

Describing Data Structures through Culture: Developing Resources for Computing Students with Large Language Models

Kevin Lemus, University of Maryland Baltimore County

Kevin Lemus is a graduate assistant researcher and Master's student at the University of Maryland, Baltimore County (UMBC). Kevin earned his Bachelor of Science in Information Systems along with a Certificate in User Experience, Web, and Mobile Development from UMBC in 2024. He is currently advancing his expertise as a candidate for a Master of Science in Human-Centered Computing, with an expected completion date of December 2026.

Stephanie Jill Lunn, Florida International University

Stephanie Lunn is an Assistant Professor in the Multidisciplinary Engineering and Computing Education, Systems and Management (ESM) department and the STEM Transformation Institute at Florida International University (FIU). She also has a secondary appointment in the Knight Foundation School of Computing and Information Sciences (KFSCIS). Previously, Dr. Lunn earned her doctorate in computer science from the KFSCIS at FIU, with a focus on computing education. She also holds B.S. and M.S. degrees in computer science from FIU and B.S. and M.S. degrees in neuroscience from the University of Miami. In addition, she served as a postdoctoral fellow in the Wallace H. Coulter Department of Biomedical Engineering at the Georgia Institute of Technology, with a focus on engineering education. Her research interests span the fields of computing and engineering education, human-computer interaction, data science, and machine learning.

Edward Dillon, University of Maryland Baltimore County

Edward Dillon is an Associate Professor in the Department of Information System at the University of Maryland, Baltimore County (UMBC). Edward received his B.A. in Computer and Informational Science from the University of Mississippi in 2007. He would go on to obtain his Masters and Ph.D. in Computer Science from the University of Alabama in 2009 and 2012, respectively. Prior to his arrival to UMBC in August 2024, Dr. Dillon served as a Computer Science Instructor at Jackson State University (2012-2013), and a Postdoctoral Researcher at Clemson University (2013-2014) and the University of Florida (2014-2016), and Assistant/Associate Professor in the Department of Computer Science at Morgan State University (2017-2024). His research focuses on human-centered computing with emphasis on computing education, broadening participation in computing, and tech workforce prep.

Dr. Krystal L Williams, University of Wisconsin-Madison | Education Policy & Equity Research Collective (Ed.PERC)

Krystal L. Williams, Ph.D. is an Associate Professor of Higher Education in the University of Wisconsin Department of Educational Leadership and Policy Analysis. She directs the Education Policy and Equity Research Collective (Ed.PERC), including the Computing Sciences Education Research Team (cSERT) where she and her team explore issues regarding race and public policy with an emphasis on Historically Black Colleges and Universities, and broadening participation in STEM. She attended the University of Michigan where she completed her doctoral studies in the Center for the Study of Higher and Postsecondary Education, and Clark Atlanta University where she earned a BS and MS in mathematics, graduating as valedictorian.

Christian Ruiz

Describing Data Structures through Culture: Developing Resources for Computing Students with Large Language Models

Abstract

Data structures are used to format, store, and access data and are a critical foundation of computing education and programming pedagogy. However, students unfamiliar with them may struggle to conceptualize these computational building blocks more abstractly. In recognition that the same descriptions may not resonate with all populations of students, we sought to explore how large language models (LLMs) can aid in contextual content creation of data structure explanations. Given that learners may come from a range of backgrounds, in this work, we focused on tailoring content to connect with a particular community, Hispanic or Latine students in the United States (U.S.). As a first step, we developed six prompts, each asking about six data structures, using them as input across six different LLMs. We applied discourse analysis to compare the ($n = 216$) outputs, utilizing established measures of culture to quantify patterns. The responses generated often centered on cuisine, followed by familial relationships, and then events. The majority of the outputs were in English, with some Spanish phrases applied, although the proportion of Spanish could increase in prompts specifying role play scenarios. Yet, given the broad framing, many of the examples followed suit. As such, the second phase applied the prompt identified as having the highest cultural relevance in the first phase to further scope to the population of a specific region, nationalities predominant in the southeastern region of the U.S. In particular, we explored possible explanations of data structures based on specifying 10 nationalities as part of the input across the six platforms. Additional analyses of the ($n = 360$) responses yielded further insight into more tailored examples to resonate with different populations with a high population density in this region. Again, culinary examples were the most prevalent; however, there was more specificity of dishes referenced. The outputs often provided step-by-step instructions, and these were frequently either entirely or mostly in English. In this paper, we share our approach, detailing the implications and potential ethical concerns. We also provide guidelines around how LLMs may aid in creating educational materials that could resonate with different groups.

1 Introduction

Data structures are considered a fundamental component of software covered in computing education curricula [1] and are a key feature of technical interviews for job attainment [2]. Yet, students tend to struggle with them, often resorting to memorization rather than “engaging core data structure concepts” [3, p. 176]. Educators have taken different approaches to aid in students’ understanding and application of these abstractions, including efforts around visualization with programs, flow charts [4], and augmented reality [5]. Generative Artificial Intelligence (AI) has

been integrated into homework as a way of improving learning outcomes and enhancing code quality [6]. In addition, scholars have described how teaching data structures based on culture by developing examples pertinent to students' lived experiences can help them to better connect with the content [7].

The process of incorporating aspects of culture into education is not new and has been shown to aid in students' engagement and understanding through the introduction of concepts in ways that may connect with their shared heritage [8]. Being able to develop content that can meet the needs of a range of learners can be challenging, especially in instances where instructors may not share the background of the students they serve. Ultimately, this work aimed to establish resources that instructors could potentially incorporate into their courses to communicate with different learners. Towards this goal, we elected to query Large Language Models (LLMs) as a source of generating examples that could aid in their conceptualization through connections with culture. While LLMs continue to evolve, computing education "students and instructors agree that LLMs should be welcome in academia" [9, p. 942] and have been described as potential tools for enhancing pedagogy. Existing literature confirms that LLMs are capable of accurate content creation, driving a strong interest among educators to use them for lesson preparation [10]. However, what is currently missing is an evaluation of how successfully these technologies can adapt that content to be culturally relevant for specific student populations.

Our approach is motivated by several challenges. First, the adoption of LLMs in education necessitates investigation of the quality of generated responses to ensure they meet the needs of diverse learners. In addition, students may lack sufficient understanding of data structures or may struggle to produce correct examples themselves, making it difficult to directly gather or validate preferred approaches. Faculty may also lack the cultural knowledge required to design examples that are broadly applicable across diverse student populations [11, 12]. Furthermore, although LLMs continue to evolve, it remains unclear how consistently culturally relevant outputs are produced across commercially available systems. Scholars note that LLMs may struggle with sociotechnical context, and suggested that there is a need to determine how well educational materials may align [13]. Finally, structured approaches to comparing tools and prompts are required when generating illustrative examples of data structure concepts.

Accordingly, we conducted a study to explore how large language models (LLMs) generate culturally relevant instructional content and to assess how they describe specific data structures. Our goal was to create examples that incorporate students' shared social identities and validate their experiences through a celebration of their unique backgrounds [14, 15]. We structured the research in two phases to compare model performance using generic prompts for individuals of Hispanic or Latine descent versus nationality-specific ones.

In this investigation, we sought to answer the following research questions (RQs):

- **RQ1:** How do the prompts applied in different LLMs impact examples provided when creating data structures for someone who is of Hispanic or Latine descent?
- **RQ2:** How do the outputs from different LLMs suggest describing varying data structures to someone who is of Hispanic or Latine descent?
- **RQ3:** How do the outputs from different LLMs suggest describing varying data structures

to someone who is of Hispanic or Latine descent from a specific nationality?

