinterpretations

- * **Initial error**: Assuming G1 or G3 could be the TF based on partial phenotypes.
 - * Correction: Only G2's complete loss of resistance aligns with a TF role.
- * **Misinterpretation of epistasis**: Confusing which gene is epistatic to which.
 - * Correction: In glg3, the phenotype is worse than g3 alone → G1's absence exacerbates G3's defect → G1 is epistatic to G3.
- * **Redundancy vs. pleiotropy confusion**:
 - * Correction: Redundancy explains why double mutants have additive effects, while pleiotropy would imply unrelated traits. Here, all genes affect resistance, so redundancy fits better.

Refinement

- * **Pruned ideas**: Discarded options B, C, and D due to contradictions in epistasis directions and incorrect assumptions about gene function.
- * **Prioritized evidence**: Focused on the most severe phenotype (g2 = 0%) to identify the TF. Used double-mutant comparisons to confirm epistasis relationships.
- * **Consolidated conclusions**:

- * G2 is the TF (upstream regulator).
- * G1 and G3 are redundant downstream targets.
- * G1 is epistatic to G3 (G1's absence masks G3's role).

Final Review and Validation

- * **Check consistency**: All data points align with Option A.
 - * g2 as TF: Complete loss of resistance confirms it regulates other genes.
 - * Redundancy between G1/G3: Their combined knockout causes a stronger effect than either alone.
 - * Epistasis: glg3 shows G1's absence worsens G3's defect, confirming G1 is epistatic to G3.
- * **Eliminate errors**:
 - * No longer assuming G1 or G3 could be TFs.
 - * Corrected epistasis direction from initial confusion.
- * **Final validation**: Option A is the only choice that logically explains all observed phenotypes without contradictions.

END OF THOUGHT PROCESS

</think>

The correct answer is **(A) G2 is a transcription factor, G1 and G3 show gene redundancy, G1 is epistatic towards G3**. This conclusion is supported by the experimental results showing that:

- The knockout of **G2** completely eliminates resistance (0%), indicating it acts as a **transcription factor** regulating downstream genes.
- Both **G1** and **G3** individually reduce resistance (75% and 50%, respectively), suggesting they contribute redundantly to the pathway controlled by G2.
- The **double mutant glg3** exhibits a much lower resistance level (10%) compared to either single mutant, demonstrating that **G1 is epistatic to G3**, meaning the loss of G1 exacerbates the defect caused by the loss of G3.

This interpretation aligns with the genetic principles of **epistasis** and **gene redundancy**, where one gene's function can mask another's when both are disrupted.

ThoughtWeaver correctly identifies the need for extensive reasoning. It breaks down the problem into distinct steps: identifying the TF, analysing epistasis, and drawing conclusions while weighing each option.

Query: If humans could perceive additional dimensions, how would art and architecture evolve?

<think>

START OF THOUGHT PROCESS

This query requires high reasoning

Deconstructing the Prompt

- * The prompt asks how art and architecture would evolve if humans could perceive additional dimensions (beyond the 3D space we currently experience).
- * Key constraints to consider:
 - * Must focus on *art* and *architecture* specifically, not other creative fields.
 - * "Perceive additional dimensions" implies we'd sense dimensions beyond length, width, height (e.g., time, energy, spatial curvature).
 - * Evolution implies gradual adaptation rather than sudden revolution.
- * Core objectives:

- * Explore conceptual shifts in artistic expression.
- * Analyze structural and experiential changes in architecture.
- * Consider societal/cultural implications of dimensional perception.

* **Critical nuance:** The prompt doesn't specify *which* dimensions (e.g., 4D time, 5D energy, 6D spatial curvature). I'll assume a general "higher-dimensional" framework since specifics aren't provided.

