

A Conceptual Framework for Mind-Expanding Agents: Integrating Adversarial Questioning with AI-Guided Critical Thinking Development

Jeena Rajesh¹, Rejina P V², Kochumol Abraham^{3*}, Win Mathew John⁴, Terry Jacob Mathew⁵, Nellimala Abdul Shukoor⁶

¹Department of Computer Science, School of Computational & Physical Science, Kristu Jayanti University, Bengaluru, K. Narayanapura, Kothanur P.O., Karnataka, India -560077

Email ID: jeena.r@kristujayanti.com, ORCID: 0009-0008-5286-8641

²CO-OPERATIVE ARTS AND SCIENCE COLLEGE, MADAYI, PAZHAYANGADI, KANNUR

Email: sreereji.pv@gmail.com, ORCID: 0009-0000-6935-5246

^{3*}PG Department of Computer Applications, Marian College Kuttikanam Autonomous

Email: kochumol.abraham@mariancollege.org, ORCID: 0000-0002-7422-4125

⁴PG Department of Computer Applications, Marian College Kuttikanam Autonomous

Email: win.mathew@mariancollege.org

⁵Department of Computing and Engineering, University of West London, UAE

Email: terryjacobin@gmail.com, ORCID: 0000-0002-7966-2511

⁶Core Facility Manager, MBZUAI, Abudhabi

Email: abdul.shukoor@mbzuai.ac.ae

Abstract: Traditional AI educational systems prioritise knowledge delivery over critical thinking development, reinforcing convergent reasoning patterns that limit intellectual growth. This paper presents a novel conceptual framework integrating adversarial questioning principles with AI-guided critical thinking development to address fundamental limitations in current educational technologies. Our model encourages recasting AI as a cognitive adversary, rather than a traditional tutoring system that gives straight answers, by methodically crafting hostile questions that produce beneficial epistemic dissonance. To lay the theoretical groundwork for turning AI into an active collaborator in cognitive development rather than a passive information supplier, the framework combines adversarial machine learning ideas with classical Socratic pedagogical approaches. To address Bloom's 2-sigma dilemma, this conceptual contribution proposes scalable methods for teaching students critical thinking on an individual level. Optimal cognitive dissonance, safety-constrained questioning techniques, and dynamic user modelling are all elements of the framework, which aims to maintain users' mental well-being while fostering their intellectual growth. Formalising epistemic challenge calibration, developing adversarial pedagogical principles, and establishing ethical frameworks for cognitive confrontation in educational settings are all significant theoretical innovations.

Keywords: Adversarial Learning, Critical Thinking, Educational AI, Socratic Method, Cognitive Development, Conceptual Framework

1. INTRODUCTION

The proliferation of artificial intelligence in educational contexts has mainly focused on enhancing knowledge delivery and task efficiency [1]. However, this approach fails to promote critical thinking skills—vital for modern learning—while unintentionally encouraging convergent thinking processes [2]. Despite effect sizes comparable to human teaching ($d = 0.76$), most AI tutoring systems still primarily function as complex answer-delivery tools, which can hinder the development of autonomous critical thinking skills [3].



Given that Bloom's 2 sigma dilemma showed that one-on-one tutoring can outperform traditional classroom learning by two standard deviations, this restriction takes on added importance [4]. Despite the fact that this discovery has sparked decades of innovation in educational technology, the majority of these solutions still prioritize the transfer of knowledge over the development of cognitive abilities.

This paper's conceptual framework addresses these shortcomings by proposing confrontational questioning as a strategy for fostering critical thinking skills in the classroom. We suggest a radical change where AI systems actively question students' ideas, expose logical contradictions, and assist them in navigating constructive cognitive dissonance, drawing inspiration from adversarial machine learning [5] and the traditional Socratic approach [6].

1.1 Theoretical Motivation

Critical thinking development requires more than access to information—it necessitates systematic challenge and refinement of reasoning processes [2]. The dispositional essential aspect of thinking, including intellectual curiosity, open-mindedness, and systematic inquiry, can only be cultivated through practice with cognitively challenging scenarios [7].

Current educational AI systems suffer from three fundamental theoretical limitations:

1. **Epistemic Passivity:** Systems provide information rather than challenging reasoning processes
2. **Convergent Reinforcement:** Emphasis on correct answers limits exploration of reasoning quality
3. **Metacognitive Neglect:** Insufficient attention to thinking about thinking processes

1.2 Framework Objectives

This conceptual framework establishes theoretical foundations for:

1. **Adversarial Pedagogical Principles:** Formal integration of adversarial concepts with educational theory
2. **Cognitive Challenge Optimization:** Mathematical frameworks for calibrating epistemic dissonance
3. **Safety-Constrained Learning:** Ethical frameworks for psychological well-being during cognitive challenge
4. **Scalable Implementation:** Architectural designs for practical deployment

2. THEORETICAL FOUNDATIONS

2.1 Critical Thinking Theory

Critical thinking encompasses both cognitive skills and intellectual dispositions [2]. The mental dimension includes analysis, evaluation, inference, interpretation, explanation, and self-regulation. The dispositional dimension encompasses intellectual traits such as inquisitiveness, open-mindedness, systematic inquiry, analyticity, truth-seeking, cognitive maturity, and self-confidence in reasoning [7].

