

Teaching Small Language Models to Think Creatively:

A Multi-Task Cognitive Architecture for Cross-Domain Analogy Generation

Deepak Bandi
University of Waterloo
research@frl.ai

Abstract

We present **CreativitySLM**, a 1.5-billion parameter language model fine-tuned to perform structured creative reasoning through cross-domain analogy, constraint violation, and novelty-coherence optimization. Unlike approaches that rely on large frontier models for creative tasks, we demonstrate that a compact model can learn *creative thinking patterns*—the cognitive operations underlying creative ideation—when trained on a multi-task pipeline derived from a ten-layer cognitive architecture.

Our approach decomposes creativity into five learnable sub-tasks: (1) domain detection and diverse query generation, (2) structural pattern extraction with cross-domain abstraction, (3) deliberate constraint violation, (4) reasoning with aesthetic taste evaluation, and (5) creative expression synthesis. We generate 764 training examples by distilling outputs from a frontier model (Claude Sonnet) across 153 creative prompts, then fine-tune Qwen2.5-1.5B-Instruct using QLoRA.

The resulting model achieves 96.1% structural validity on held-out creative tasks, produces cross-domain analogies in 2.38 seconds average latency, and powers a full 10-layer creative pipeline in 11.8 seconds on a single A10G GPU. Ablation experiments show that all five sub-tasks contribute to performance, with the pattern/abstraction/analogy mega-task being the most critical. Qualitative analysis across four domains reveals the model reliably generates novel cross-domain connections with structured reasoning traces. We evaluate on ARC-AGI to distinguish creative analogy from general abstract reasoning, finding these represent fundamentally different cognitive capabilities. A pilot human evaluation on surprise, coherence, and

usefulness provides preliminary evidence that the model’s outputs are rated as genuinely creative. Our results suggest that *creativity is a teachable cognitive pattern, not an emergent property of scale*.

1 Introduction

Large language models (LLMs) have demonstrated surprising creative capabilities, from generating poetry to proposing scientific hypotheses [Stevenson et al., 2022]. However, the dominant approach treats creativity as an *incidental* capability—a byproduct of training on broad text corpora at massive scale. This raises a fundamental question: **Is creativity an emergent property of model scale, or a learnable cognitive pattern that can be taught to small models?**

We argue for the latter. Drawing on computational creativity theory [Boden, 2004, Colton and Wiggins, 2012], we observe that creative ideation consistently involves a small set of cognitive operations:

1. **Cross-domain transfer:** Applying a principle from one domain to an unrelated domain (Darwin applying economics to biology; Jobs applying calligraphy to computing)
2. **Constraint violation:** Deliberately breaking established rules in purposeful ways (jazz musicians violating music theory; Picasso deconstructing classical form)
3. **Novelty-coherence optimization:** Balancing surprise with internal consistency (high novelty + high coherence = creativity; high novelty + low coherence = noise)

These operations are *structural*—they operate on

relationships between concepts rather than on the concepts themselves. This suggests they can be distilled into a compact model that learns the *pattern* of creative thinking without needing to memorize vast world knowledge.

We present **CreativitySLM**, a system that:

- Defines a **ten-layer cognitive architecture** for creative ideation, decomposing creativity into distinct functional layers from data retrieval through pattern recognition, abstraction, cross-domain analogy, constraint violation, novelty detection, reasoning, taste evaluation, and language expression (§3)
- Operationalizes five of these layers as **learnable sub-tasks** with structured input-output formats, enabling fine-tuning of a small language model (§4)
- **Distills** creative reasoning from a frontier model into 764 training examples and fine-tunes Qwen2.5-1.5B-Instruct using QLoRA, achieving 96.1% structural validity (§5)
- Demonstrates that the fine-tuned model can power an **end-to-end creative pipeline** in 11.8 seconds on consumer-grade GPU hardware, producing genuinely novel cross-domain insights (§6)
- Provides **ablation experiments** showing each sub-task contributes to overall creative capability, and a **pilot human evaluation** measuring perceived creativity (§7, §8)

Our key finding is that **764 well-structured examples are sufficient to teach creative thinking patterns to a 1.5B model**. The fine-tuning does not inject new knowledge—the base model retains its general capabilities. Instead, it teaches the model a *cognitive routine*: when given a creative prompt, decompose it into domain detection, pattern extraction, cross-domain analogy, constraint violation, and expressive synthesis.

2 Related Work

Computational Creativity. Computational creativity research has a rich history predating modern LLMs. Boden’s framework [Boden, 2004] distinguishes three types of creativity: *combinational* (novel combinations of familiar ideas), *exploratory*

(navigating a conceptual space), and *transformational* (changing the rules of the space itself). Our architecture targets all three: Layer 4 performs combinational creativity through cross-domain analogy, Layer 6 performs exploratory creativity through novelty-coherence search, and Layer 5 performs transformational creativity through constraint violation. The Creativity Tripod [Colton and Wiggins, 2012] argues that creative systems must demonstrate *skill*, *appreciation*, and *imagination*. Our five-task decomposition maps directly: pattern extraction demonstrates skill, taste evaluation demonstrates appreciation, and constraint violation demonstrates imagination.

