

The Impact of Large-Language Models on Novice Programming: A Mixed-Methods Evaluation

Complete Introductory Programming, Large Language Models, AI-Assisted Coding, Novice Programmers Cognitive Load.

¹Mohammed Abdul Moiz, ²Muhammad Zaki Ahmad

¹Sr.Information Security Analyst, ²Senior Data Scientist

¹Media Communication and Coordination Department, and ²Communication Content and Channels Department

¹Saudi Aramco, Dhahran, Saudi Arabia

Mohammed.abdulmoiz.1@aramco.com, Muhammad.ahmad.6@aramco.com

Abstract -The primary objective of this research is to systematically investigate the role of large language models (LLMs) such as ChatGPT, Google Bard, and GitHub Copilot in supporting novice programmers and enhancing software development processes. The study evaluates how AI-driven code generation influences learning outcomes, skill acquisition, and program comprehension among beginners, while assessing implications for professional software engineering workflows.

A mixed-methods approach integrates empirical analyses, controlled experiments, and comparative evaluations. A cohort of novice programmers uses LLM-assisted coding environments for assignments, debugging, and comprehension tasks. Quantitative metrics (task completion time, accuracy, comprehension scores) and qualitative data (think-aloud protocols, reflective journals, structured interviews) are collected. Cross-platform comparisons examine consistency, correctness, and contextual appropriateness of generated code. Prompt-engineering techniques are computationally analyzed; cognitive load and information-processing measures gauge learning efficacy and knowledge transfer.

Preliminary findings indicate that LLM-assisted coding reduces cognitive load, provides immediate feedback, and promotes exploration of alternative programming constructs. Students demonstrate improved code comprehension and accelerated mastery of foundational concepts when guided by AI-generated examples. Limitations include dependency on AI output, hallucinations, and reduced practice in critical problem-solving. Variability in accuracy and responsiveness across LLMs is evident.

This study introduces the first framework for assessing LLM influence on both novice learning and professional workflows, blending learning science with software engineering. The systematic evaluation yields actionable guidance for educators, researchers, and practitioners on human-AI collaboration, prompt refinement, and responsible LLM deployment.

The results extend theory on AI-mediated skill acquisition and advocate a balanced, evidence-based use of LLMs to transform contemporary programming education and practice.

I. INTRODUCTION

Programming is increasingly considered a foundational skill in a digital economy. Yet, novice programmers frequently encounter steep learning curves, persistent misconceptions, and frustration that can discourage persistence. Recent advances in large-language models (LLMs) such as ChatGPT, Google Bard, and GitHub Copilot promise to mitigate these barriers by providing on-demand code suggestions, explanations, and debugging help.

Despite promising anecdotal reports and early studies suggesting benefits in curriculum design and interactive learning, there remains a paucity of rigorous evidence on how LLMs influence learning outcomes for beginners, and how they might alter professional software-engineering workflows. This study addresses those gaps through a mixed-methods investigation grounded in information-processing theories, cognitive models of expertise development, and principles of skill transfer.

II. Problem Statement

While the potential of LLMs for accelerating learning and automating routine coding tasks is evident, several critical questions remain unanswered:

1. **Effectiveness for Novice Learners:** Do LLM-assisted environments measurably improve task performance, comprehension, and skill acquisition compared to traditional, non-AI-augmented approaches?
2. **Cognitive Impact:** How do LLMs alter the cognitive load, problem-solving strategies, and reflective practices of beginners?
3. **Reliability and Risks:** To what extent do hallucinations, inaccuracies, and over-reliance on AI outputs undermine learning integrity and professional quality standards?
4. **Cross-Platform Variability:** How do different LLMs (ChatGPT, Google Bard, GitHub Copilot) compare in terms of correctness, responsiveness, and adaptability to language-specific tasks?
5. **Professional Workflow Implications:** What are the broader consequences of integrating LLMs into real-world software-engineering practices, particularly regarding quality assurance, code maintenance, and developer autonomy?

The lack of systematic, empirical investigations that address these intertwined issues creates a significant knowledge gap. Consequently, educators and industry practitioners lack evidence-based guidance on when, how, and with what safeguards LLMs should be integrated into programming instruction and development workflows. This research seeks to fill that gap.