2 Motivation

The federal designation of Hispanic-Serving Institution (HSI) in the U.S. is defined by enrollment numbers, specifically institutions with at least 25% Hispanic full-time equivalent students. However, this label does not guarantee a curriculum tailored to students’ distinct cultural and linguistic backgrounds. We argue that the HSI designation often obscures a lack of substantive “servingness” in the classroom [14, 15], particularly in computer science, where pedagogical examples often rely on generic or majority-centric norms. This gap highlights a disconnect between institutional designation and instructional practice, motivating the need for more culturally responsive educational approaches.

Addressing this gap requires understanding the practical challenges involved in developing such instruction. This work is motivated by several interrelated considerations:

1. The increasing adoption of large language models (LLMs) in education necessitates systematic investigation of the quality of their outputs to ensure they meet the needs of diverse learners.
2. Students may lack sufficient understanding of data structures or may struggle to produce correct examples themselves, limiting the feasibility of directly eliciting or validating preferred approaches.
3. Faculty may not have sufficient knowledge of different cultural contexts to consistently design examples that are applicable across a broad range of students.
4. Although LLMs continue to evolve, it remains unclear how reliably they generate culturally relevant outputs across commercially available systems.
5. There is a need for structured approaches to comparing tools and prompts when developing illustrative cases for data structure concepts.

While HSIs enroll large and diverse Hispanic and Latine student populations, this population is not homogeneous, and students’ cultural identities vary significantly across nationality, region, and lived experience. As a result, curricular approaches that assume a uniform Hispanic identity risk overlooking important cultural distinctions.

Building on the goal of creating culturally responsive instructional resources that validate students’ lived experiences, including familial structures and regional traditions [14, 15], our work translates this motivation into practical computing education contexts.

3 Background

3.1 Theoretical Framework

Our work was guided by culturally responsive pedagogy (CRP), which centers on enhancing teaching through engagement of students’ cultural identities to draw on their knowledge, experiences, and perspectives [16, 17]. It has been described as “a pedagogy that empowers students intellectually, socially, emotionally, and politically by using cultural referents to impart knowledge, skills, and attitudes” [18, pp. 17-18]. Prior studies have shown CRP approaches can

be particularly valuable for supporting Hispanic or Latine students' sense of belonging and ethnic-racial identity [19].

In computing, CRP has been applied in the context of professional development as a way of promoting practices that could further enable support for more learners to engage and empower them in this lucrative field [20]. Related to this concept is Culturally Responsive Computing, which is a framework that extends culturally responsive pedagogy to computing education by emphasizing instructional practices that connect computing content to students' cultural identities and lived experiences and is suggested to support inclusion in the field, particularly for Black and Hispanic or Latine students' "desire for learning rooted in how they see the world from their viewpoint" [21, p. 15]. We applied CRP in our data collection through specific prompts developed in our analysis and in the interpretation of the outputs.

3.2 Drawing on Experience in Pedagogy

Different spaces inhabited by cultural influences can make it easier to make connections to structures already present in the brain [22]. Context can influence understanding by building on existing knowledge [23]. Such ideas are rooted in longstanding theories of cognition and education [24, 25].

Piaget's learning theory has previously described the active role that individuals play in constructing meaning within their own learning process through interaction with the environment [24, 25]. Assimilation of new information or experiences is integrated into a "schema" of our own existing knowledge and experiences through the interaction [26]. Wadsworth later described that such "schemata" (the plural of "schema") are akin to a cognitive filing system of "index cards" that can enable individuals to react to information [27]. As such, the theory emphasizes the value of context, scaffolding, and how educators can aid in building knowledge by drawing on existing experience, using differentiated teaching to suit the needs of unique individuals with their own experiences.

3.3 Considerations for LLMs

LLMs create text through the application of an initial notion and then apply statistics to make predictions about what the next word would be to develop phrases, sentences, and larger units of meaning [28]. Scholars have argued that LLMs are trained on human-developed text, but that the line between what is "natural" versus machine-made is blurry [28]. Creating responses in this way has been described as "discourse *sui generis*" [28, p. 7], since it is classified as "of its own kind" rather than being human-developed.

LLMs are based on the input data they are trained on, and such information can be incorrect or biased [29]. The fundamental approach taken by such algorithms can be impacted by majority views, which influence data training, and that this will, in turn, impact aspects of power dynamics and privilege [30]. Moreover, LLMs have been described as hallucinating, creating misleading outputs, lacking agency, having the potential for bias, lacking transparency in decision-making, and raising concerns around sustainability and the amount of energy they need, among other issues [31, 32, 33]. However, they also represent scalability, extensibility, and the potential to query large amounts of data to create content that could innovate on culture in different ways than a single individual may. Accordingly, we incorporated LLMs in our approach to creating contextually pertinent content rooted in CRP for Hispanic or Latine students.

4 Methods

We employed a two-phase research design structured to progress from broad pedagogical exploration to specific cultural contextualization, see Figure 1. We began Phase 1 with a broad scope, evaluating six prompt variations across six LLMs and six fundamental data structures ($n = 216$). This phase utilized generic descriptors of Hispanic or Latine identity to observe the baseline outputs. Building on these initial findings, Phase 2 narrowed down the focus to ten specific national identities across the southeast part of the U.S. This phase evaluated the ten nationalities with the same six LLMs and six data structures ($n = 360$). This phase investigated whether specifying nationality produced more culturally relevant examples. To interpret the results, we utilized discourse analysis, examining not only the correctness of the technical analogies but also the depth, accuracy, and nuance of the cultural references employed [34, 35].

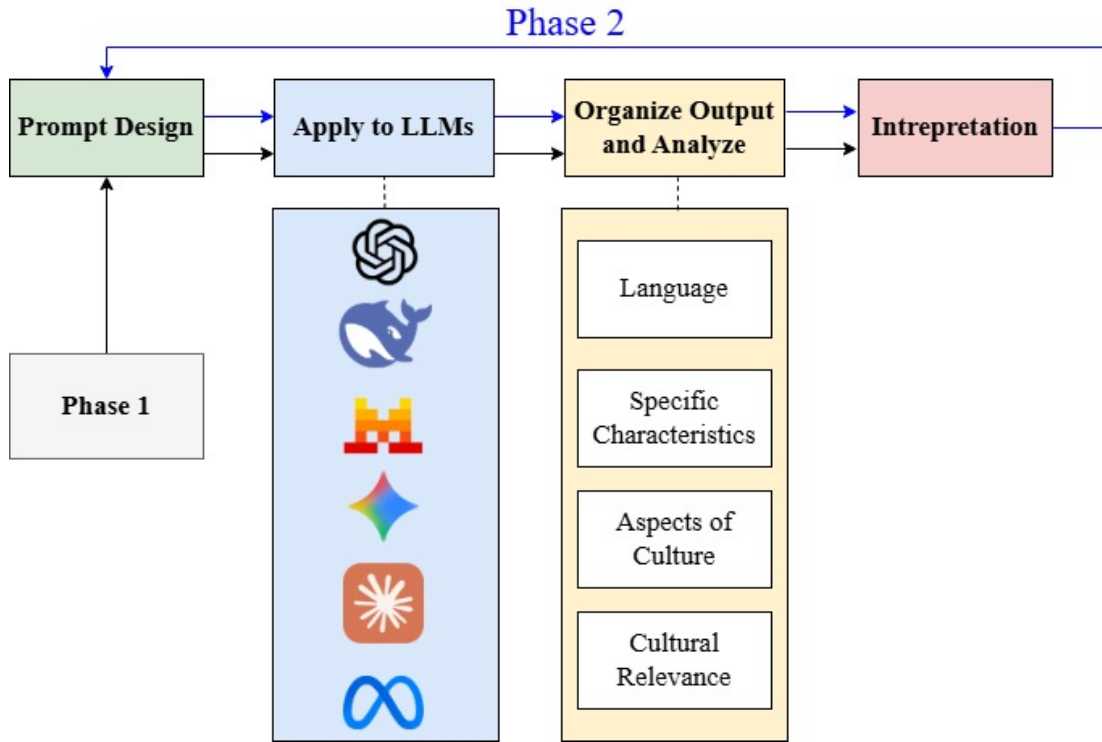


Figure 1: Overview of the two phase research approach

4.1 LLMs and Prompts

Six LLMs were compared in this study, selected based on their access and rate of recent citations. While each LLM possesses its own distinct technical specifications [36, 37], the selected set included OpenAI’s GPT-4, Google’s Gemini 2.0, Anthropic’s Claude 3.5 Sonnet, Meta’s Llama 3, Mistral (via Le Chat), and DeepSeek. We examined examples across six fundamental data structures: a list (or linked list), stack, queue, priority queue (PQ) or heap, binary search tree, and hash table.