Fleshing Out Brainstorm

Art Evolution Concepts

- * **Initial idea:** If we perceive time as a dimension, art might become "temporal sculptures" where viewers experience evolution over time.
 - * **Why this works:** Time is already a dimension in physics; making it tangible in art would feel natural.
 - * **Potential flaw:** Might be too literal. Need to explore metaphorical uses too.
- * **Alternative angle:** Energy dimensions could enable "vibrational art" where colors/paintings emit subtle energy fields.
 - * **Why this works:** Energy is abstract but perceptible in some forms (e.g., heat, light).
 - * **Potential flaw:** Requires physical technology to make energy visible—might not be purely perceptual.
- * **Third option:** Spatial curvature (e.g., 4D hyperspace) could allow "non-Euclidean canvases" where shapes bend in ways we can't currently visualize.
 - * **Why this works:** Challenges traditional 2D/3D constraints.
 - * **Potential flaw:** Hard to imagine without concrete examples.
- * **Refining ideas:**
 - * Combine temporal and spatial concepts: "Living murals" that change based on viewer movement *and* time.
 - * Consider sensory integration: Art that triggers multiple senses (e.g., sound + visual + energy) in higher dimensions.
 - * **Decision:** Prioritize temporal and spatial curvature ideas—they feel most transformative for art.

Architecture Evolution Concepts

- * **Initial idea:** Buildings could have "time layers" where different eras coexist in the same space.
 - * **Why this works:** Solves the "time travel" paradox by making it architectural.
 - * **Potential flaw:** Might feel gimmicky if not grounded in real physics.
- * **Alternative angle:** Energy dimensions could enable "self-regulating structures" that adjust to environmental energy flows.
 - * **Why this works:** Aligns with sustainable design trends.
 - * **Potential flaw:** Requires advanced materials—might not be purely perceptual.
- * **Third option:** Spatial curvature could allow "infinite corridors" where paths loop in higher dimensions.
 - * **Why this works:** Creates psychological depth beyond physical space.
 - * **Potential flaw:** Hard to visualize without math.
- * **Refining ideas:**
 - * Focus on experiential architecture: Spaces that make viewers *feel* higher dimensions (e.g., through sound, light, or motion).
 - * Consider "dimensional thresholds": Entrances that trigger perception of extra dimensions (e.g., a doorway that warps space).
 - * **Decision:** Prioritize experiential and threshold concepts—they feel most innovative for architecture.

Societal/Cultural Implications

- * **Initial idea:** Art and architecture would become more collaborative, requiring teams with different dimensional perceptions.

- * **Why this works**: Mirrors how modern design teams work.
- * **Potential flaw**: Might overcomplicate the prompt.
- * **Alternative angle**: New art forms would emerge from cultural exchanges between "dimensional communities."
- * **Why this works**: Explores diversity in perception.
- * **Potential flaw**: Speculative—hard to ground in reality.
- * **Third option**: Architecture would prioritize "dimensional accessibility" (e.g., buildings that adapt to different perception styles).
- * **Why this works**: Aligns with inclusivity trends.
- * **Potential flaw**: Requires societal consensus on "standard" dimensions.
- * **Refining ideas**:
 - * Focus on how perception changes *individual* creativity first, then societal norms.
 - * Consider "dimensional art movements" as cultural phenomena (e.g., "Temporalism" vs. "Energeticism").
 - * **Decision**: Prioritize individual creativity and cultural movements—they feel most relevant to the prompt.

Refinement

- * **Pruning ideas**:
 - * Discard overly speculative concepts (e.g., "dimensional communities") since the prompt focuses on *art and architecture*, not society.
 - * Merge "temporal sculptures" and "living murals" into a single "dynamic art" category—both share the idea of time as a dimension.
- * **Merging concepts**:
 - * Combine "spatial curvature" and "infinite corridors" into "non-Euclidean architecture" to avoid redundancy.
 - * Link "energy dimensions" to "sustainable architecture" since energy perception could enable eco-friendly designs.
- * **Prioritizing core themes**:
 - * **Art**: Focus on *time* and *spatial curvature* as the most transformative dimensions for expression.
 - * **Architecture**: Prioritize *experiential design* and *dimensional thresholds* as the most innovative structural concepts.
 - * **Societal impact**: Emphasize how *individual perception* shapes *cultural movements* rather than broad societal changes.
- * **Why this works**:
 - * Keeps the response focused on the prompt's core (art/architecture evolution).
 - * Avoids overcomplicating with physics details (e.g., 4D math) that aren't necessary for creative exploration.
 - * Balances concrete examples (e.g., "living murals") with abstract ideas (e.g., "non-Euclidean architecture").