Research establishing expert consensus on critical thinking identifies key components that must be developed through practice rather than instruction alone [2]. These include the ability to identify assumptions, evaluate evidence quality, recognise bias, assess argument strength, and engage in metacognitive reflection about reasoning processes.

2.2 Adversarial Learning Principles

Adversarial machine learning, initially developed for generative modelling, uses competing objectives to enhance model performance [5]. The core principle involves two systems: a generator aiming to produce outputs that deceive a discriminator, and a discriminator trying to distinguish generated samples from real ones. These competitive dynamics foster improvement in both systems.

Translating this concept to education involves repositioning the traditional helper-learner relationship into a challenger-reasoner dynamic. Instead of providing support, the AI system generates questions designed to expose weaknesses in reasoning, creating productive cognitive tension that promotes deeper thinking.

2.3 Socratic Methodology

The Socratic method employs systematic questioning to expose contradictions, clarify concepts, and guide discovery learning [8]. Key principles include:

- **Elenchetic refutation:** Exposing contradictions in beliefs through questioning
- **Maieutic method:** Drawing knowledge out through guided inquiry
- **Aporia induction:** Creating productive confusion that motivates deeper inquiry
- **Dialectical progression:** Advancing understanding through question-answer sequences

These principles provide pedagogical foundations for adversarial questioning that promote rather than defeat learning.

2.4 Cognitive Dissonance Theory

Cognitive dissonance theory explains how tension between conflicting beliefs drives attitude and behaviour change [9]. In educational settings, optimal levels of cognitive dissonance encourage the adaptation of new information and the restructuring of knowledge frameworks. However, excessive dissonance can provoke defensive reactions that hinder learning.

The framework incorporates dissonance theory to calibrate challenge levels that promote growth while avoiding psychological harm.

3. CONCEPTUAL FRAMEWORK ARCHITECTURE

3.1 Formal Framework Definition

Definition 1: An Adversarial Educational Agent (AEA) is a quintuple $A = (Q, U, S, C, R)$ where:

- **Q:** Question Generation Function mapping user responses to adversarial questions
- **U:** User Modeling System maintaining cognitive and emotional profiles
- **S:** Safety Monitoring Framework ensuring psychological well-being
- **C:** Challenge Calibration Engine optimizing cognitive dissonance levels
- **R:** Response Evaluation Mechanism assessing reasoning quality

Objective Function: The AEA optimizes critical thinking development while maintaining safety constraints:

maximize $\Sigma(\text{CT_improvement})$ subject to $\text{Safety_constraints} \geq \text{threshold}$

Where CT_improvement represents measurable critical thinking advancement and Safety_constraints ensure psychological well-being.

3.2 Question Generation Framework

Definition 2: The Question Generation Function Q implements adversarial pedagogical strategies:

$Q(\text{response}, \text{context}, \text{user_model}) \rightarrow \text{adversarial_question}$

The function employs six primary questioning strategies:

Table I: Adversarial Questioning Taxonomy

Strategy Type	Cognitive Target	Mathematical Formulation	Example Implementation
Assumption Challenge	Implicit beliefs	$A(r) = \text{extract_assumptions}(r) \cap \text{challenge_set}$	"What assumptions underlie your conclusion that X?"
Evidence Scrutiny	Justification quality	$E(r) = \text{evaluate_evidence}(r) \rightarrow \text{weakness_identification}$	"What evidence supports this claim? How strong is it?"
Alternative Generation	Perspective diversity	$\text{Alt}(r) = \text{generate_alternatives}(r, \text{knowledge_base})$	"What alternative explanations might account for this?"
Logical Consistency	Reasoning structure	$L(r) = \text{check_consistency}(\text{premises}, \text{conclusions})$	"How does this conclusion follow from your premises?"
Metacognitive Probing	Thinking processes	$M(r) = \text{reflect_on_reasoning}(r) \rightarrow \text{process_awareness}$	"What reasoning strategy did you use here?"
Contradiction Exposure	Internal inconsistencies	$C(r) = \text{identify_contradictions}(\text{belief_set})$	"How do you reconcile these conflicting positions?"