4.2 Phase 1: Evaluation of Generalized Cultural Contexts

To determine the most effective phrasing, we initially designed six prompt variations ranging from generic student explanations to instructor-led, culturally specific scenarios:

1. *If you were to describe a [DATA STRUCTURE] to a student who is of Hispanic or Latin descent who is unfamiliar with them using an example that you think would help them to understand what it is, how would you do so?*
2. *If you were to describe a [DATA STRUCTURE] to a student who is of Hispanic or Latin descent who is unfamiliar with them using a culturally relevant example that you think would help them to understand what it is, how would you do so?*
3. *Pretend you are a computing student who is of Hispanic or Latin descent trying to explain a [DATA STRUCTURE] to a family member or friend who is unfamiliar with them using an example that you think would help them to understand what it is. How would you do so?*
4. *Pretend you are a computing student who is of Hispanic or Latin descent trying to explain a [DATA STRUCTURE] to a family member or friend who is unfamiliar with them using a culturally relevant example that you think would help them to understand what it is. How would you do so?*
5. *Pretend you are a computing instructor who is trying to explain a [DATA STRUCTURE] to a student of Hispanic or Latin descent who is unfamiliar with it. How would you do so?*
6. *Pretend you are a computing instructor who is trying to explain a [DATA STRUCTURE] to a student of Hispanic or Latin descent who is unfamiliar with them using a culturally relevant example that you think would help them to understand what it is. How would you do so?*

The data was gathered in the U.S. during February 2025, which can impact the outputs produced [38]. The combination of 6 prompts x 6 data structures x 6 LLMs yielded a total of $n = 216$ outputs used in the analysis. Note that the data collection was done twice with these combinations to determine if variability between runs could impact outcomes. However, an analysis of both sets of outputs illustrated that they were relatively consistent in terms of the aspects presented. We also used the default parameters, although memory was turned off in settings, and a fresh window or session was used between queries to avoid historical influence.

4.3 Phase 2: Nationality-Specific Application

Based on the outputs and cultural relevance of the outputs observed in Phase 1, we elected to modify the highest scoring of role-playing prompts to answer RQ3 by exploring specific nationalities: *Pretend you are a computing instructor who is trying to explain a [DATA STRUCTURE] to a student of [NATIONALITY] descent who is unfamiliar with it. Using a culturally relevant example that you think would help them understand what it is, explain how you would do so.*

We iterated this prompt across the six data structures (List, Stack, Queue, Priority Queue, BST, and Hash Table). For the demographic variable, we focused on the top 10 nationalities described as most frequent in the southeastern region of the United States (U.S.) by the 2023 Census Bureau, as shown in Table 1. In particular, our exploration prioritized Florida and Puerto Rico as defined by CAHSI's southeastern region of the U.S.

Nationality	Percent of Population
Puerto Rican	46.7%
Cuban	17.7%
Mexican	8.4%
Colombian	5.2%
Venezuelan	4.4%
Dominican (Dominican Republic)	4.0%
Nicaraguan	2.1%
Honduran	1.8%
Guatemalan	1.7%
Peruvian	1.6%

Table 1: Top Nationalities According to the U.S. Census Bureau, 2023

This second phase of data collection was conducted in April 2025. The combination of 10 nationalities $\times$ 6 data structures $\times$ 6 LLMs yielded a total of $n = 360$ outputs for the analysis in Phase 2.

4.4 Data Analysis

The researchers have previously described how analyzing outputs across LLMs can be approached [39]. In our study, the first and fourth authors applied the prompts to each LLM, generating a combined total of $n = 576$ outputs ($n = 216$ from the exploratory Phase 1 and $n = 360$ from the nationality-specific Phase 2). All outputs were stored and analyzed manually using a qualitative content analysis. The analysis was guided by the study’s broader research interest in understanding how LLMs represent and adapt cultural content in explanations of computing concepts. More specifically, we inductively examined whether the generated responses incorporated culturally grounded references [40], how those references were expressed, and whether they appeared meaningfully connected to the intended audience rather than simply translated linguistically. They categorized each response across four primary dimensions: Language (e.g., English vs. Spanish), Specific Characteristics (e.g., code snippets), Aspects of Culture (e.g., culinary, folklore), and Cultural Relevance (high/medium/low). To ensure reliability, the authors reviewed the classifications separately and then discussed them together to resolve ambiguities and establish consensus on their interpretations of the categories.

We assessed discourse through descriptions of surface culture (what is observable) and deep culture from shared experiences, previously [40]. The taxonomy defined aspects of culture in the context of arts, folklore, food (food and culinary contributions), heroes/personalities, history, and holidays (including patriotic holidays, religious observances, and personal rites and celebrations). Compared to others, deep culture was defined through many aspects such as ceremony and celebration, communication, ethics and discipline, and traditions. These distinctions helped guide how cultural references were identified and interpreted within the LLM outputs. Specifically, our analysis included evaluating each result for the “Language” of the response (English Only, Predominant English, Predominant Spanish, or only Spanish); “Specific Characteristics” pertinent to computing pedagogy (Step-By-Step Instruction, Code Snippets, Diagrams/Imagery); the “Cultural Relevance” (High, Medium, or Low); and “Aspects of Culture” (Arts, Folklore,

Culinary, Family & Social Structures, Holidays and Celebrations, Symbolism, Storytelling & Oral Traditions).

In this paper, we consider cultural relevance as the extent to which an output meaningfully incorporates cultural references, values, practices, or contexts in a way that goes beyond surface-level language to meaningfully connect to the cultural experiences of Hispanic or Latine learners. This meant that a response was not considered culturally relevant just because it was written in Spanish. Outputs were coded as high in cultural relevance when they incorporated specific and meaningful cultural references that clearly supported the explanation of the computing concept; medium when they included some relevant cultural elements but in a more limited or generic way; and low when they relied mainly on Spanish language use or contained little meaningful cultural grounding. Meanwhile, “cultural relevance” was subjectively determined by the first and fourth authors. The authors first reviewed outputs independently and then discussed their interpretations together to resolve ambiguity and reach a consensus.

In this paper, we consider cultural relevance as the extent to which an output meaningfully incorporates cultural references, values, practices, or contexts in a way that goes beyond surface-level language and beyond the use of Spanish alone, meaningfully connecting to the cultural experiences of Hispanic or Latine learners.

Outputs were coded as:

- **High in cultural relevance:** When they incorporated specific and meaningful cultural references that clearly supported the explanation of the computing concept (e.g., explaining a linked list by describing the specific ingredients and sequential preparation steps of a regional dish like Guatemalan *Fiambre*, where each element conceptually points to the next).
- **Medium in cultural relevance:** When they included some relevant cultural elements but in a more limited or generic way (e.g., using common foods like “tacos” or “empanadas” to explain a stack, without a deeper connection to the culture).
- **Low in cultural relevance:** When they relied mainly on Spanish language use or contained little meaningful cultural grounding (e.g., using a standard, generic textbook example of waiting in a line at a bank but simply translating the text to Spanish or swapping in names like “Maria” and “Jose”).