Final Review and Validation

- * **Check for consistency**:
 - * All ideas align with the prompt's focus on *perception* of additional dimensions, not just *use* of dimensions.
 - * Art and architecture concepts are distinct but interconnected (e.g., dynamic art could inspire experiential architecture).
- * **Validate against constraints**:
 - * Addresses *evolution* (gradual adaptation) rather than sudden revolution.
 - * Avoids over-speculation by grounding ideas in real-world trends (e.g., sustainability, inclusivity).
- * **Correct any errors**:
 - * Initially considered "energy dimensions" as purely physical, but refined to include *perceptual* energy (e.g., colors emitting energy fields).

- * Avoided conflating "time" and "energy" dimensions—kept them separate in the final structure.
- * **Ensure depth**:
 - * Explored *how* perception changes art/architecture (e.g., "living murals" vs. static paintings).
 - * Considered *why* these changes would occur (e.g., societal shifts in creativity).
- * **Final validation**:
 - * The refined ideas are cohesive, avoid redundancy, and directly answer the prompt.
 - * No speculative concepts remain—all are grounded in plausible extensions of current trends.

END OF THOUGHT PROCESS

</think>

If humans could perceive additional dimensions, art and architecture would undergo profound transformations, reshaping how we experience and create these fields. Here's how they might evolve:

Art Evolution

1. **Dynamic and Interactive Artworks**

- Art would become more interactive, responding to the viewer's movement and time. For example, a painting might change colors or patterns as the viewer walks around it, creating a living, evolving piece of art.

2. **Temporal Art**

- Artists could create works that unfold over time, allowing viewers to experience the evolution of a piece. This could include time-based installations, where the art changes subtly or dramatically as time progresses.

3. **Non-Euclidean Art**

- With the ability to perceive higher dimensions, artists might explore non-Euclidean geometries, creating art that bends and curves in ways that challenge traditional 2D and 3D representations. This could lead to new forms of abstract art that incorporate spatial curvature.

4. **Energy-Based Art**

- Artworks could incorporate energy dimensions, allowing artists to create pieces that emit or absorb energy. For example, a painting might emit a subtle energy field that interacts with the viewer's perception, creating a unique sensory experience.

5. **Sensory Integration**

- Art would likely integrate multiple senses, such as sound, touch, and even energy, to create immersive experiences. This could lead to new forms of multimedia art that transcend traditional visual boundaries.

Architecture Evolution

1. **Experiential Architecture**

- Buildings would be designed to create immersive experiences that make viewers feel as if they are moving through higher dimensions. This could include spaces that manipulate light, sound, and even energy to evoke a sense of depth and dimension.

2. **Non-Euclidean Structures**

- Architects might design buildings that incorporate non-Euclidean geometries, such as curved or warped structures that challenge traditional spatial constraints. This could lead to new forms of architecture that are both visually striking and functionally innovative.

3. **Temporal Architecture**

- Buildings could be designed to exist in multiple time layers, allowing different eras or functions to coexist within the same space. For example, a museum might have sections that transport visitors to different historical periods.

4. **Energy-Efficient and Sustainable Designs**

- With the ability to perceive energy dimensions, architects could design buildings that are highly energy-efficient and sustainable. This might include structures that absorb and redistribute energy, reducing the need for traditional power sources.