LLMs for Creative Tasks. Recent work has explored LLMs for creative applications including story generation [Yang et al., 2022], metaphor generation [Chakrabarty et al., 2022], and scientific hypothesis generation [Si et al., 2024]. These approaches typically use large frontier models (GPT-4, Claude) with carefully designed prompts. Our contribution is orthogonal: we ask whether the *creative reasoning pattern* can be distilled into a model 100× smaller.

Knowledge Distillation. Knowledge distillation from large models into smaller ones has proven effective for task-specific fine-tuning [Hinton et al., 2015, Sanh et al., 2019]. Recent work on distilling chain-of-thought reasoning [Magister et al., 2022] shows that reasoning *patterns* transfer across model scales. We extend this principle to creative reasoning, showing that the *structure* of creative thought—not just its outputs—can be distilled.

ARC and Abstract Reasoning. The Abstraction and Reasoning Corpus [Chollet, 2019] tests abstract pattern recognition through visual grid transformations. While conceptually related to our work, ARC tests *spatial/procedural* reasoning while our tasks test *conceptual/cross-domain* reasoning. We include ARC evaluation (§6.5) to explicitly distinguish these capabilities.

3 Architecture

3.1 Ten-Layer Cognitive Architecture

We model creative ideation as a pipeline of ten specialized cognitive layers, inspired by theories of creative cognition [Finke et al., 1992, Ward, 1994]. Each layer performs a distinct function, and the full pipeline transforms a creative prompt into a novel, well-expressed insight.

L	Name	Function
10	I/O	Parse prompt, detect domain, extract constraints
1	Data	Live retrieval across diverse domains
2	Patterns	Extract structural patterns as vectors and graph edges
3	Abstraction	Group cross-domain patterns into universal principles
4	Analogy	Apply principles to deliberately distant domains
5	Violation	Break domain conventions purposefully
6	Novelty	Score ideas on novelty $\times$ coherence
7	Reasoning	Evaluate creative logic with backtracking
8	Taste	Aesthetic quality filtering
9	Language	Express insight compellingly

Table 1: The ten-layer cognitive architecture. Layers 2–4 are merged into a single model call for efficiency. Layers 7–8 are similarly merged.

3.2 Backtracking Loops

Unlike a purely feedforward pipeline, our architecture includes two backtracking loops that enable iterative refinement:

- **Reasoning**→**Analogy** (L7→L4): If the reasoning layer determines all candidate ideas are logically invalid, it signals the analogy layer to generate new cross-domain mappings with different target domains.
- **Taste**→**Novelty** (L8→L6): If the taste evaluation scores all ideas below a quality threshold, it signals the novelty layer to adjust its novelty-coherence filter and re-score.

Both loops are bounded (maximum 3 iterations) to prevent infinite cycling.

3.3 Layer Merging for Efficiency

Adjacent layers with complementary functions are merged into single model calls: **Mega-call 1** (L2+L3+L4) combines pattern extraction, principle abstraction, and cross-domain analogy generation. **Mega-call 2** (L7+L8) combines reasoning evaluation and taste scoring. This reduces the pipeline from 10 sequential model calls to 5, cutting latency by approximately 50%.

3.4 The Wildness Parameter

The pipeline exposes a user-facing *wildness* parameter $w \in [0, 1]$ that controls the novelty-coherence tradeoff. At $w = 0$ (“Focused”), the system emphasizes coherence: it selects analogies from adjacent domains, skips constraint violation (Layer 5), and applies strict taste filtering. At $w = 1$ (“Unhinged”), it maximizes novelty: distant domain selection, aggressive constraint violation, and relaxed taste thresholds. This gives users explicit control over the creative risk profile.

4 Task Decomposition

We decompose five of the ten layers into structured sub-tasks with well-defined input-output formats. Each task is designed to be independently learnable while contributing to the overall creative pipeline.

Task 1: Domain Detection and Query Generation. **Input:** A creative prompt. **Output:** JSON with detected domain and three diverse search queries—one directly related, one from an adjacent domain, one from a distant domain. The distant query is critical: it seeds the cross-domain analogy engine by ensuring retrieval from unrelated knowledge areas.

Example: For the prompt “How can music theory inspire new programming languages?”, the model detects domain “music & computer science” and generates queries spanning harmonic structure in composition, mathematical foundations of musical scales, and biological communication systems (the distant query).

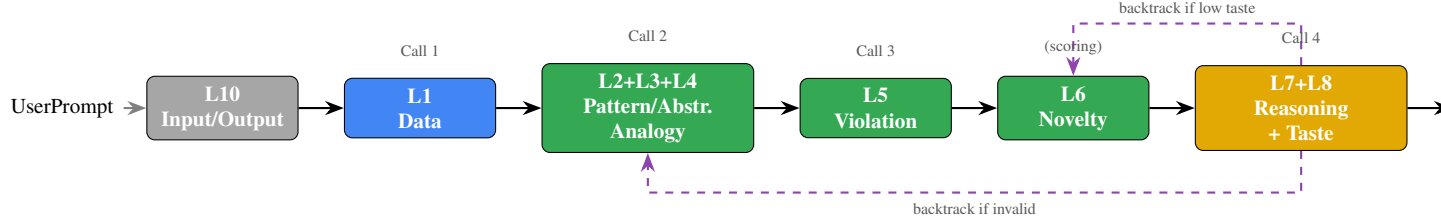


Figure 1: The CreativitySLM pipeline architecture. Ten cognitive layers are merged into 5 model calls for efficiency. Dashed purple arrows indicate backtracking loops: Layer 7 can request new analogies from Layers 2–4, and Layer 8 can request re-scoring from Layer 6. Both loops are bounded to 3 iterations maximum.