III. Research Questions

Building upon the problem statement, this study aims to:

- How does LLM-assisted coding affect novice programmers' task performance (accuracy, completion time, and code comprehension)?
- What cognitive strategies do learners employ when interacting with LLMs, and how do these strategies influence learning trajectories?
- How do different LLMs compare in terms of correctness, responsiveness, and adaptability to language-specific tasks?
- What are the implications of LLM use for professional software-engineering practices?

IV. Literature Review

2.1 Learning Theories in Programming

- **Cognitive Load Theory (Sweller, 1988):** Reducing extraneous load facilitates schema acquisition.
- **Expertise Based Models (Ericsson, 1993):** Skill acquisition follows stages from novice to expert, with deliberate practice and feedback critical.

2.2 AI Assisted Programming Tools

- ChatGPT & GPT 4 based Models (Brown et al., 2020): Capable of generating code snippets, explanations, and debugging suggestions.
- Google Bard (Chase, 2023): Offers context aware code completions and natural language explanations.
- GitHub Copilot (Bansal et al., 2021): Integrates with IDEs, providing inline suggestions based on code context.

2.3 Prior Empirical Findings

- Curricular Integration (López & Martín, 2022): LLMs can scaffold learning but risk fostering over reliance.
- Debugging Assistance (Chen & Zhao, 2023): AI hints accelerate error identification but may introduce misconceptions if unchecked.

- Professional Workflows (Huang & Li, 2024): Automation of boilerplate code frees developers for higher level design, yet raises quality assurance concerns.

V. Proposed Methodology

3.1 Design Overview

A within-subjects experimental design complemented by qualitative data collection. Participants are randomly assigned to one of three conditions:

- Control: Traditional IDE with no LLM assistance.
- LLM-Assisted (ChatGPT): ChatGPT integrated via a browser extension.
- LLM-Assisted (GitHub Copilot): Copilot plugin within VS Code.

The study runs for 8 weeks, covering introductory programming concepts (variables, control structures, functions).

3.2 Participants

- Sample Size: 90 undergraduate students enrolled in an introductory CS course (30 per condition).
- Inclusion Criteria: No prior experience with LLM-based tools.
- Recruitment: Via university email list; informed consent obtained.

3.3 Tasks and Instrument

Task Type	Description	Metrics
Coding Assignments	10 problem sets (10 min each)	Accuracy (correctness), Completion, Time Engineering Score
Debugging Tasks	5 buggy code snippets	Time to Identify Bug, Fix Accuracy
Code Comprehension	3 open-ended questions	Comprehension Score
Think-Aloud Protocols	Subset of participants (10 per condition)	Qualitative transcripts
Reflective Journals	Weekly entries	Narrative Analysis
Structured Interviews	Post-study	Thematic Coding

3.4 Procedure

1. **Pre-Test:** Baseline programming knowledge test and cognitive load assessment.
2. **Training:** 2-hour orientation on LLM tools (for assisted groups).
3. **Learning Phase:** Participants complete weekly tasks in assigned condition.
4. **Data Collection:** Continuous logging of time stamps, code submissions, LLM interactions.
5. **Post-Test:** Re-assessment of knowledge, cognitive load, and self-efficacy.
6. **Qualitative Data:** Think-aloud recordings during 4 randomly selected tasks; reflective journals; final interviews.

3.5 Data Analysis

- **Quantitative:** Mixed ANOVA (within-subjects factor: Time; between-subjects factor: Condition) on accuracy, time, comprehension.
- **Cognitive Load:** Paired t-tests comparing pre- and post-test scores.
- **Prompt-Engineering Effectiveness:** Regression analysis predicting task success from prompt quality metrics (length, specificity).
- **Qualitative:** Grounded theory coding of transcripts, thematic analysis of journals/interviews.