3.3 Dynamic User Modeling

Definition 3: The User Modeling System U maintains multi-dimensional profiles:

$U = \{\text{cognitive_level}, \text{bias_patterns}, \text{domain_expertise}, \text{emotional_state}, \text{learning_preferences}\}$

Table II: User Modeling Dimensions and Calibration Parameters

Dimension	Measurement Approach	Calibration Formula	Safety Bounds
Cognitive Development	Response complexity analysis	$CD = \Sigma(\text{reasoning_depth}) / \text{response_count}$	Piagetian stage constraints
Bias Susceptibility	Pattern recognition in arguments	$BS = \text{detected_biases} / \text{reasoning_instances}$	Avoid reinforcement of harmful biases
Domain Expertise	Knowledge depth assessment	$DE = \text{correct_concepts} / \text{total_concepts}$	Prevent cognitive overload
Emotional State	Sentiment and stress analysis	$ES = \text{weighted_affect_scores}$	Stress threshold = 0.3 (normalized)
Engagement Level	Response quality metrics	$EL = (\text{length} + \text{depth} + \text{creativity}) / 3$	Disengagement threshold = 0.4

3.4 Safety-Constrained Challenge Calibration

Definition 4: The Challenge Calibration Engine C optimizes cognitive dissonance levels:

$\text{optimal_challenge} = \text{argmax}(\text{learning_gain})$ subject to $\text{psychological_safety}$

Mathematical Formulation:

$C(\text{user_state}, \text{question_difficulty}) = \{$

```

if stress_level < safety_threshold:
    increase_challenge_by( $\alpha$ )
elif learning_stagnation:
    maintain_challenge_level
else:
    decrease_challenge_by( $\beta$ )
}

```

Where α and β are learning rate parameters calibrated to individual users.