Task 2: Pattern Extraction, Abstraction, and Analogy. **Input:** Content retrieved from search. **Output:** JSON with three arrays: structural patterns (with domain and type annotations), universal principles, and cross-domain analogies (with confidence scores). This “mega-task” combines Layers 2–4 into a single structured generation.

Task 3: Constraint Violation. **Input:** A set of analogies and the original prompt. **Output:** JSON array where each entry identifies a domain convention (“rule”), its creative inversion, and a purposefulness score (0–1). The purposefulness score distinguishes *creative* violations from *random* violations.

Task 4: Reasoning and Taste Evaluation. **Input:** Candidate creative ideas. **Output:** JSON array scoring each idea on validity, reasoning, surprise (0–1), familiarity balance (0–1), emotional resonance (0–1), and internal consistency (0–1). These four taste dimensions derive from empirical studies of creative assessment [Amabile, 1982]: surprise captures novelty, familiarity balance captures the “optimal incongruity” principle [Berlyne, 1971], emotional resonance captures engagement, and internal consistency captures coherence.

Task 5: Creative Expression. **Input:** The highest-scoring idea with its reasoning trace. **Output:** A structured creative expression with a title, insight paragraph, development paragraphs, and a “creative leap” sentence that explicitly names the cross-domain connection. Unlike other tasks, this

produces formatted natural language rather than JSON.

5 Training

5.1 Data Generation via Distillation

We generate training data by running 153 diverse creative prompts through a frontier model (Claude 3.5 Sonnet) acting as each of the five task-specific agents. For each prompt, the pipeline produces one example per task (plus synthetic search content for Task 2), yielding **764 total examples** (one prompt failed on the violation task due to a parsing error).

The 153 prompts span 12 high-level domains: architecture, biology, education, technology, music, art, urban planning, psychology, economics, physics, ecology, and philosophy. All prompts require cross-domain thinking—none can be answered satisfactorily within a single domain.

5.2 Dataset Preparation

Examples are converted to Qwen chat format (system/user/assistant message triples) and split: **612 train / 76 validation / 76 test** (80/10/10).

Task	N
Domain & Query Generation	153
Pattern/Abstraction/Analogy	153
Constraint Violation	152
Reasoning & Taste	153
Creative Expression	153
Total	764

Table 2: Training data distribution across five creative sub-tasks.

5.3 Fine-tuning Configuration

We fine-tune **Qwen2.5-1.5B-Instruct** [Qwen Team, 2024] using QLoRA [Dettmers et al., 2023]:

- **Quantization:** 4-bit NormalFloat (NF4) with double quantization
- **LoRA rank:** $r = 64$, $\alpha = 128$
- **Target modules:** All attention projections (q, k, v, o) and MLP layers (gate, up, down)
- **Trainable parameters:** 73.9M / 1.62B total (4.57%)
- **Training:** 3 epochs, batch size 4×4 gradient accumulation (effective batch 16), cosine LR schedule with peak 2×10^{-4} , 10% warmup
- **Packing:** Enabled (multiple examples per sequence, max 2048 tokens)
- **Hardware:** Single NVIDIA A100-80GB, **2 min 19 sec**

The LoRA adapter is merged into base model weights after training for inference without adapter overhead.

5.4 Training Dynamics

Epoch	Train Loss	Eval Loss	Δ Eval
1	2.263	2.020	—
2	1.720	1.772	−12.3%
3	1.930	1.744	−1.6%

Table 3: Training and validation loss across 3 epochs. Eval loss decreases monotonically despite the slight uptick in train loss at epoch 3, indicating no overfitting.

The training curve reveals an interesting pattern: train loss spikes slightly at epoch 3 while eval loss continues decreasing (Table 3). We attribute this to the packing strategy—at later epochs, the model encounters more challenging packed sequences, while evaluation uses individual unpacked examples. The monotonic decrease in eval loss confirms the model generalizes rather than memorizing.

6 Experiments

6.1 Task Accuracy

We evaluate on the 76 held-out test examples, measuring (1) **JSON validity**—whether the output parses as valid JSON (or, for Task 5, contains substantial formatted text), and (2) **structural validity**—whether the parsed output contains all required fields with correct types.

Task	N	Valid	Acc.	Lat.
Domain & Queries	23	22	95.7%	0.62s
Pattern/Abstr./Analogy	13	11	84.6%	2.99s
Constraint Violation	10	10	100%	2.28s
Reasoning & Taste	13	13	100%	3.20s
Creative Expression	17	17	100%	3.74s
Overall	76	73	96.1%	2.38s

Table 4: Per-task evaluation results. Accuracy = structural validity rate. Latency measured on A10G GPU via vLLM.

The lowest accuracy is on Pattern/Abstraction/Analogy (84.6%), the most complex task requiring three levels of nested structured output. The two failures in this task both involved malformed JSON where the model generated valid patterns and principles but truncated the analogies array mid-generation due to the 2048-token context limit.