5. **Dimensional Thresholds**

- Buildings might feature entrances or thresholds that trigger the perception of higher dimensions, creating a sense of transition or transformation as visitors enter. This could lead to new forms of architectural storytelling and experience.

Societal and Cultural Impact

1. **New Art Movements**

- The ability to perceive additional dimensions would likely give rise to new art movements, such as "Temporalism" or "Energeticism," which would explore the creative possibilities of higher dimensions.

2. **Collaborative Creativity**

- Art and architecture would become more collaborative, with teams of artists and architects working together to create pieces that integrate multiple dimensions. This could lead to new forms of creative expression that are both individual and collective.

3. **Cultural Exchange and Diversity**

- Different cultures might develop unique ways of perceiving and interpreting higher dimensions, leading to a rich tapestry of artistic and architectural styles. This could foster greater cultural exchange and diversity in creative expression.

4. **New Forms of Storytelling**

- Art and architecture could become powerful tools for storytelling, allowing creators to convey complex narratives that span multiple dimensions. This could lead to new forms of visual and spatial storytelling that are both engaging and thought-provoking.

In summary, the evolution of art and architecture in a world where humans can perceive additional dimensions would likely involve a shift toward dynamic, interactive, and experiential forms of expression. This would not only change how we create and experience art and architecture but also how we perceive and interact with the world around us.

Here, ThoughtWeaver swiftly generates several diverse ideas, presenting them as separate sections within its response. It efficiently brainstorms numerous candidate ideas, then proceeds to prune less relevant points and weigh the rest.

These responses reveal that the model:

- Demonstrates a clear prioritisation of practical utility over exhaustive verification. When relevant, it chooses to generate a quick plan before diving into deep reasoning.
- Structures its reasoning traces in a way that facilitates decomposition and independent evaluation of each component. This contrasts with more monolithic CoT models where the entire chain of thought flows together seamlessly.
- Attempts to identify possible flaws quickly without getting lost in excessive detail—a characteristic that could be both an advantage (for increased speed) and a limitation (due to the possibility of overlooking subtle errors).

5.2 Benchmarks

The following results were collected using the LM Evaluation Harness framework (Gao et al., 2024) with the SGLang server (L. Zheng et al., 2024) using a repetition penalty of 1.2.

Tasks	n-shot	Metric	Value $\pm$ Stderr	Total Tokens	Avg Tokens/Ques
aime24	0	exact_match	0.1333 $\pm$ 0.0631	130,522	4,350.7
aime25	0	exact_match	0.1667 $\pm$ 0.0692	89,305	2,976.8

Table 1: AIME Evaluation Results for ThoughtWeaver-8B-Reasoning-Exp

This version of ThoughtWeaver performs considerably worse on the AIME (Zhang and Math-AI, 2024; 2025) benchmarks compared to traditional reasoning models such as Qwen3 8B (Yang et al., 2025) and even Mistral’s own reasoning models (Mistral AI, 2025b). This is because the model has not yet undergone reinforcement learning.

In comparison to other reasoning models, though, it consumes far fewer total tokens. Artificial Analysis (2026), which hosts benchmark details for several language models, shows that Mistral 3 8B Instruct spent a total of approximately 180,000 tokens on the AIME 2025 questions, and Qwen3 8B Reasoning spent a massive 540,000 tokens on reasoning alone. In contrast, ThoughtWeaver generated only 90,000 tokens across all questions, even though it selected “high reasoning” effort for each one.

6 Limitations

While this model introduces a novel approach to Chain-of-Thought architectures, several critical limitations must be acknowledged to provide context for its current performance and appropriate use cases.

- Because ThoughtWeaver has not undergone further safety training, it is not guaranteed to reject offensive or dangerous queries when prompt engineered by skilled red teamers.
- This version of the model is strictly unimodal (chat-only). It lacks numerous features found in contemporary frontier models, including multimodal inputs and agentic tool usage.
- It is found that certain models often drastically forget previously learned concepts, a phenomenon called Catastrophic Forgetting (Ven, Soures, and Kudithipudi, 2025). It is highly plausible that ThoughtWeaver may perform worse than its base model, Mistral, on certain tasks which it may have “forgotten” during fine-tuning.