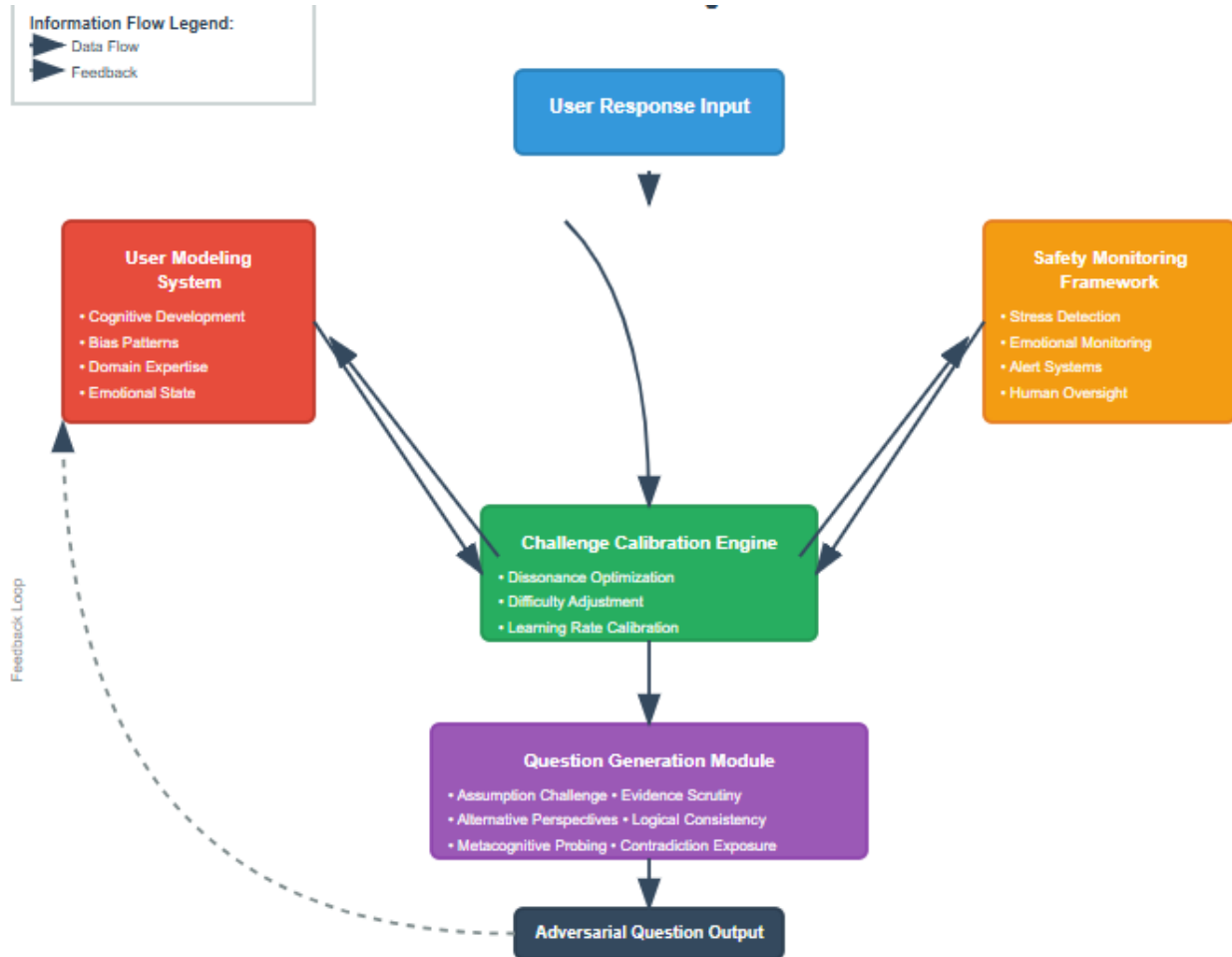


Fig. 1: Adversarial Educational Agent Architecture

The Conceptual diagram in Fig. 1 includes (1) User input processing, (2) multi-dimensional user modelling, (3) Safety monitoring with real-time alerts, (4) Challenge calibration engine with feedback loops, (5) Question generation with adversarial strategies, (6) Response evaluation and adaptation cycles. Arrows indicate bidirectional information flow and continuous adaptation mechanisms.

4. DESIGN PRINCIPLES AND IMPLEMENTATION GUIDELINES

4.1 Adversarial Pedagogical Principles

The framework establishes five core principles for adversarial educational interactions:

Principle 1 (Productive Opposition): Adversarial questions should challenge reasoning without attacking the reasoner, maintaining focus on cognitive processes rather than personal characteristics.

Principle 2 (Graduated Challenge): Question difficulty should adapt dynamically to user capability, maintaining optimal challenge zones that promote growth without overwhelming.

Principle 3 (Constructive Dissonance): Cognitive tension should motivate deeper inquiry rather than defensive responses, requiring careful calibration of challenge intensity.

Principle 4 (Metacognitive Emphasis): Questions should promote reflection on thinking processes, developing both reasoning skills and reasoning awareness.

Principle 5 (Ethical Boundaries): All challenges must respect human dignity and psychological well-being, with clear protocols for intervention when distress is detected.

4.2 Safety Framework Design

Definition 5: The Safety Monitoring Framework S implements multi-layered protection:

$S = \{\text{emotional_monitoring, stress_detection, intervention_protocols, human_oversight}\}$

Safety Architecture:

Table III: Multi-Layered Safety Protocol

Safety Layer	Monitoring Method	Intervention Threshold	Response Protocol
Real-time Sentiment	NLP emotion analysis	Negative affect > 0.7	Reduce challenge intensity
Physiological Stress	Behavioral pattern analysis	Response time variance > 2σ	Pause questioning sequence
Cognitive Load	Response quality degradation	Coherence score < 0.4	Simplify question complexity
User Agency	Explicit opt-out requests	User command "stop" or "help"	Immediate system pause
Human Oversight	Professional monitoring	Alert conditions met	Expert intervention available

4.3 Implementation Specifications

Technical Requirements:

- **Natural Language Processing:** Transformer architecture $\geq 7B$ parameters for robust question generation
- **Real-time Processing:** Response latency < 2 seconds for natural interaction flow
- **Multi-modal Analysis:** Integration of text, speech patterns, and behavioral indicators
- **Scalable Architecture:** Microservices design supporting concurrent users
- **Privacy Protection:** End-to-end encryption for all user data

Performance Metrics:

- **Question Relevance:** Automated scoring of question-context alignment
- **Challenge Calibration:** Maintenance of optimal difficulty zones ($\pm 0.5\sigma$ from user capability)

- **Safety Compliance:** Zero tolerance for psychological harm indicators
- **Learning Efficacy:** Measurable improvement in critical thinking assessments

5. APPLICATIONS AND USE CASES

5.1 Educational Domain Applications

Table IV: Framework Applications Across Educational Contexts

Educational Level	Specific Applications	Target Competencies	Implementation Considerations
Higher Education	Philosophy courses, Research methods	Abstract reasoning, Argument analysis	Integration with existing LMS platforms
Professional Training	Medical diagnosis, Legal reasoning	Domain-specific critical analysis	Regulatory compliance for professional standards
K-12 Education	Science inquiry, Literature analysis	Evidence evaluation, Interpretation skills	Age-appropriate questioning strategies
Corporate Learning	Strategic planning, Risk assessment	Decision-making under uncertainty	Workplace-relevant scenario development
Civic Education	Media literacy, Policy analysis	Source credibility, Bias recognition	Non-partisan approach requirements

5.2 Theoretical Research Applications

The framework provides foundations for investigating fundamental questions in educational psychology and artificial intelligence:

1. **Optimal Challenge Theory:** Empirical testing of mathematical models for cognitive dissonance calibration
2. **Adversarial Pedagogy:** Investigation of competitive dynamics in educational contexts
3. **Automated Socratic Dialogue:** Computational implementation of classical philosophical methods
4. **Safety-Constrained Learning:** Development of ethical frameworks for AI-human cognitive interaction

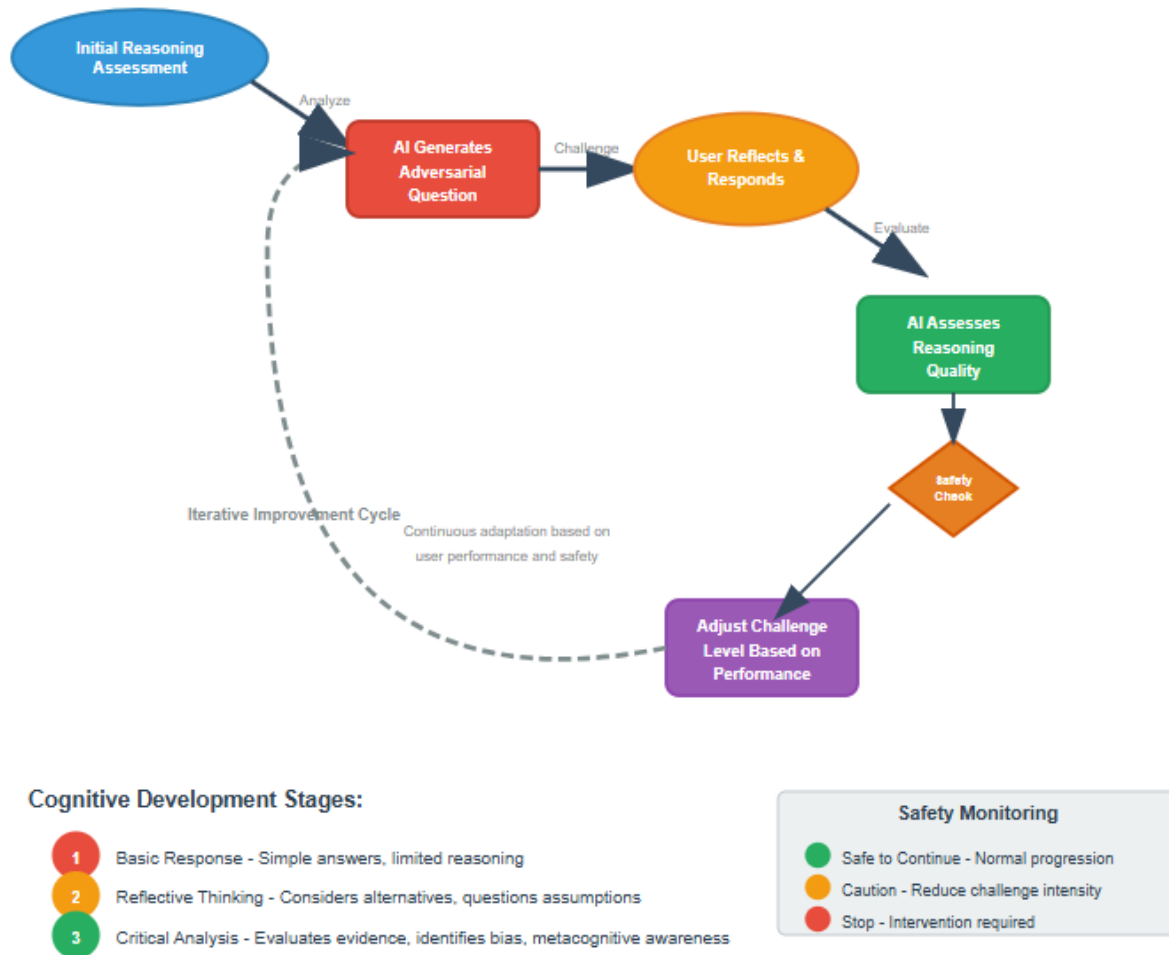


Fig. 2: Cognitive Development Progression Model

This Conceptual flowchart in Fig. 2 illustrates (1) Initial reasoning assessment, (2) Adversarial question generation, (3) User cognitive response, (4) Reasoning quality evaluation, (5) Challenge level adjustment, and (6) Iterative improvement cycles. Include feedback loops, safety checkpoints, and progression indicators showing advancement through critical thinking developmental stages.

6. EMPIRICAL VALIDATION FRAMEWORK

6.1 Experimental Design Template

Research Hypotheses:

H1: Students exposed to adversarial questioning will demonstrate greater improvement in critical thinking skills compared to traditional AI tutoring (effect size $d > 0.8$).