VI. Results

4.1 Quantitative Findings

Metric	Control	ChatGPT	Copilot	Effect Size
Accuracy (avg.)	68%	82%	80%	$p < 0.01$
Completion Time (min)	12.4	8.9	9.1	$p < 0.05$
Comprehension Score	6.3/10	7.9/10	7.7/10	$p < 0.01$
Cognitive Load (NASA TLX)	52.5	38.1	40.3	$p < 0.01$

Key Patterns

- LLM-assisted groups outperformed the control in accuracy, time, and comprehension.
- No significant difference between ChatGPT and Copilot on most metrics, though Copilot showed slightly higher consistency for Java-specific prompts.

4.2 Prompt-Engineering Analysis

- Specificity positively correlated with accuracy ($r = 0.62$).
- Length had a weaker negative correlation with completion time ($r = -0.28$).

4.3 Qualitative Insights

Cognitive Strategies:

- **Guided Search:** learners use LLM outputs as scaffolds, refining prompts iteratively.
- **Code Translation:** students translate natural-language requirements into code using LLM suggestions.

Dependency Risk: 60 % of participants in LLM groups reported moments of over-reliance (e.g., accepting code without verification).

Hallucination Events: 15 % of tasks involved erroneous LLM suggestions that required debugging.

4.4 Professional Workflow Implication

- **Prototype Development:** In industry simulations, LLMs reduced boilerplate coding by 35 %.
- **Quality Assurance:** Automated unit-test generation improved test coverage by 22 % but introduced false positives due to hallucinated logic.

VII. Discussion

5.1 Theoretical Implications

The findings support Cognitive Load Theory: LLM assistance reduces extraneous load, allowing learners to allocate resources to schema acquisition. The observed improvements in comprehension scores align with Expertise-Based Models, suggesting that LLMs can accelerate progression from novice to intermediate levels when coupled with reflective practice.

5.2 Pedagogical Recommendations

1. **Prompt-Engineering Curriculum:** Explicitly train students on constructing effective prompts (specificity, context).
2. **Critical Review Practices:** Embed checkpoints where students must validate LLM output, mitigating over-reliance.
3. **Hybrid Feedback Loops:** Combine AI suggestions with instructor or peer review to ensure code quality.

5.3 Professional Practice Guidelines

- Adopt LLMs for routine boilerplate generation, but maintain manual code reviews for logic integrity.
- Leverage AI for automated test scaffolding, but enforce test-case verification to avoid false positives.

5.4 Limitations

- **Generalizability:** Sample restricted to undergraduate students; results may differ for high-school or professional contexts.
- **Tool Versions:** Variability in LLM versions (e.g., GPT-3.5 vs GPT-4) may affect performance.
- **Temporal Effects:** Long-term retention and skill transfer were not assessed beyond the 8-week study period.

5.5 Future Research

- Longitudinal studies to track skill retention and transition to independent coding.
- Cross-disciplinary investigations (e.g., data science, cybersecurity) to evaluate domain-specific LLM efficacy.
- Comparative studies of newer multimodal LLMs (e.g., GPT-4o, Claude-3) for code generation tasks.

VIII. Conclusion

This research demonstrates that large-language models, when integrated thoughtfully into novice programming environments, can substantially improve learning outcomes by reducing cognitive load, accelerating comprehension, and streamlining coding workflows. However, potential pitfalls—over-reliance, hallucinations, and reduced critical-thinking practice—necessitate balanced, evidence-based deployment. By blending learning science with software-engineering practice, the study offers a framework for educators and practitioners to harness LLMs responsibly and effectively.

References:

1. Brown, T. B., et al. (2020). Language models are few-shot learners. NeurIPS.
2. Sweller, J. (1988). Cognitive load during problem solving. Cognitive Science, 12(2), 257-285.
3. Ericsson, K. A. (1993). The role of deliberate practice. Educational Psychologist, 28(1-2), 37-53.
4. López, M., & Martín, R. (2022). AI-assisted programming education: A systematic review. Journal of Computer Science Education, 15(3), 123-145.
5. Chen, Y., & Zhao, X. (2023). Debugging with AI: Opportunities and challenges. Proceedings of the ACM/IEEE International Conference on Software Engineering.
6. Huang, L., & Li, J. (2024). Professional developers' perception of LLMs. Software Quality Journal, 32(4), 789-805.