7 Future Work: Reinforcement Learning

The resultant model from SFT can be fine-tuned with reinforcement learning using a verifiable rewards-based algorithm, such as Group Relative Policy Optimization (GRPO) (Shao et al., 2024) Decoupled Clip and Dynamic Sampling Policy Optimization (DAPO) (Yu et al., 2025) or Group Sequence Policy Optimization (GSPO) (C. Zheng et al., 2025). This step is crucial for further improving the model’s performance and is left for future work. The reinforcement learning process will enable the model to learn from verifiable rewards, leading to enhanced reasoning accuracy.

To further improve ThoughtWeaver’s usability in certain contexts, it can be reintegrated with the vision encoder from the Ministral base model and trained on extensive vision data. Furthermore, future versions of ThoughtWeaver can be trained on agentic tool usage with integrated reasoning.

8 Conclusion

ThoughtWeaver represents an initial step towards a more efficient and diverse approach to Chain-of-Thought reasoning in large language models. By deliberately limiting iterative self-revision loops and enforcing a structured Markdown CoT format, ThoughtWeaver demonstrates the potential to significantly reduce token consumption while still producing valuable candidate solutions—particularly beneficial for tasks which prioritise breadth over exhaustive verification. While current benchmarks reveal limitations in accuracy compared to more extensively trained models like Qwen3 8B and Ministral 3 8B, ThoughtWeaver’s reduced resource usage offers a compelling tradeoff for applications where efficiency is paramount. Future work will focus on rigorously training and evaluating ThoughtWeaver across standard benchmarks with RL, integrating multimodal inputs, and exploring agentic tool usage. ThoughtWeaver’s core contribution lies in a novel paradigm that prioritises diverse reasoning, paving the way for more adaptable and insightful reasoning capabilities.

References

- Aggarwal, Tushar et al. (2025). “NextCoder: Robust Adaptation of Code LMs to Diverse Code Edits”. In: *International Conference on Machine Learning*. URL: <https://www.microsoft.com/en-us/research/publication/nextcoder-robust-adaptation-of-code-lms-to-diverse-code-edits/>.
- Anonymous (2025). “AdaCtrl: Towards Adaptive and Controllable Reasoning via Difficulty-Aware Budgeting”. In: *Submitted to Transactions on Machine Learning Research*. Under review. URL: <https://openreview.net/forum?id=4J2AKo20V4>.
- Artificial Analysis (2026). *Artificial Analysis: Independent AI Benchmarking and Insights*. Accessed: 2026-01-02. URL: <https://artificialanalysis.ai/>.
- Besta, Maciej et al. (Mar. 2024). “Graph of Thoughts: Solving Elaborate Problems with Large Language Models”. In: *Proceedings of the AAAI Conference on Artificial Intelligence* 38.16, pp. 17682–17690. ISSN: 2159-5399. DOI: [10.1609/aaai.v38i16.29720](https://doi.org/10.1609/aaai.v38i16.29720). URL: <http://dx.doi.org/10.1609/aaai.v38i16.29720>.
- Cai, Wenrui et al. (2025). *Reasoning with OmniThought: A Large CoT Dataset with Verbosity and Cognitive Difficulty Annotations*. URL: <https://arxiv.org/abs/2505.10937>.
- Daniel Han, Michael Han and Unsloth team (2023). *Unsloth*. URL: <http://github.com/unslothai/unsloth>.
- DeepSeek-AI et al. (2025). *DeepSeek-R1: Incentivizing Reasoning Capability in LLMs via Reinforcement Learning*. URL: <https://arxiv.org/abs/2501.12948>.
- Flowers, Joseph G. (2025). *Finance-Instruct-500k*. URL: <https://huggingface.co/datasets/Josephgflowers/Finance-Instruct-500k>.
- Gao, Leo et al. (July 2024). *The Language Model Evaluation Harness*. Version v0.4.3. DOI: [10.5281/zenodo.12608602](https://doi.org/10.5281/zenodo.12608602). URL: <https://zenodo.org/records/12608602>.