4.5 Reflexivity

Given that there are many ways to interpret text [35], which may influence observations, we want to describe the authors’ backgrounds. The first author co-led the analysis and interpretation and led the writing of the manuscript. He is a Master’s student in Human-Centered Computing who identifies as Hispanic (Guatemalan) and a Spanish speaker, which influences how he perceives and interprets the data in this study. His cultural background and Spanish assist him in noticing when AI-generated explanations reflect meaningful cultural references versus when they rely on generalized or stereotypical depictions. Recognizing the diversity among Hispanic and Latine students in the U.S., this author acknowledges that their experiences do not represent the full range of perspectives reflected in the data.

The second author was involved with the study planning, interpretation, writing, and editing. She is an assistant professor at an HSI in the southeastern region of the U.S. She identifies as a White and non-Hispanic woman. She has previously created and led a course on AI ethics. Her perspective as an educator shaped her perspective and regard for the potential pedagogical applications and implications of LLMs. It also drove her interest to explore how content could be tailored to better serve students at HSIs.

The third author contributed to the study planning, interpretation, and editing. He is an African-American, non-Hispanic scholar whose perspective is shaped by cultural experiences and the barriers faced as a member of a marginalized group in computing. He is an associate professor with experience teaching at two Historically Black Colleges and Universities (HBCUs) and one minority-serving institution (MSI). These experiences also led him to explore opportunities to tailor curricular content for HSIs, drawing on practices from both his current and prior institutions.

The fourth author co-led the analysis and interpretation. He is a Hispanic (Dominican) undergraduate computing student currently attending a Hispanic-Serving Institution (HSI) and is bilingual in English and Spanish. Their personal experience as a computing student attending an HSI, as well as being a fluent Spanish speaker, impacted the interpretation of the quality of generated responses.

The fifth author contributed to editing the manuscript. She is an African-American, non-Hispanic scholar whose work is driven by a deep commitment to research equity in STEM education and expanding participation in STEM. She is an associate professor at a predominantly white institution in the Midwest.

5 Results

This section details the findings from the exploratory Phase 1 and the nationality-specific Phase 2. We present the discourse analysis of the outputs generated by the six LLMs across the varied prompts.

5.1 Phase 1

An overview of the findings around the discourse analysis produced by different prompts and LLM models is shared in Table 2. Each corresponding prompt and data structure examined is described in terms of the analysis of the output. Although all models were prompted in English, responses were presented as “English Only” (EO) outputs. Instead, the responses given often included terms or phrases in Spanish, as indicated by the high prevalence of responses in the “Mostly English” (ME) category. As indicated, in some cases, there was also a mix of both English and Spanish (BES), “Mostly Spanish” (MS), or even “Spanish Only” (SO) responses. The complete set of data outputs can be shared upon request to the authors.

The specific characteristics related to pedagogy included instances when an output went beyond giving a strict definition. Although such examples were not always included, the most common approach offered was to provide step-by-step instructions. Diagrams and/or imagery were the least likely option to appear.

In addition to the specific ratings, the authors summarized their interpretation as it pertained to the aspects of culture. It was noted that all models used similar language, providing outputs in mostly

English with some Spanish phrases, except for a few prompts. Responses tended to follow a similar format: 1) Give an example for the data structure in question; 2) Explain how the example applies to the data structure; and 3) Sometimes provide further technical details to give more context as to the usefulness of the data structure.

ChatGPT. Gave a broad representation of Latin cultures (Puerto Rico, Colombia, Dominican Republic, etc.), but a majority of the cuisines mentioned came from Mexican culture (e.g., “Tamales,” “Tortillas,” or “Tacos”). For instance, an output described a linked list in the context of waiting in line for tamales at a festival.

When discussing a binary search tree, ChatGPT mentioned family across all six prompts, such as describing examples with “Abuelita” (grandmother) or “Primo Luis” (Cousin Luis). Most cultural celebrations also came from Mexican culture. When providing explanations, ChatGPT offered technical steps alongside cultural analogies.

Mistral. A majority of the cuisines mentioned come from Mexican culture (e.g., “Tamales,” “Tacos,” and “Carne Asada”), but also include mentions of Puerto Rican dishes like “Arroz con Gandules” and “Bistec Encebollado.” Some symbolism described included “Dominoes,” “Loteria” (played in various Latin countries), and “Alebrijes” (which are folk art sculptures that originated from Mexico). It also mentioned “Don Quixote” (a novel by Miguel de Cervantes). Overall, it offered logical and technical clarity with structured explanations. The dialogue was given in Spanish and English, running scenarios. Although the steps provided were thorough, it was noted that “it felt like I was reading them from a textbook.”

Meta’s AI. For a majority of the responses, the ending gave a Spanish word meaning “understand,” and also wrote a response in Spanish stating that hopefully this was helpful. Meta AI provided other information apart from cuisine, but some of the technical terms were a bit simplified. Mexican culture was mostly represented in the responses, something that felt more pronounced with the names used (e.g., Diego, Maria, and Luis). In terms of cuisine, “Tortillas” and “Tacos” were mentioned, as well as “Empanadas.” Some answers around music and literature were given, e.g., “Bad Bunny,” “J Balvin,” “Shakira,” “Cien Años de Soledad by Gabriel García Márquez” (One Hundred Years of Solitude).

DeepSeek. The cuisine examples featured “Tortillas,” “Tacos,” and “Tamales.” In terms of familial relationships, ideas were shared like “Priority Queue: Children get priority to eat.” DeepSeek placed a greater emphasis on explaining why the examples it provided were appropriate, rather than elaborating on the data structure itself (like Gemini did). It did occasionally give multiple examples in a single response, and some outputs were entirely in Spanish.

Gemini. A majority of cuisine items again came from Mexican culture (e.g., “Tacos” and “Tortillas”) but also referenced “Sancocho” (a general term for stew common throughout the Caribbean and Latin America). In terms of celebrations, it mentioned “Quinceañeras,” a traditional Latin American celebration marking a girl’s 15th birthday. It included a lot of storytelling within responses relative to other models and sometimes included multiple examples within one response. It was seen as providing “good examples and depth” and the greatest amount of technical details beyond the example.