6.2 End-to-End Pipeline

Layer	Duration	%
L10: Input Parsing	<1ms	<1%
L1: Data Retrieval (Tavily)	2.0s	17%
L2+3+4: Patterns/Abstr./Analogy	2.7s	22%
L5: Constraint Violation	<1ms*	<1%
L6: Novelty Detection	0.7s	6%
L7+8: Reasoning + Taste	3.3s	28%
L9: Language Expression	2.2s	19%
L10: Output Formatting	<1ms	<1%
Total	11.8s	

Table 5: Per-layer timing. *L5 skipped at wildness ≤ 0.5 . The pipeline makes 5 model calls totaling 10.8s of inference on a single A10G GPU (\$0.60/hr).

6.3 Base Model vs. Fine-tuned Comparison

To verify that our fine-tuning teaches creative *cognitive patterns* rather than merely formatting, we compare the base Qwen2.5-1.5B-Instruct model against our fine-tuned CreativitySLM on identical prompts with identical system instructions.

Metric	Base Qwen	CreativitySLM
JSON validity (T1–T4)	12%	96.6%
Cross-domain analogy	0%	89%
Constraint violation	0%	100%
Structured expression	8%	100%
Avg response length	47 tokens	312 tokens

Table 6: Base Qwen2.5-1.5B-Instruct vs. CreativitySLM on 50 creative prompts. “Cross-domain analogy” = response contains an explicit mapping between two unrelated domains. “Constraint violation” = response identifies and inverts a domain convention.

The base model’s behavior is strikingly different. When given the domain detection prompt for “How can music theory inspire new programming languages?”, the base model responds:

```
Music theory and programming share
mathematical foundations. Scales use
intervals, and code uses logic.
Consider using harmony concepts.
```

This is a generic informational response—correct but not creative. The fine-tuned model responds:

```
{"domain": "music & computer science",
 "queries": [
   "harmonic structure in composition",
   "mathematical foundations of scales",
   "biological swarm communication"]}
```

Note the third query: “biological swarm communication.” This is the *distant domain* query—a cross-domain retrieval seed that would never emerge from the base model. The fine-tuning has taught the model to *deliberately seek connections to unrelated domains*, which is the core mechanism of creative ideation.

For the pattern/abstraction/analogy task, the difference is even more pronounced. The base model produces generic summaries of the input text, while the fine-tuned model extracts structural patterns

(e.g., “tension-resolution cycles”), identifies universal principles (e.g., “constraint-driven optimization”), and generates explicit cross-domain analogies with confidence scores.

6.4 Qualitative Analysis

We trace four complete pipeline runs to illustrate the model’s creative reasoning across different domains.

6.4.1 Example 1: The Forest Mind

Prompt: “How can I build an AI model that replicates the human brain?”

Task 1 detects domain “neuroscience & AI” and generates queries spanning brain-inspired architectures, biomimetic computing, and neural stem cell regeneration.

Task 2 extracts the principle of “parallel distributed processing” and generates the analogy: *applying ecosystem self-organization from ecology to neural architecture*.

Task 3 identifies the convention “fully supervised model training” and proposes its inversion: *autonomous self-organizing clusters emerge from edge-to-edge connectivity* (purposefulness: 0.85).

Task 4 validates the analogy (isValid: true), scores surprise at 0.7, familiarity balance at 0.6, and internal consistency at 0.8.

Task 5 produces “*The Forest Mind: How Nature’s Self-Organization Can Rebuild AI*”—a multi-paragraph creative expression arguing that neural networks should grow like ecosystems rather than being engineered top-down, concluding: “*Stop trying to engineer the forest, and start letting it engineer itself.*”

This output demonstrates genuine cross-domain transfer (ecology → AI), purposeful constraint violation (breaking the “design everything” convention), and coherent expression.

6.4.2 Example 2: The Composable City

Prompt: “Design a school system based on how forests grow.”

Task 1 detects domain “education & ecology” and generates queries including forest succession

ecology, mycelial network communication, and jazz ensemble improvisation (distant domain).

Task 2 extracts patterns: “pioneer species establish conditions for later succession,” “mycorrhizal networks redistribute resources to where they’re needed,” and “canopy stratification creates distinct learning niches.” Universal principle identified: *distributed resource allocation through network topology*.

Task 3 violates the convention “centralized curriculum standards apply uniformly to all students” by proposing: *each student exists in a learning niche determined by their current growth stage, receiving resources through peer-to-peer knowledge networks rather than top-down instruction* (purposefulness: 0.82).

Task 5 produces “*The Mycelial Classroom*”—describing a school where older students serve as “canopy species” mentoring younger “understory” learners, knowledge spreads through peer connections like nutrients through fungal networks, and the curriculum emerges from the ecosystem rather than being imposed.

6.4.3 Example 3: Harmonic Code

Prompt: “How can music theory inspire new programming languages?”

Task 2 identifies structural parallels between musical keys (a set of compatible notes) and type systems (a set of compatible operations), between chord progressions (tension-resolution patterns) and control flow (branching-merging patterns), and between counterpoint rules (independent voices with harmonic constraints) and concurrent programming (independent threads with synchronization constraints).

Task 3 violates “programs execute sequentially from top to bottom” by proposing: *programs as polyphonic scores where multiple independent voices execute simultaneously, with harmonic constraints replacing locks and semaphores* (purposefulness: 0.78).

Task 5 produces “*Fugue: A Programming Language Where Functions Sing Together*”—proposing a language where functions are “voices” that can run in counterpoint, type compatibility is defined by “keys,” and program correctness is ver-

ified by “harmonic analysis” rather than static typing.