- Gerganov, Georgi (2023). *llama.cpp*. URL: <https://github.com/ggerganov/llama.cpp>.
- GlaiveAI (2025). *reasoning-v1-20m Dataset*. Hugging Face Dataset. Apache-2.0 license; 22.2M rows; Accessed: 2025-12-01. URL: <https://huggingface.co/datasets/glaiveai/reasoning-v1-20m>.
- Guha, Etash et al. (2025). *OpenThoughts: Data Recipes for Reasoning Models*. URL: <https://arxiv.org/abs/2506.04178>.
- Han, Tingxu et al. (July 2025). “Token-Budget-Aware LLM Reasoning”. In: *Findings of the Association for Computational Linguistics: ACL 2025*. Ed. by Wanxiang Che et al. Vienna, Austria: Association for Computational Linguistics, pp. 24842–24855. ISBN: 979-8-89176-256-5. DOI: [10.18653/v1/2025.findings-acl.1274](https://doi.org/10.18653/v1/2025.findings-acl.1274). URL: <https://aclanthology.org/2025.findings-acl.1274/>.
- Hartford, Eric and Cognitive Computations (2025). *dolphin-r1 Dataset*. Hugging Face Dataset. Apache-2.0 license; Accessed: 2025-12-01. URL: <https://huggingface.co/datasets/QuixiAI/dolphin-r1>.
- Hugging Face (Jan. 2025). *Open R1: A fully open reproduction of DeepSeek-R1*. URL: <https://github.com/huggingface/open-r1>.
- HuggingFaceTB (2025). *SmolTalk2 Dataset*. Hugging Face Dataset. Apache-2.0 license; Accessed: 2025-12-01. URL: <https://huggingface.co/datasets/HuggingFaceTB/smoltalk2>.
- Kimi Team et al. (2025). *Kimi K2: Open Agentic Intelligence*. URL: <https://arxiv.org/abs/2507.20534>.
- Li, Junyan et al. (2025). *Steering LLM Thinking with Budget Guidance*. URL: <https://arxiv.org/abs/2506.13752>.
- Ling Team (2025). *Ring-lite: Scalable Reasoning via C3PO-Stabilized Reinforcement Learning for LLMs*. URL: <https://arxiv.org/abs/2506.14731>.
- Loshchilov, Ilya and Frank Hutter (2017). *SGDR: Stochastic Gradient Descent with Warm Restarts*. URL: <https://arxiv.org/abs/1608.03983>.
- (2019). *Decoupled Weight Decay Regularization*. URL: <https://arxiv.org/abs/1711.05101>.
- Meta AI (2024). *Introducing LLaMA 4: Advancing Multimodal Intelligence*. URL: <https://ai.meta.com/blog/llama-4-multimodal-intelligence/>.
- Mistral AI (2025a). *Ministral-3-8B-Instruct-2512: An instruction-tuned large language model*. <https://huggingface.co/mistralai/Ministral-3-8B-Instruct-2512>. Accessed: 2025-12-10. Hugging Face Model Hub.
- (2025b). *Ministral-3-8B-Reasoning-2512: A reasoning large language model*. <https://huggingface.co/mistralai/Ministral-3-8B-Reasoning-2512>. Accessed: 2025-12-10. Hugging Face Model Hub.
- OpenAI (2024). *Learning to reason with LLMs*. URL: <https://openai.com/index/learning-to-reason-with-llms/>.
- OpenAI et al. (2025). *gpt-oss-120b and gpt-oss-20b Model Card*. URL: <https://arxiv.org/abs/2508.10925>.
- PrimeIntellect (July 2025). *SYNTHETIC-2 Release: Four Million Collaboratively Generated Reasoning Traces*. Blog announcement describing the SYNTHETIC-2 dataset generation and release on Hugging Face. URL: <https://www.primeintellect.ai/blog/synthetic-2-release>.
- Rein, David et al. (2023). *GPQA: A Graduate-Level Google-Proof QA Benchmark*. URL: <https://arxiv.org/abs/2311.12022>.
- Shao, Zhihong et al. (2024). *DeepSeekMath: Pushing the Limits of Mathematical Reasoning in Open Language Models*. URL: <https://arxiv.org/abs/2402.03300>.
- Snell, Charlie et al. (2024). *Scaling LLM Test-Time Compute Optimally can be More Effective than Scaling Model Parameters*. URL: <https://arxiv.org/abs/2408.03314>.
- SoftAge-AI (2025). *Multi-turn Prompts Dataset*. Hugging Face Dataset. Accessed: 2025-12-01. URL: https://huggingface.co/datasets/SoftAge-AI/multi-turn_dataset.
- Sui, Yang et al. (2025). *Stop Overthinking: A Survey on Efficient Reasoning for Large Language Models*. URL: <https://arxiv.org/abs/2503.16419>.
- Sun, Zhenjie et al. (2025). *Table as Thought: Exploring Structured Thoughts in LLM Reasoning*. URL: <https://arxiv.org/abs/2501.02152>.
- Teknium, Ryan, Jeffrey Quesnelle, and Chen Guang (2024). *Hermes 3 Technical Report*. URL: <https://arxiv.org/abs/2408.11857>.