Prompts		Language (English versus Spanish)*					Specific Characteristics**			Aspects of Culture							Cultural Relevance		
		EO	ME	BES	MS	SO	SbS	CS	DI	Arts	Folklore	Culinary	Family & Social Structures	Holidays & Celebrations	Symbolism	Storytelling & Oral Traditions	High	Medium	Low
1	LL	MI	CG, MT, DS, GE, CL							MI, DS		CG, MT	GE	DS, GE	CL		MI, MT, DS, GE		CG, CL
	Stack	MI	CG, MT, DS, GE, CL									CG, MI, MT, DS, GE, CL		DS			MI, MT, DS, GE	CL	CG
	Queue		CG, MI, MT, DS, GE, CL									CG, MI, MT, DS	GE		CL		MI, MT, DS	CG, GE, CL	
	PQ/Heap	MI	CG, MT, DS, GE, CL									CG, MI, MT, DS	MI, DS, GE	MI	CL		MI, MT	CG, DS	GE, CL
	BST	MI	CG, MT, DS, GE, CL				CG, MI, MT	MT	CG			DS	CG, DS, CL	MI	GE		MI, MT	DS, CL	CG, GE
	HT	MI	CG, MT, DS, GE, CL									CG, MI, MT, DS	DS, GE		GE, CL		MI	CG, MT, DS, GE	CL
2	LL		CG, MI, MT, DS, GE, CL				CL					CG, MI, DS, GE	DS	MT, DS, CL			MI, MT, DS	GE, CL	CG
	Stack		CG, MI, MT, DS, GE, CL				CL					CG, DS, GE, CL	DS, CL				MI, MT, DS, GE	CL	CG
	Queue	MI	CG, MT, DS, GE, CL				CL					CG, MI, MT, DS, CL	GE	DS			MI, MT, DS	CG, GE, CL	
	PQ/Heap	MI	CG, MT, DS, GE, CL				CG, MT					CG, DS, GE	DS, GE, CL	CL			MT, DS, GE	CG, MI, CL	
	BST	MI	CG, MT, DS, GE, CL				CG, MI, MT					CG, DS, GE	CG, DS, CL				MI, MT, DS	CG, GE, CL	
	HT	MI	CG, MT, DS, GE, CL				CL					CG, MI, MT, DS, CL	DS, GE	GE			MI, MT, DS, GE	CG, CL	
3	LL	MI	CG, MT, DS, GE, CL				CG, MI, MT					GE, CL	CG, DS, GE, CL		MI	GE	MT, DS, GE, CL	CG, MI	
	Stack	MI	CG, MT, DS, GE, CL				CG, MI, MT					CG, MT, DS, GE, CL	DS, GE			GE	MI, MT, DS, GE, CL		CG
	Queue		CG, MI, MT, DS, GE, CL				CG, MI, MT					CG, MT, GE, CL	DS, GE			GE	MI, MT, GE, CL	CG, DS	
	PQ/Heap		CG, MT, DS, GE, CL		MI		CG, MI, MT					GE	CG, MI, DS, CL	MI	GE	GE	MI, MT, DS	GE, CL	CG
	BST	MI	CG, MT, DS, GE, CL				CG, MI, MT					MT, GE	CG, DS, CL	MI	MI	GE	MI, MT, GE	DS, CL	CG
	HT		CG, MI, MT, DS, GE, CL				CG, MI, MT			MI		CG, MT, DS, CL	DS, GE		GE	GE	GE	CG, MI, MT, DS, CL	
4	LL		CG, MT, DS, GE, CL	MI			CG, MI, MT					CG, MI, MT, GE, CL	DS, CL	DS		GE	MI, MT, DS, GE, CL		CG
	Stack		CG, MT, DS, GE, CL		MI		CG, MI, MT					CG, MI, MT, DS, GE, CL				GE	MI, MT, DS, GE	CG, CL	
	Queue		CG, MT, DS, GE, CL		MI		CG, MI, MT					CG, MI, MT, DS, GE, CL	DS, GE, CL	DS		GE	MI, MT, DS, GE	CG, CL	
	PQ/Heap		CG, MT, DS, GE, CL		MI		CG, MI, MT					CG, MT, DS, GE, CL	MI, DS, CL	DS		GE	CG, MI, MT, DS, GE	CL	
	BST		CG, MT, DS, GE, CL		MI		CG, MI, MT					CG, MI, MT, GE	CG, DS, GE, CL	DS		GE	MI, MT, DS, GE	CG, CL	
	HT		CG, MT, DS, GE, CL		MI		CG, MI, MT					CG, MT	MI, DS	CG, DS	GE, CL	GE	MI, MT, DS, CL	CG, GE	
5	LL	CL	CG, MT, GE		MI	DS	CG, MI, MT	CL				MT, CL	DS, GE	DS, GE	CL		MI, MT, DS, GE	CG	CL
	Stack	CL	CG, MT, GE		MI	DS	CG, MI, MT, DS					CG, MI, MT, DS, GE, CL					CG, MI, MT, DS	GE, CL	
	Queue	CL	CG, MT, GE		MI	DS	CG, MI, MT, DS					CG, MT, GE	DS, GE		CL		CG, MI, MT, DS, GE		CL
	PQ/Heap	CL	CG, MT, GE		MI	DS	CG, MI, MT, DS	DS	DS			CG	CG, DS		GE, CL		MI, MT, DS	CG	GE, CL
	BST	CL	CG, MT, GE		MI	DS	CG, MI, MT, DS	DS	DS			MI, GE, CL	CG, DS, GE, CL	CG, DS			MI, MT, DS, GE, CL	CG	
	HT	MI, CL	CG, MT, GE			DS	CG, MI, MT	CL				CG, MI, GE	DS, GE	DS	CL		CG, MI, MT, DS, GE		CL
6	LL	CL	CG, MI, MT, GE			DS	CG, MI, MT			CG		CL	CG, GE, CL	CG, MI, GE	MI		CG, MI, MT, DS, GE, CL		
	Stack	MI, CL	CG, MT, GE			DS	CG, MI, MT	CL				CG, MI, MT, GE, CL	DS	DS, GE			CG, MI, MT, DS, GE, CL		
	Queue	MI, CL	CG, MT, GE			DS	CG, MI, MT, DS					CG, MI, MT, DS, CL	GE				CG, MI, MT, DS, GE	CL	
	PQ/Heap	MI, CL	CG, MT, GE			DS	CG, MI, MT, DS					CG, MT, DS, GE, CL	DS	DS, GE	GE		MI, MT, DS, GE	CG, CL	
	BST	MI, CL	CG, MT, GE			DS	CG, MI, MT, DS					GE	CG, DS, GE, CL	CG, MI, DS, CL			MI, MT, DS, GE, CL	CG	
	HT	MI, CL	CG, MT, GE			DS	CG, MI, MT, DS					CG, MI, GE	MI, DS, GE, CL	DS, GE, CL			CG, MI, MT, DS, GE, CL		

*EO=English Only; ME = Mostly English; BES = Both English & Spanish; MS = Mostly Spanish; SO = Spanish Only;

**SbS = Step-by-Step Instruction; CS = Code Snippets; DI = Diagrams/Imagery

Corresponding LLM Output: CG = ChatGPT; MI = Mistral; MT = Meta; DS = DeepSeek; GE = Gemini; CL = Claude

Table 2: Comparing LLM outputs for each prompt in relation to language, specific characteristics valuable for pedagogy, aspects of culture, and cultural relevance (**Phase 1**)

Claude. Again, a majority of the cuisines mentioned were Mexican, e.g., “Tortilla” and “Taco.” Responses would occasionally switch between English and Spanish within a response. For instance: Prompt 1, Q3 Response begins in English: “I’d use the example of a line.” The following paragraph was then in Spanish: “Imagina la fila en una panadería...” (“Imagine the line at a bakery”). The next paragraph was again in English: “This is exactly how a queue works...” It was noted that it is unclear how useful it is to provide a scenario in Spanish but to provide the technical details in English. As a result, Claude was subjectively perceived as providing less helpful responses, with some confusing language choices and a lack of technical details, compared to the other models.

Overall Phase 1 Summary. Generally, there was a relative lack of culturally relevant examples for specific Hispanic populations that are not Mexican/Central American. Responses typically used either Mexican cultural references or broadly applicable Hispanic cultural references, such as music or ubiquitous dishes, like rice and beans. At times, responses bordered on being stereotypical, with frequent mentions of “Marias,” “Carmens,” “Tacos,” and “Tamales.” DeepSeek, Gemini, and Claude were seen as providing less unique examples than Meta AI, Mistral, and ChatGPT. Representative examples of these model outputs are shared in Table 3.