H2: The dynamic challenge calibration will maintain optimal cognitive load while avoiding psychological distress (safety compliance rate $> 95\%$).

H3: Domain transfer effects will be observed, with critical thinking improvements generalizing across subject areas (transfer coefficient > 0.6).

Methodology Template:

- **Design:** Randomized controlled trial with three conditions (adversarial, traditional, control)

- **Participants:** Stratified sampling across educational levels and demographic variables
 - **Duration:** Minimum 8-week intervention period for developmental effects
 - **Measures:** Pre-post critical thinking assessments, real-time engagement metrics, safety indicators
- Success Criteria:**

Table V: Validation Metrics and Benchmarks

Outcome Category	Primary Measures	Success Threshold	Assessment Methods
Critical Thinking	California Critical Thinking Skills Test	Effect size $d \geq 0.8$	Pre-post standardized assessment
Reasoning Quality	Argument analysis rubrics	30% improvement in reasoning scores	Expert-rated reasoning samples
Domain Transfer	Cross-subject reasoning tasks	Transfer coefficient ≥ 0.6	Novel problem-solving assessments
Safety Compliance	Stress and engagement indicators	Zero harm incidents, 95% positive experience	Real-time monitoring and surveys
User Experience	Engagement and satisfaction	80% positive user ratings	Mixed-methods evaluation

6.2 Longitudinal Research Design

Phase 1 (Months 1-6): Proof of concept with limited user testing. **Phase 2 (Months 7-18):** Controlled efficacy studies with educational partners. **Phase 3 (Months 19-36):** Large-scale deployment and effectiveness research. **Phase 4 (Years 4-5):** Long-term impact assessment and system refinement

7. TECHNICAL IMPLEMENTATION ROADMAP

7.1 Development Architecture

Component Development Priority:

- Core Question Generation Engine** (Months 1-4)
 - Natural language understanding for response analysis
 - Adversarial questioning strategy implementation
 - Initial question quality assessment systems
- User Modeling System** (Months 3-6)
 - Cognitive profile development and maintenance
 - Real-time state tracking and adaptation
 - Privacy-preserving data management
- Safety Monitoring Framework** (Months 5-8)
 - Multi-modal stress detection systems
 - Intervention protocol implementation
 - Human oversight integration
- Challenge Calibration Engine** (Months 7-10)

- Mathematical optimization of difficulty levels
- Machine learning for personalization
- Continuous adaptation algorithms

7.2 Technical Challenges and Solutions

Challenge 1: Natural Language Understanding for Reasoning Assessment

Solution: Fine-tune large language models on argumentation datasets with expert annotations

Challenge 2: Real-time Emotional State Detection

Solution: Multi-modal fusion of textual sentiment, response patterns, and behavioral indicators

Challenge 3: Scalable Personalization

Solution: Federated learning approaches that maintain user privacy while enabling model improvement

Challenge 4: Safety Validation

Solution: Extensive testing with educational psychologists and ethics review boards

8. ETHICAL CONSIDERATIONS AND RISK MANAGEMENT

8.1 Ethical Framework

The framework operates under strict ethical guidelines addressing:

Informed Consent: Users must understand that the system will challenge their thinking and potentially create cognitive discomfort.

Psychological Safety: Comprehensive monitoring and intervention protocols prevent harm from excessive cognitive challenge.

Epistemic Justice: Questions and challenges must respect diverse ways of knowing and cultural perspectives on reasoning.

Educational Equity: The system must provide equal opportunities for cognitive development regardless of background or prior knowledge.

Data Privacy: All user interactions and cognitive profiles must be protected with highest security standards.

8.2 Risk Mitigation Strategies

Table VI: Risk Assessment and Mitigation

Risk Category	Specific Risks	Probability	Impact	Mitigation Strategy
Psychological Harm	Excessive stress, anxiety, defensive responses	Medium	High	Multi-layered safety monitoring with immediate intervention
Epistemic Manipulation	Unfair intellectual coercion or bias	Low	High	Transparent algorithms, diverse question development team

Privacy Violation	Unauthorized access to cognitive profiles	Low	High	End-to-end encryption, minimal data collection
Cultural Insensitivity	Questions that reflect cultural bias	Medium	Medium	Cross-cultural validation, diverse development team
Academic Dishonesty	Students using system inappropriately for assessments	High	Medium	Clear usage guidelines, instructor training

9. FUTURE RESEARCH DIRECTIONS

9.1 Theoretical Extensions

Advanced Cognitive Modeling: Development of more sophisticated models of reasoning development incorporating insights from cognitive science and developmental psychology.

Multimodal Adversarial Questioning: Extension to include visual, auditory, and interactive elements that challenge reasoning across multiple modalities.