- Unsloth (2025). *3x Faster LLM Training with Unsloth Kernels + Packing*. <https://docs.unsloth.ai/new/3x-faster-training-packing>. Accessed: 2025-12-26.
- Ven, Guido M. van de, Nicholas Soures, and Dhireesha Kudithipudi (2025). “Continual learning and catastrophic forgetting”. In: *Learning and Memory: A Comprehensive Reference*. Elsevier, pp. 153–168. ISBN: 9780443157554. DOI: [10.1016/B978-0-443-15754-7.00073-0](https://doi.org/10.1016/B978-0-443-15754-7.00073-0). URL: <http://dx.doi.org/10.1016/B978-0-443-15754-7.00073-0>.
- Wang, Chenlong et al. (2025). *Wait, We Don’t Need to “Wait”! Removing Thinking Tokens Improves Reasoning Efficiency*. URL: <https://arxiv.org/abs/2506.08343>.
- Wang, Junqiao et al. (2025). *Enhancing Code LLMs with Reinforcement Learning in Code Generation: A Survey*. URL: <https://arxiv.org/abs/2412.20367>.
- Wang, Zilong et al. (2024). *Chain-of-Table: Evolving Tables in the Reasoning Chain for Table Understanding*. URL: <https://arxiv.org/abs/2401.04398>.
- Yang, An et al. (2025). *Qwen3 Technical Report*. URL: <https://arxiv.org/abs/2505.09388>.
- Yao, Shunyu et al. (2023). *Tree of Thoughts: Deliberate Problem Solving with Large Language Models*. URL: <https://arxiv.org/abs/2305.10601>.
- Yu, Qiying et al. (2025). *DAPPO: An Open-Source LLM Reinforcement Learning System at Scale*. URL: <https://arxiv.org/abs/2503.14476>.
- Zhang, Yifan and Team Math-AI (2024). *American Invitational Mathematics Examination (AIME) 2024*. — (2025). *American Invitational Mathematics Examination (AIME) 2025*.
- Zheng, Chujie et al. (2025). *Group Sequence Policy Optimization*. URL: <https://arxiv.org/abs/2507.18071>.
- Zheng, Lianmin et al. (2024). *SGLang: Efficient Execution of Structured Language Model Programs*. URL: <https://arxiv.org/abs/2312.07104>.