Phase 1		
Prompt #4 (Linked List)		
LLMs:	<i>ChatGPT</i>	<i>Gemini</i>
Examples:	“Okay, imagine we’re at a big family party. You know how we sometimes have a line for the food, right? Maybe the first person in line is Tía Carmen, the second person is Tío Juan, and the third is Cousin Luis. Now, here’s the thing: In this line, each person knows who’s next, right? Tía Carmen knows that after her, Tío Juan is next, and Tío Juan knows that after him, it’s Cousin Luis. It’s like they’re all holding hands in a line, passing the baton, one after the other.” Now, a linked list in computer science works kind of like that. It’s a series of elements (like people in the line), where each one ‘knows’ who comes next. The first element is the head (Tía Carmen), and it points to the next element (Tío Juan), who in turn points to the next (Cousin Luis), and so on. Each one is ‘linked’ to the next one, and if we want to get to the third person in the line, we have to follow the links: first go to Tía Carmen, then Tío Juan, and finally Cousin Luis.”	Me: “Okay, so you know how we make tamales for Christmas, right? It’s a whole process, everyone helps out, and we have to do things in a certain order.” Family Member: “Yeah, like soaking the corn husks, then making the masa, then the filling...” Me: “Exactly! Think of a linked list like that recipe. Each step in the recipe is like a ‘node’ in the list. It has the instructions for that step – like ‘spread the masa on the husk’ – and then it also has a little note that says, ‘next, add the filling.’ That note is like a ‘pointer’ – it links that step to the next one.” Family Member: “So it’s like a chain of instructions?” Me: “Precisely! And the best part is, if we decide to add a new step, like ‘add a special salsa,’ we don’t have to rewrite the whole recipe. We just find the right spot in the chain, add the new instruction, and adjust the pointers. It’s much easier than trying to rearrange a whole list of instructions.” Family Member: “Okay, I get it. So each step has the instruction and the link to the next one.” Me: “Yep! And just like with the tamales, there’s a first step – the ‘head’ of the list – and a last step – the ‘tail.’ That’s how we know where the recipe starts and ends.” Family Member: “So this is how computers keep track of things in order?” Me: “Exactly! Linked lists are super useful for all sorts of things, like keeping track of a playlist, or managing a list of customers. It’s all about having those little pointers that connect everything in the right order.”

Table 3: Example outputs for Phase 1

5.2 Phase 2

Phase 2 examined pedagogical examples and the depth of cultural integration with relation to particular nationalities identified as prevalent in the southeastern region of the U.S., as shown in Table 4, which spans across the next two pages. The specific characteristics related to pedagogy included instances when an output went beyond giving a strict definition. Although such examples were not always given, the most common approach offered was to provide step-by-step instructions. Diagrams and/or imagery were the least likely option to appear.

ChatGPT. The quantitative analysis highlights a heavy reliance on culinary themes appearing in 65.00% of its responses. The data revealed that this model produced the most nuanced, regionally specific metaphors when prompted with a specific nationality. For example, to explain a linked list, the LLM discussed the Guatemalan salad, *Fiambre* [41], where individual ingredients represented connected nodes. For Nicaraguan contexts, it explained Stacks using a pile of *Vigorón* plates at a festival [42]. Additionally, it mapped hash tables to the *Colectivo* bus routes in Honduras, using route numbers as keys. Despite these rich narratives, it rarely supported its explanations with visual aids, generating diagrams in only 15.00% of cases.

Mistral. While the model referenced culinary dishes specific to each nationality, such as citing “Baleadas” [43] for Honduras or “Ajiaco” [44] for Colombia, it struggled with repetition and variety. Ultimately, it received the highest percentage of “Low” cultural relevance ratings at 51.67%. Researchers noted the delivery felt textbook-like because it reused identical metaphors across different cultures. For instance, it applied the analogy of “organizing a recipe book” to explain hash tables for Puerto Rican, Cuban, Mexican, Venezuelan, and Nicaraguan contexts by merely swapping the dish names. However, when deviating from these templates, it produced highly specific cultural markers. Notable examples include explaining Binary Search Trees using a collection of Salsa CDs referencing Héctor Lavoe for Puerto Rico or sorting national team soccer jerseys referencing Paolo Guerrero for Peru.

Meta’s AI. The model achieved a “High” cultural relevance rating in 56.67% of responses, distinguishing itself by incorporating the highest frequency of “Arts” references (26.67%) among all platforms. It adopted a conversational, familial tone, often addressing the learner as “Primo” or “Prima.” Beyond cuisine, it utilized highly specific geographical and cultural markers. For instance, in the Colombian context, it referenced the Vallenato festival in Valledupar [45],” explicitly naming artists like Leandro Díaz and Silvestre Dangond to explain Binary Search Trees.

DeepSeek. This model excelled in pedagogical clarity by consistently employing the “Concept-Definition-Analogy” structure. It often began with a relatable real-world scenario before introducing the technical term, a strategy known as scaffolding. For instance, when explaining Linked Lists, it used the analogy of a “Treasure Hunt” where each clue points to the next location, directly mapping the concept of nodes and pointers. DeepSeek also demonstrated high cultural adaptability. For the Mexican context, it explained Hash Tables using a “Posada” gift exchange where nicknames serve as keys to find specific gift bags [46]. In the Colombian context, it used the “Feria de las Flores” parade to explain Priority Queues, where floats are ordered by importance rather than arrival time [44]. This ability to seamlessly integrate deep cultural markers with rigorous technical explanations contributed to its top-tier performance ratings.