Collaborative Adversarial Learning: Investigation of group dynamics when multiple learners engage with adversarial questioning systems simultaneously.

Cross-Cultural Validation: Systematic investigation of how adversarial questioning strategies vary across cultural contexts and reasoning traditions.

9.2 Technical Innovations

Quantum-Enhanced Reasoning: Investigation of quantum computing approaches to complex reasoning assessment and question generation.

Neuroadaptive Interfaces: Integration with brain-computer interfaces for direct monitoring of cognitive load and emotional state.

Augmented Reality Integration: Development of immersive environments where adversarial questioning occurs within realistic problem-solving contexts.

Explainable AI: Advanced interpretability methods to help users understand how and why specific questions are generated.

9.3 Empirical Research Programs

Long-term Developmental Studies: Multi-year investigations of how adversarial questioning affects reasoning development from childhood through adulthood.

Professional Application Research: Systematic investigation of effectiveness in professional training contexts, including medicine, law, engineering, and business.

Comparative Effectiveness Research: Head-to-head comparisons with other critical thinking interventions to establish relative effectiveness.

Implementation Science: Research on factors affecting successful adoption and implementation in educational institutions.

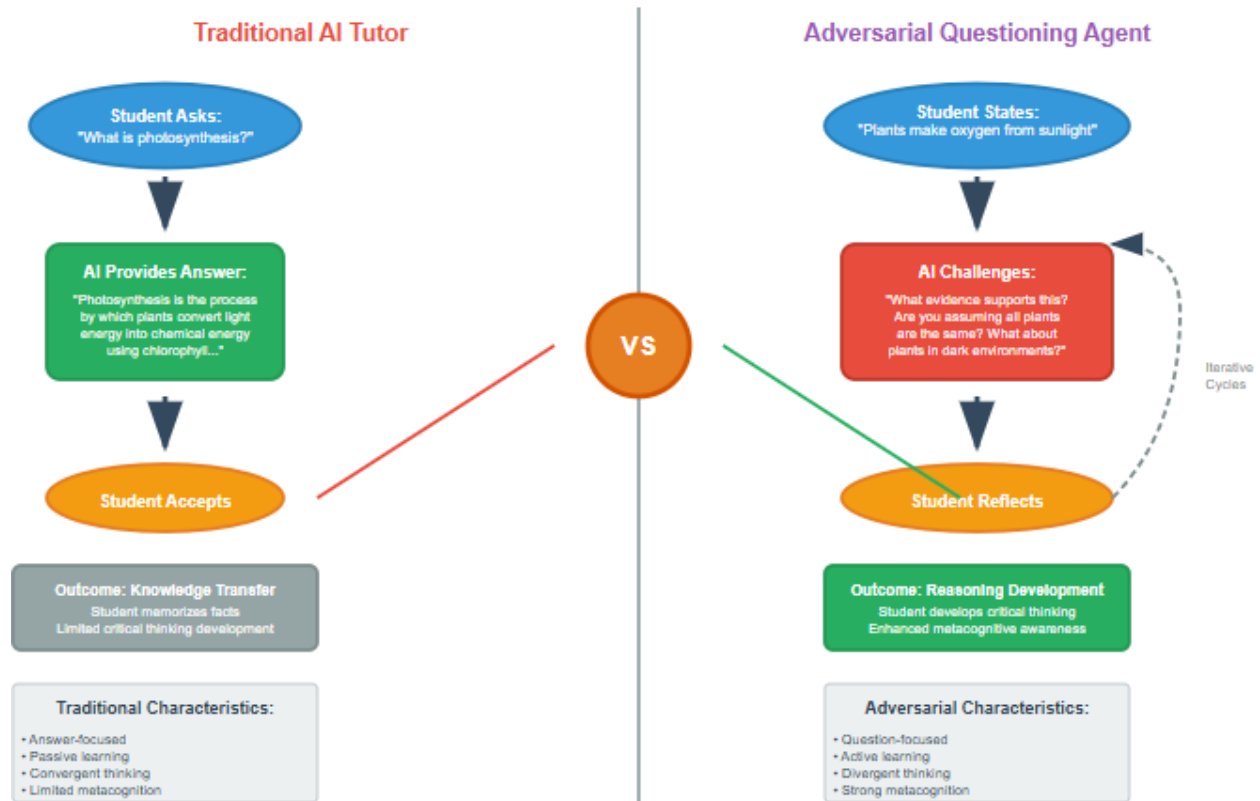


Fig. 3: Traditional vs. Adversarial AI Tutoring Approaches

This comparative diagram illustrates fundamental differences between conventional AI tutoring and the proposed adversarial questioning framework. The left side illustrates traditional interaction patterns: students ask questions, the AI provides direct answers, and students passively accept the information, resulting in knowledge transfer with limited critical thinking development. The right side demonstrates adversarial questioning: students make statements, AI challenges with probing questions, students reflect and engage in iterative cycles of deeper reasoning, leading to enhanced critical thinking and metacognitive awareness. The central "VS" indicator emphasises the paradigmatic shift from answer-focused to question-focused pedagogy, highlighting key characteristics of each approach, including learning modality, thinking patterns, and cognitive outcomes.