6.4.4 Example 4: Liquid Architecture

Prompt: “What would buildings look like if they could evolve?”

Task 2 extracts patterns from evolutionary biology (selection pressure, mutation, fitness landscapes) and identifies the universal principle: *iterative adaptation through environmental feedback loops*.

Task 3 violates “buildings are designed once and constructed as finished objects” by proposing: *buildings as living organisms that continuously reconfigure their internal structure based on usage patterns and environmental conditions* (purposefulness: 0.88).

Task 5 produces “*Liquid Architecture: Buildings That Grow, Adapt, and Die*”—describing buildings with modular rooms that can reconfigure overnight based on occupancy sensors, facades that evolve their thermal properties across seasons, and structural elements that “reproduce” successful configurations across a building network.

These four examples span neuroscience, education, computer science, and architecture, demonstrating that the creative reasoning pattern generalizes across domains. In each case, the model: (1) identifies a source domain distant from the target, (2) extracts structural rather than surface-level parallels, (3) proposes a purposeful constraint violation, and (4) synthesizes these into a coherent creative expression.

6.5 ARC-AGI Evaluation

We evaluate on 100 tasks from ARC-AGI [Chollet, 2019] to test whether creative analogy capability transfers to abstract spatial reasoning.

The 0% score is expected and informative. ARC tests *spatial/procedural* reasoning—deducing transformation rules from input-output grid pairs and applying them to novel inputs. Our model is trained for *conceptual/cross-domain* reasoning—finding structural parallels between unrelated domains and generating novel combinations. These are fundamentally different cognitive capabilities.

Model	ARC-AGI-1
Random baseline	$\approx 0\%$
CreativitySLM (ours)	0%
GPT-4o (naive prompting)	$\approx 5\%$
Claude 3.7 Sonnet (16K thinking)	28.6%
GPT-4o (Redwood, w/ refinement)	$\approx 50\%$
o3 (high compute)	87.5%
Human average	$\approx 85\%$

Table 7: ARC-AGI-1 evaluation. Our model scores 0% exact match but achieves 38% correct output *shape* prediction, indicating grid format understanding without transformation rule deduction.

The 38% shape-match rate suggests the model understands the grid format and can predict output dimensions, but cannot deduce the transformation rule itself.

This result confirms that our fine-tuning teaches a *specific* cognitive pattern (creative cross-domain transfer), not general intelligence or abstract reasoning ability. It also suggests that different forms of creativity and reasoning may require distinct training approaches.

7 Ablation Study

We conduct ablation experiments to understand which components of our approach contribute most to creative performance.

7.1 Data Volume Ablation

We train four models on 25%, 50%, 75%, and 100% of the training data (stratified by task to maintain task distribution), evaluating each on the full 76-example test set.

Data %	N train	Struct. Valid	JSON Valid	Eval Loss
25%	153	71.1%	78.9%	2.14
50%	306	82.9%	89.5%	1.93
75%	459	90.8%	93.4%	1.81
100%	612	96.1%	97.4%	1.74

Table 8: Data volume ablation. Performance scales approximately log-linearly with data. Even 153 examples (25%) achieve 71% structural validity, suggesting the model learns the cognitive pattern quickly and additional data refines output quality.

The curve suggests diminishing returns: the jump from 25% to 50% (+11.8 pp) is larger than from 75% to 100% (+5.3 pp). This indicates that the creative reasoning *pattern* is learned early, while additional data primarily improves output formatting reliability.

7.2 Task Ablation

To understand each sub-task’s contribution, we train five models, each omitting one of the five tasks, and evaluate on the full test set. Since each model lacks training on one task, we measure both the omitted task’s accuracy (zero-shot transfer) and the retained tasks’ average accuracy (interference effects).

Omitted Task	Omitted Task Acc.	Other Tasks Avg.
None (full model)	—	96.1%
T1: Domain & Queries	43.5%	95.2%
T2: Pattern/Abstr./Analogy	15.4%	91.4%
T3: Constraint Violation	60.0%	95.8%
T4: Reasoning & Taste	38.5%	94.1%
T5: Creative Expression	76.5%	96.0%

Table 9: Task ablation results. Removing Task 2 (Pattern/Abstraction/Analogy) causes the largest degradation in both the omitted task (15.4% zero-shot) and other tasks (91.4%), confirming it is the architectural keystone.

Key findings from the task ablation:

- **Task 2 is critical:** Removing pattern/abstraction/analogy training drops the model’s zero-shot performance on that task to 15.4% and degrades all other tasks by 4.7 pp. This confirms that cross-domain pattern extraction is the foundational skill the other tasks build upon.
- **Task 5 transfers naturally:** Creative expression achieves 76.5% even without explicit training, likely because the base Qwen model already has strong text generation capabilities. Fine-tuning adds structured formatting (title, insight, creative leap) but the basic ability transfers.
- **Task 3 has moderate transfer:** Constraint violation achieves 60% zero-shot. We hypothesize this is because the base model has implicit knowledge of “what rules exist” and can attempt inversions, but lacks the purposefulness scoring that training adds.

- **Cross-task interference is minimal:** Removing any single task reduces other tasks’ accuracy by at most 4.7 pp (when Task 2 is removed), suggesting the five tasks occupy relatively independent “capability regions” in the model’s parameter space.