Prompts	Nationalities	Data Structures	Language (English versus Spanish)				Specific Characteristics			Aspects of Culture						Cultural Relevance					
			EO	ME	BES	MS	SO	Sbs	CS	DI	Arts	Folklore	Culinary	Family & Social Structures	Holidays & Celebrations	Symbolism	Storytelling & Oral Traditions	High	Medium	Low	
Puerto Rico		LL	CG.MLMT	GL,CL			DS	CG.MLMT,GE,DS,CL					CG,GE,DS	DS	CG.MLMT,GE,CL				CG.ML,GE,DS,CL		MT
		Stack	MLMT	CG,GE,CL			DS	CG.MLMT,GE,DS,CL					CG.MLMT,GE,DS,CL	DS	CG,CL				CG.MLMT,GE,DS,CL		
		Queue	MLMT	CG,GE,CL			DS	CG.MLMT,GE,DS,CL					CG.MLMT,GE,DS,CL		MI				CG.MLMT,GE,DS,CL		
		PQ/Heap	CG.MLMT	CL			GE,DS	CG.MLMT,GE,DS,CL			DS		GE	CG,GE	GE		CL		GE,DS	CL	CG.MLMT
		BST	CG.MLMT	CL			GE,DS	CG.MLMT,GE,DS,CL		CL			GE	CG,GE,DS,CL	CG,DS,CL		MT		MT,GE,DS	CG,CL	MI
		HT	MLMT	CG,CL			GE,DS	CG.MLMT,GE,DS,CL					CG.MLMT,GE	DS			DS,CL		CG.MLMT,GE,DS	CL	
Cuban		LL	CG.MI,GE	MT,CL			DS	CG.MLMT,GE,DS,CL	DS	MT	MT,GE,CL		MLDS				CG		MT,GE	CG.MLDS,CL	
		Stack	CG.MI	MT,GE,CL			DS	CG.MLMT,GE,DS,CL	DS	CG,MT			CG.MLMT,GE,DS,CL	DS				MLMT,GE,DS,CL	CG		
		Queue	MI	CG.MT,GE,CL			DS	CG.MLMT,GE,DS,CL	DS	CG,MT			CG.MT,GE,DS,CL				DS		DS	CG,GE,CL	MLMT
		PQ/Heap	CG.MLMT	GE,CL			DS	CG.MLMT,GE,DS,CL	DS	CG	DS		CL	CL	CL		GE,DS		CL	DS	CG.MLMT,GE
		BST	CG.MI	MT,GE,CL			DS	CG.MLMT,GE,DS,CL	DS	MT,DS	MT		CG.MT,DS	ML,CL			GE,DS		MT	DS	CG.MLGE,CL
		HT	CG.MI	MT,GE,CL			DS	CG.MLMT,GE,DS,CL	DS	DS			CG,MT	GE,DS	GE		CL		MT,DS	CG,GE,CL	MI
Mexican		LL	CG.MI	MT,GE,CL			DS	CG.MLMT,GE,DS,CL					CG.MT,GE,DS,CL		GE,DS				GE,DS,CL	CG	MLMT
		Stack	CG.MI	MT,GE,CL			DS	CG.MLMT,GE,DS,CL					CG.MLMT,GE,DS,CL					MT,GE,CL	CG.MLDS		
		Queue	CG.MI	MT,GE,CL			DS	CG.MLMT,GE,DS,CL	DS				CG.MLMT,GE,DS,CL				GE		MT,GE,DS	ML,CL	CG
		PQ/Heap	CG.MI	MT,GE,CL			DS	CG.MLMT,GE,DS,CL	DS				MT,DS	GE	CG,GE		CL		GE	CG,DS	MLMT,CL
		BST	MI	CG.MT,GE,CL			DS	CG.MLMT,GE,DS,CL	DS	DS	MT		GE,DS	CG.MLDS,CL	DS				MT,GE,DS	CG	ML,CL
		HT	CG.MI	MT,GE,CL			DS	CG.MLMT,GE,DS,CL	DS	DS	MI		MT,GE,DS		CG		GE,CL		MT,GE,DS	CG,CL	MI
Colombian		LL	MLMT	CG,GE,CL			DS	CG.MLMT,GE,DS,CL	CL	MLDS	MLMT		CG,DS				GE,CL		ML,GE,CL	CG,DS	MT
		Stack	MI	CG.MT,GE,CL			DS	CG.MLMT,GE,DS,CL	CL	MLDS			CG.MLMT,GE,DS,CL					CG.MLMT,GE,DS,CL	CG.MLMT		
		Queue	MI	CG.MT,GE,CL			DS	CG.MLMT,GE,DS,CL	CL	CG,DS			MLMT,GE		CG		GE,DS,CL		MT,GE,DS	CG,MI	
		PQ/Heap	CG.MLMT	GE,CL			DS	CG.MLMT,GE,DS,CL	CL	CG.MT,DS	DS						GE,CL		DS		CG.MLMT,GE,CL
		BST	CG.MI	MT,GE,CL			DS	CG.MLMT,GE,DS,CL	CL	MLMT,DS,CL	MT		GE	CG,MI	CL		DS		MT,CL	GE,DS	CG,MI
		HT	CG.MLMT	GE,CL			DS	CG.MLMT,GE,DS,CL	CL	CG.MT,DS	MLMT		DS,CL				GE,CL		MLDS	GE,CL	CG,MT
Venezuelan		LL	CG.MLMT	GE			DS,CL	CG.MLMT,GE,DS,CL					CG.MLMT,GE,DS,CL				DS		MLMT,DS,CL	CG,GE	
		Stack	CG.MLMT	GE			DS,CL	CG.MLMT,GE,DS,CL	DS	CL			CG.MLMT,GE,DS,CL						MLMT,DS,CL	CG,GE	
		Queue	CG.MLMT	GE			DS,CL	CG.MLMT,GE,DS,CL	DS	CL			CG.MLMT,GE,DS,CL						MLMT,DS,CL	CG,GE	
		PQ/Heap	CG.MLMT	GE			DS,CL	CG.MLMT,GE,DS,CL	DS	CL			CG.MLMT,GE,CL				DS		MLMT,DS,CL	CG,GE	
		BST	CG.MLMT	GE			DS,CL	CG.MLMT,GE,DS,CL	DS	DS,CL	MT						GE,DS,CL		DS	CL	CG.MLMT,GE
		HT	CG.MLMT	GE			DS,CL	CG.MLMT,GE,DS,CL	DS	DS,CL	MT		GE,DS	CG,MT,GE,DS	DS,CL		CL		DS,CL	GE	CG.MLMT

6

*EO=English Only; ME = Mostly English; BES = Both English & Spanish; MS = Mostly Spanish; SO = Spanish Only;

**Sbs = Step-by-Step Instruction; CS = Code Snippets; DI = Diagrams/Imagery

Corresponding LLM Output: CG = ChatGPT; MI = Mistral; MT = Meta; DS = DeepSeek; GE = Gemini; CL = Claude

Table 4: Comparing LLM outputs for nationality-specific prompts (Phase 2)

Prompts	Nationalities	Language (English versus Spanish)						Specific Characteristics			Aspects of Culture						Cultural Relevance		
		EO	ME	BES	MS	SO	Sbs	CS	DI	Arts	Folklore	Culinary	Family & Social Structures	Holidays & Celebrations	Symbolism	Storytelling & Oral Traditions	High	Medium	Low
Dominican	LL	MLMT,CL	GE	CG		DS	CG,MI,MT,GE,DS,CL		CL	MT,CL		CG,MI		CG	GE,DS		CG,MLDS	GE	MT,CL
	Stack	CG,MLMT,CL	GE			DS	CG,MLMT,GE,DS,CL	CG	ML,CL		CG,MLMT,GE,DS,CL						CG,MLMT,DS,CL	GE	
	Queue	CG,MLMT,CL	GE			DS	CG,MLMT,GE,DS,CL	CG	ML,CL		CG,MLMT,GE,DS,CL				GE,CL		MT,CL	DS	CG,MI,GE
	PQ/Heap	CG,MLMT,CL	GE			DS	CG,MLMT,GE,DS,CL		MI		CG,MLMT,GE,DS,CL		CG		GE,DS,CL			CG,CL	MLMT,GE,DS
	BST	CG,MT,CL	ML,GE			DS	CG,MLMT,GE,DS,CL	CG	CG,MLDS,CL		GE,DS	MT,DS,CL	DS		GE,DS,CL		MLDS	GE,CL	CG,MT
	HT	CG,MLMT,CL	GE			DS	CG,MLMT,GE,DS,CL	CG	MI		MT,GE	GE,DS	DS		GE,DS,CL		MT,GE	DS	ML,CL
	LL	MI	CG,MT,CL			GE,DS	CG,MLMT,GE,DS,CL				CG,MI	CG,MT,CL			GE,DS		CG,MLMT,DS,CL	GE	
	Stack	CG,MLMT	CL			GE,DS	CG,MLMT,GE,DS,CL				CG,MLMT,GE,DS,CL	CL			GE,DS		CG,MLMT,DS,CL		ML,MT
Nicaraguan	Queue	CG,MLMT	CL			GE,DS	CG,MLMT,GE,DS,CL				CG,MLMT,GE,DS,CL				CG,GE,DS,CL		ML,MTDS,CL		CG,GE
	PQ/Heap	CG,MLMT	CL			GE,DS	CG,MLMT,GE,DS,CL		DS		CL				CG,GE,DS,CL			CL	CG,MLMT,GE,DS
	BST	CG,MT	CL	MI		GE,DS	CG,MLMT,GE,DS,CL		MLDS	MT	CG,MI	CG,CL			GE,DS		MI	GE,DS,CL	CG,MT
	HT	CG	ML,MT,CL			GE,DS	CG,MLMT,GE,DS,CL		MLDS		CG,MLMT,GE,DS,CL				GE,DS,CL		MLMT	CG,GE,DS	
	LL	CG,MLMT	CL			GE,DS	CG,MLMT,GE,DS,CL			MT	CG,MLMT,GE,DS,CL				GE,CL		CG,MLGE,DS,CL		MT
	Stack	CG,MLMT	CL			GE,DS	CG,MLMT,GE,DS,CL	DS		GE,DS	CG,MLMT,GE,DS,CL				CL		DS,CL	GE	CG,MLMT
	Queue	CG,MLMT	CL			GE,DS	CG,MLMT,GE,DS,CL	DS			CG,MLMT,GE,DS,CL				GE,CL		MT,CL	CG,GE,DS	MI
	PQ/Heap	CG,MLMT	CL			GE,DS	CG,MLMT,GE,DS,CL	DS		MI	CL				GE,DS,CL			CL	CG,MLMT,GE,DS
Honduran	BST	CG,MI	MT,CL			GE,DS	CG,MLMT,GE,DS,CL	DS	CG,DS		CG,GE,DS	ML,CL			MT		MTDS,CL	GE	CG,MI
	HT	CG,MI	MT,CL			GE,DS	CG,MLMT,GE,DS,CL	DS	DS	MI	CG,MT,GE		GE		DS,CL		MT,GE	DS,CL	CG,MI
	LL	CG,ML,GE	MT,CL			DS	CG,MLMT,GE,DS,CL	CL	CL	GE	CG,MLDS		MT		CL		CG,MT,GE,DS,CL		MI
	Stack	MLMT,GE	CG,CL			DS	CG,MLMT,GE,DS,CL	CG,CL		MT	CG,MLGE,DS,CL						MTDS,CL	CG,MLGE	
	Queue	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL	CL			CG,MLMT,GE,DS,CL				DS,CL		CG,MT,CL	ML,GE,DS	
	PQ/Heap	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL	CG,CL			CG,MLMT,GE,DS,CL				GE,DS,CL				CG,MLMT,GE,DS
	BST	CG,MI	MT,CL			GE,DS	CG,MLMT,GE,DS,CL	DS	CG,DS		CG,GE,DS	ML,CL			MT		MTDS,CL	GE	CG,MI
	HT	CG,MI	MT,CL			GE,DS	CG,MLMT,GE,DS,CL	DS	DS	MI	CG,MT,GE		GE		DS,CL		MT,GE	DS,CL	CG,MI
Guatemalan	LL	CG,ML,GE	MT,CL			DS	CG,MLMT,GE,DS,CL	CL	CL	GE	CG,MLDS		MT		CL		CG,MT,GE,DS,CL		MI
	Stack	MLMT,GE	CG,CL			DS	CG,MLMT,GE,DS,CL	CG,CL		MT	CG,MLGE,DS,CL						MTDS,CL	CG,MLGE	
	Queue	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL	CL			CG,MLMT,GE,DS,CL				DS,CL		CG,MT,CL	ML,GE,DS	
	PQ/Heap	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL	CG,CL	CL	MT					GE,DS,CL				CG,MLMT,GE,DS,CL
	BST	CG,MI	MT,CL			GE,DS	CG,MLMT,GE,DS,CL	CL	CG,CL	MT					GE,DS,CL				CG,MLMT,GE,DS,CL
	HT	CG,MLMT,GE	CL			DS	CG,MLMT,GE,DS,CL	CL	CG,CL	MT	CG	CL	CL		GE,DS		ML,MTDS	CG,CL	GE
	LL	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL	CG,DS,CL	CL	MI	CG		MT		GE,DS,CL		ML,MTDS	CG,CL	
	Stack	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL		CL	CG	CG,MLMT,GE,DS,CL		CL				ML,MTDS,CL	CG,GE	
Peruvian	Queue	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL		DS	MT					GE,CL		MLDS	GE,CL	CG,MT
	PQ/Heap	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL				CG,MLMT,GE,DS,CL				GE,CL		MLGE	CG,DS,CL	MT
	BST	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL			MT					GE,DS,CL		DS	CL	CG,MLMT,GE,DS,CL
	HT	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL		DS,CL	CG	CG,MLMT,GE,CL				DS		CG,MT,GE,CL	DS	MI
	LL	CG,MLMT	GE,CL			DS	CG,MLMT,GE,DS,CL		MTDS	CF,MT	CG,MLDS				GE,CL		CG	GE,DS,CL	MLMT