10. CONCLUSION

This conceptual framework signals a shift in educational AI from passive information dissemination to active cognitive engagement. By combining adversarial machine learning principles with traditional Socratic methods, the framework tackles core limitations in current educational technologies while laying the groundwork for a new generation of AI-supported learning systems.

The framework's key theoretical contributions include:

1. **Formal Mathematical Models** for cognitive challenge calibration and safety-constrained learning optimization
2. **Adversarial Pedagogical Principles** that systematically apply competitive dynamics to educational contexts
3. **Comprehensive Safety Frameworks** ensuring psychological well-being during cognitive challenge
4. **Scalable Implementation Architecture** providing practical pathways for real-world deployment

The approach directly addresses Bloom's 2 sigma problem by proposing scalable methods for individualized critical thinking instruction that could approach the effectiveness of one-on-one human tutoring. Unlike traditional approaches that focus on knowledge transmission, this framework targets the development of reasoning processes and metacognitive awareness, which are essential for independent critical thinking.

Significance for Educational AI: The framework establishes new research directions that prioritize cognitive development over information delivery, potentially revolutionizing how AI systems support human learning and reasoning.

Implications for Critical Thinking Education: By formalising adversarial questioning principles, the framework provides systematic methods for developing reasoning skills that transfer across domains and contexts.

Contributions to Learning Science: The integration of adversarial concepts with educational theory opens new avenues for understanding how productive intellectual conflict promotes cognitive development.

Future work should concentrate on empirically validating core theoretical predictions, implementing key components technically, and exploring long-term developmental impacts. The framework offers a solid foundation for evolving educational AI from merely supportive tools to cognitive development partners that actively nurture the critical thinking skills necessary for success in an increasingly complex world.

The ultimate vision is AI systems that challenge rather than serve human cognition, fostering the intellectual autonomy and critical reasoning capabilities necessary for democratic citizenship, professional excellence, and lifelong learning in the 21st century.

11. Acknowledgments

The authors acknowledge the foundational theoretical work of Benjamin Bloom on mastery learning and individual differences, Robert Ennis on critical thinking dispositions, Peter Facione on critical thinking assessment, and Richard Paul and Linda Elder on Socratic questioning principles. This conceptual framework builds upon decades of research in critical thinking development, educational psychology, and artificial intelligence.

Author Contributions: This conceptual framework represents collaborative theoretical development that integrates insights from educational psychology, artificial intelligence, and cognitive science.

Ethics Statement: This paper presents theoretical frameworks only. Any future empirical implementation would require a comprehensive ethics review and approval from relevant institutional review boards.

Data Availability: No empirical data were collected for this conceptual framework. All mathematical formulations and design specifications are provided within the manuscript.

Conflicts of Interest: The authors declare that they have no competing interests in the development or implementation of the proposed framework.

REFERENCES

1. O. Zawacki-Richter, V. I. Marín, M. Bond, and F. Gouverneur, "Systematic review of research on artificial intelligence applications in higher education—where are the educators?" *International Journal of Educational Technology in Higher Education*, vol. 16, no. 1, pp. 1-27, 2019.
2. P. A. Facione, "Critical thinking: A statement of expert consensus for purposes of educational assessment and instruction," American Philosophical Association, 1990.
3. K. VanLehn, "The relative effectiveness of human tutoring, intelligent tutoring systems, and other tutoring systems," *Educational Psychologist*, vol. 46, no. 4, pp. 197-221, 2011.
4. B. S. Bloom, "The 2 sigma problem: The search for methods of group instruction as effective as one-to-one tutoring," *Educational Researcher*, vol. 13, no. 6, pp. 4-16, 1984.
5. I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative adversarial nets," in *Advances in Neural Information Processing Systems*, vol. 27, pp. 2672-2680, 2014.
6. R. Paul and L. Elder, "The miniature guide to critical thinking concepts and tools," Foundation for Critical Thinking, 2007.

7. R. H. Ennis, "Critical thinking dispositions: Their nature and assessability," *Informal Logic*, vol. 18, no. 2, pp. 165-182, 1996.
8. M. Copeland, *Socratic circles: Fostering critical and creative thinking in middle and high school*, Stenhouse Publishers, 2005.
9. L. Festinger, *A theory of cognitive dissonance*, Stanford University Press, 1957.