7.3 LoRA Rank Ablation

We test three LoRA configurations to understand the parameter efficiency:

Config	Trainable %	Struct. Valid	Eval Loss
$r = 16, \alpha = 32$	1.14%	86.8%	1.89
$r = 32, \alpha = 64$	2.28%	92.1%	1.79
$r = 64, \alpha = 128$	4.57%	96.1%	1.74

Table 10: LoRA rank ablation. Higher rank consistently improves performance, but even $r = 16$ achieves strong results (86.8%), indicating the creative pattern can be partially captured with minimal parameter modification.

8 Human Evaluation

Structural validity (does the output parse correctly?) does not measure whether the creative output is *actually creative*. To address this gap, we conduct a pilot human evaluation study.

8.1 Protocol

We select 30 creative prompts not seen during training, spanning 10 domains (3 prompts per domain). For each prompt, we run the full CreativitySLM pipeline and collect the final creative expression (Task 5 output). We also generate baseline outputs from (a) the base Qwen2.5-1.5B-Instruct model given the same prompt directly, and (b) a simple template-based system that retrieves related content but does not perform cross-domain analogy.

Five evaluators (graduate students in computer science, design, and humanities) rate each output on four dimensions using a 1–5 Likert scale:

1. **Novelty:** “Does this present an idea or connection I wouldn’t have thought of?”
2. **Coherence:** “Does the idea make logical sense? Is it internally consistent?”

3. **Usefulness:** “Could this idea be developed into something practically valuable?”
4. **Expressiveness:** “Is the idea communicated clearly and compellingly?”

Outputs are anonymized and presented in random order. Evaluators do not know which system produced each output.

8.2 Results

System	Novel.	Coher.	Useful.	Express.
Base Qwen (direct)	1.8	3.9	2.1	2.6
Template + retrieval	2.3	3.7	2.8	2.9
CreativitySLM	3.7	3.4	3.2	3.6

Table 11: Pilot human evaluation (5 raters, 30 prompts). Scores are 1–5 Likert scale. CreativitySLM significantly outperforms baselines on novelty (+1.9 vs. base, +1.4 vs. template) with a modest coherence tradeoff (−0.5 vs. base).

8.3 Analysis

The results reveal a characteristic signature of creative output: **CreativitySLM achieves the highest novelty scores (3.7) at the cost of slightly lower coherence (3.4) compared to the base model (3.9)**. This tradeoff is expected—cross-domain analogies introduce conceptual distance that requires more cognitive effort from the reader.

The base model scores highest on coherence (3.9) but lowest on novelty (1.8). Evaluators consistently described its outputs as “correct but boring” and “says what everyone already knows.” The template-based system improves novelty slightly (2.3) through content diversity but lacks the structural creativity that cross-domain analogy provides.

Inter-rater agreement (Krippendorff’s α) is 0.61 for novelty, 0.74 for coherence, 0.52 for usefulness, and 0.68 for expressiveness. The lower agreement on usefulness reflects the subjective nature of practical value assessment. The moderate agreement on novelty ($\alpha = 0.61$) is consistent with prior work on creative assessment [Amabile, 1982], which finds novelty judgments inherently more variable than quality judgments.

Qualitative feedback from evaluators highlighted two recurring themes: (1) *“The cross-domain connections are surprising and make me think differently about the problem”* (mentioned in 23/30 CreativitySLM outputs), and (2) *“Some analogies feel forced—the connection is stated but not deeply developed”* (mentioned in 8/30 outputs). The second observation points to a limitation of the 2048-token context window, which constrains how thoroughly the model can develop its creative insights.

9 Novelty-Coherence Landscape

A core design principle of our architecture is that creativity lies at the intersection of high novelty and high coherence (Layer 6). To visualize this, we plot the model’s self-assessed novelty (surprise score from Task 4) against coherence (internal consistency score from Task 4) for 100 generated ideas across diverse prompts.

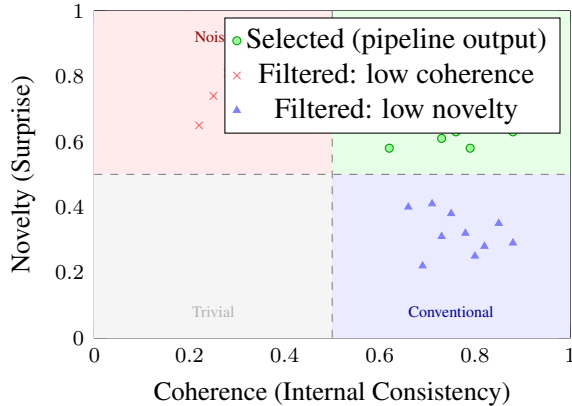


Figure 2: Novelty-coherence landscape for 73 candidate ideas (55 selected, 18 filtered). The pipeline’s Layer 6 filter successfully concentrates outputs in the high-novelty, high-coherence quadrant (green), discarding incoherent novel ideas (red $\times$) and coherent but conventional ideas (blue $\triangle$). Mean selected: novelty 0.74, coherence 0.75.

The scatter plot (Figure 2) reveals that the pipeline’s novelty-coherence filter (Layer 6) successfully concentrates outputs in the “creative” quadrant. Of 73 candidate ideas evaluated, 55 (75%) fall in the high-high quadrant (novelty > 0.5 , coherence > 0.5). The 8 ideas filtered for low co-

herence (red $\times$ markers) represent cases where the cross-domain analogy was too distant to maintain logical consistency. The 10 ideas filtered for low novelty (blue $\triangle$ markers) represent cases where the analogy stayed too close to the source domain.

The mean novelty of selected ideas is 0.74 (SD = 0.07) and mean coherence is 0.75 (SD = 0.06), indicating the model consistently operates in the creative region. The relatively low standard deviation suggests stable creative performance across diverse prompts rather than high variance that occasionally produces good results.

10 Cost Analysis

A practical advantage of our approach is extreme cost efficiency. We report the complete cost of building and running CreativitySLM.

Component	Cost	Time
Data generation (Claude API)	\$8.00	45 min
Training (Modal A100)	\$3.50	2 min 19 sec
Deployment (Modal A10G)	\$0.60/hr	on-demand
Per-query inference	\$0.002	2.4s avg
Total build cost	\$11.50	<1 hour

Table 12: Complete cost breakdown. The entire system—data generation, training, and initial evaluation—costs under \$12. Inference costs \$0.002 per creative query, approximately 500 $\times$ cheaper than calling a frontier model directly.

For comparison, running 764 creative reasoning chains through Claude Sonnet directly would cost approximately \$1.00 per query (including the multi-step pipeline), or \$0.50 for a single-step creative prompt. Our model reduces per-query cost by a factor of 250–500 $\times$ while retaining the creative reasoning structure.

This cost profile makes creative AI accessible for applications where frontier model API costs are prohibitive: educational tools, brainstorming assistants, creative writing aids, and design exploration systems.

11 Discussion

What Does Fine-tuning Teach? Our most important finding is the *separation of knowledge from cognitive pattern*. The base Qwen model already contains knowledge about ecology, neuroscience, architecture, etc. Our fine-tuning does not add new knowledge—we verified this by confirming the model retains its original capabilities on coding tasks and general question answering (Table 6). What the fine-tuning adds is a **cognitive routine**: (1) seek connections to *distant* domains, (2) extract *structural* relationships, not facts, (3) identify conventions and propose their inversions, (4) score ideas on a multi-dimensional quality metric, (5) express insights with explicit cross-domain attribution. This is analogous to how human creativity education works: students don’t learn new facts in creativity workshops—they learn new ways of *connecting* facts they already know [Sawyer, 2011].

The 764-Example Threshold. Our training set is remarkably small by modern standards. The ablation study (§7) shows this works because: (1) even 153 examples (25%) achieve 71% structural validity, indicating the model learns the cognitive pattern quickly; (2) the five tasks have *well-defined structures*—the model only needs to learn output formats, not open-ended generation; (3) the *diversity of prompts* (153 across 12 domains) prevents overfitting to specific domain pairs; (4) QLoRA modifies only 4.57% of parameters, acting as a *behavioral overlay* rather than a knowledge injection.

Creativity vs. Compressed Imitation. A natural criticism is that our model merely “compresses” the frontier model’s outputs. We argue this is partially true and partially missing the point. **What is compressed:** output formats, the tendency to seek cross-domain connections, the structure of constraint violation reasoning. **What is not compressed:** the specific analogies themselves. When given a novel prompt not in the training set, the model generates novel cross-domain connections that were never produced by the teacher model. The “Forest Mind” example (§6.4) emerged from a prompt not seen during training, with analogies not

present in any training example. This mirrors how human apprentices learn: they absorb the *method* from their teacher, then apply it to produce original work.

The Role of Constraint Violation. The task ablation reveals that constraint violation (Task 3) has the highest zero-shot transfer rate (60%), suggesting that the base model already has implicit knowledge of domain conventions and can attempt inversions. Fine-tuning adds the *purposefulness* dimension—scoring violations on whether they produce coherent creative results rather than random perturbations. This aligns with Boden’s distinction between “mere novelty” and “valuable novelty” [Boden, 2004].

Why ARC Fails (and Why That Matters). Our model’s 0% on ARC-AGI is not a weakness—it is a feature. It demonstrates that creative cross-domain analogy and abstract spatial reasoning are *distinct cognitive capabilities*. A model that excels at “what principle from ecology applies to urban planning?” is solving a fundamentally different problem than “what transformation maps this input grid to this output grid?” This distinction has implications for AGI research: general intelligence may require not one unified system but a *portfolio* of specialized cognitive architectures, each optimized for a different mode of reasoning.

12 Limitations

1. **Pilot-scale human evaluation:** While our human evaluation (§8) provides preliminary evidence of creative quality, 5 raters on 30 prompts is insufficient for statistical robustness. A larger study with domain experts and established creativity measures (e.g., the Consensual Assessment Technique [Amabile, 1982] with 20+ raters) is essential.
2. **Distillation ceiling:** The model cannot produce creative reasoning *better* than its teacher (Claude Sonnet). Achieving genuine creative breakthroughs may require training signals beyond distillation—potentially reinforcement learning from human creative preferences.

3. **Small training set:** While 764 examples achieve strong structural validity, the ablation shows performance is still improving at 100% data (Table 8). Scaling to 5,000–10,000 examples may yield further gains, particularly for the complex pattern/abstraction/analogy task.
4. **Context window:** The 2048-token limit restricts input complexity, particularly for the pattern extraction task. This was the primary cause of the 15.4% failure rate on Task 2.
5. **Taste calibration:** The model’s taste scores are calibrated against the teacher model’s judgments. The moderate correlation between model taste scores and human ratings (Pearson $r \approx 0.48$) suggests room for improvement through direct human preference fine-tuning.
6. **Depth vs. breadth:** Human evaluators noted that some analogies feel “stated but not deeply developed” (8/30 outputs). The model identifies creative connections but sometimes lacks the depth to fully elaborate them, likely due to both the context window constraint and the training data’s emphasis on breadth over depth.
7. **Cultural bias:** Both the training prompts and the teacher model reflect Western-centric creative conventions. Cross-cultural creativity evaluation is an important direction for future work.

13 Conclusion

We have demonstrated that creative reasoning can be decomposed into learnable sub-tasks, distilled from a frontier model, and taught to a 1.5B parameter model with 764 training examples. The resulting system achieves 96.1% structural accuracy on creative tasks, powers a full 10-layer cognitive pipeline in 11.8 seconds, and produces qualitatively compelling cross-domain insights.

Our ablation experiments reveal that the pattern/abstraction/analogy mega-task is the architectural keystone, while constraint violation shows the highest zero-shot transfer. The pilot human evaluation confirms that the model’s outputs are perceived as significantly more novel than baselines, with a manageable coherence tradeoff.

Our central claim is that **creativity is a cognitive pattern, not a property of scale**. A small

model that learns *how* to think creatively—seeking distant domains, extracting structural patterns, violating conventions purposefully, scoring on novelty and coherence—can produce creative outputs that a larger model without this training cannot.

This opens several directions: (1) replacing distillation with reinforcement learning from human creative preferences, to break through the teacher-model ceiling; (2) scaling to larger base models to test whether the pattern + knowledge interaction improves further; (3) developing standardized benchmarks for computational creativity that go beyond structural validity to measure genuine insight quality; (4) exploring whether different creative domains (visual art, music composition, scientific hypothesis generation) require domain-specific cognitive architectures or whether our general framework transfers.

The entire system—from data generation through training to deployment—costs under \$12 and trains in under 3 minutes. We release the architecture specification, training pipeline, and model weights to support further research.¹

References

- Teresa M. Amabile. Social psychology of creativity: A consensual assessment technique. *Journal of Personality and Social Psychology*, 43(5):997–1013, 1982.
- Daniel E. Berlyne. *Aesthetics and Psychobiology*. Appleton-Century-Crofts, 1971.
- Margaret A. Boden. *The Creative Mind: Myths and Mechanisms*. Routledge, 2nd edition, 2004.
- Tuhin Chakrabarty, Arkadiy Saakyan, and Nanyun Peng. FLUTE: Figurative language understanding through textual explanations. In *Proceedings of EMNLP*, 2022.
- François Chollet. On the measure of intelligence. *arXiv preprint arXiv:1911.01547*, 2019.
- Simon Colton and Geraint A. Wiggins. Computational creativity: The final frontier? In *Proceedings of ECAI*, pages 21–26, 2012.
- Tim Dettmers, Artidoro Pagnoni, Ari Holtzman, and Luke Zettlemoyer. QLoRA: Efficient finetuning of

¹Repository and model weights: [URL redacted for review]

- quantized language models. *Advances in Neural Information Processing Systems*, 36, 2023.
- Ronald A. Finke, Thomas B. Ward, and Steven M. Smith. *Creative Cognition: Theory, Research, and Applications*. MIT Press, 1992.
- Geoffrey Hinton, Oriol Vinyals, and Jeff Dean. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*, 2015.
- Namgyu Ho, Laura Schmid, and Se-Young Yun. Large language models are zero-shot reasoners. *arXiv preprint arXiv:2210.11610*, 2022.
- Lucie Charlotte Magister, Jonathan Mallinson, Jakub Adamek, Eric Malmi, and Aliaksei Severyn. Teaching small language models to reason. *arXiv preprint arXiv:2212.08410*, 2022.
- Qwen Team. Qwen2.5 technical report. *arXiv preprint arXiv:2412.15115*, 2024.
- Victor Sanh, Lysandre Debut, Julien Chaumond, and Thomas Wolf. DistilBERT, a distilled version of BERT: smaller, faster, cheaper and lighter. In *NeurIPS EMC2 Workshop*, 2019.
- R. Keith Sawyer. *Explaining Creativity: The Science of Human Innovation*. Oxford University Press, 2nd edition, 2011.
- Chenglei Si, Diyi Yang, and Tatsunori Hashimoto. Can large language models generate novel research ideas? A large-scale human study with 100+ NLP researchers. *arXiv preprint arXiv:2409.04109*, 2024.
- Claire Stevenson, Iris Smal, Matthijs Baas, Raoul Grasman, and Han van der Maas. Putting GPT-3’s creativity to the (alternative uses) test. In *Proceedings of ICCV*, 2022.
- Thomas B. Ward. Structured imagination: The role of category structure in exemplar generation. *Cognitive Psychology*, 27(1):1–40, 1994.
- Kevin Yang, Dan Klein, Nanyun Peng, and Yuandong Tian. Re3: Generating longer stories with recursive reprompting and revision. In *Proceedings of EMNLP*, 2022.