*EO=English Only; ME = Mostly English; BES = Both English & Spanish; MS = Mostly Spanish; SO = Spanish Only;

**SbS = Step-by-Step Instruction; CS = Code Snippets; DI = Diagrams/Imagery

Corresponding LLM Output: CG = ChatGPT; MI = Mistral; MT = Meta; DS = DeepSeek; GE = Gemini; CL = Claude

Gemini. This model distinguished itself through a heavy reliance on narrative storytelling, achieving a “High” cultural relevance rating in 40.00% of its responses. Rather than simple definitions, it often explained explanations as complete scenarios—such as describing a “Baking Disaster” to explain exceptions—weaving technical concepts into relatable stories. For specific cultural contexts, it moved beyond standard Mexican tropes to reference broader Latin American traditions, such as “Quinceañeras” [47]. The Caribbean stew “Sancocho” [48]. At the same time, its narrative depth was strong; however, quantitative analysis reveals a significant limitation: Gemini produced zero diagrams or imagery (0.00%) across all analyzed outputs, relying entirely on text to convey spatial data structures, such as trees and graphs.

Claude. Claude achieved a strong “High” cultural relevance rating of 48.33%, but its performance was mixed due to linguistic inconsistencies. The model frequently provided code snippets (80.00% of responses) and utilized diagrams (28.33%), often employing clear, visual metaphors—such as comparing “classes to cookie cutters” and “Objects to cookies.” For specific cultural contexts, it referenced traditions like the Colombian “Feria de las Flores” and the Guatemalan “Fiesta de la Virgen de la Asunción” [49].

Overall Phase 2 Summary. The analysis shows a clear split between using general prompts versus specific nationality prompts. When we used general prompts, the models tended to default to Mexican cultural references (like tacos and tamales) or broad Hispanic stereotypes (like salsa music) for every group. However, the specific instructions in Prompt 6 proved that this is often an issue with how we ask the question, rather than a lack of knowledge in the model. When we asked for specific identities, the best models moved past simple stereotypes to include distinct local details—such as the Turrón de Doña Pepa for Peru, Fiambre for Guatemala, or the concept of Fiao in Dominican stores.

The models’ outputs (see examples in Table 5) also presented their explanations using very different styles. DeepSeek and Claude’s responses stood out by including visual aids like diagrams and code snippets, making these outputs very effective for technical learning. In contrast, the outputs from Gemini and ChatGPT focused on rich storytelling to explain concepts, while Mistral’s responses often repeated the same templates in the analogies given. Ultimately, the data suggests that while these models’ outputs were biased towards a “Mexican default,” these outputs could also include accurate, specific content for diverse Hispanic groups if the user applied more precise, persona-based prompts.

Phase 2		
Prompt #4 (Linked List)		
Puerto Rico		
LLMs:	ChatGPT	MetalAI
Examples:	Alright, imagine we're at La Fiesta de la Calle San Sebastián, and you're helping organize a group of friends who want to enjoy the festival together. You know how crowded it gets, so sticking together is key. Now, let's say each of your friends is holding hands in a line—this is like a linked list. Each friend in the group represents a node in the list. The hand they're holding is like a pointer, which connects them to the next person. If someone needs to leave (maybe they got too tired from all the dancing), they just let go of the hands, and the people on either side can connect directly. Now, if it were just an array list, imagine everyone is sitting on a bench with numbered seats. If someone leaves, you have to shift everyone down. With a linked list, you don't need to shift anyone—you just update the connections.	Explaining technical concepts to family and friends can be tough, but using something they're familiar with can make all the difference. So, primo, let me explain what a list, or linked list, is in a way that's relatable. Imagine we're at the Fiestas de la Cruz in Arecibo, and we're dancing to the rhythm of bomba y plena. Now, think of a conga line – you know, where each person is connected to the one in front of them, and we're all moving together. In a conga line, each person can be thought of as a “node” in the line. Just like in a linked list, each node (or person) has a connection to the next node in the sequence. When we move forward in the line, we're essentially “traversing” the linked list, one node at a time. Here's the cool part: just as we can add new people to the conga line or remove existing ones, a linked list allows us to insert or delete nodes at specific points in the sequence. This makes it super efficient for storing and managing data in a particular order.
Colombia		
LLMs:	ChatGPT	MistralAI
Examples:	Mami ¿Te acuerdas cuando vamos a la panadería en la mañana y hay una fila larga de gente esperando comprar pan de yuca, buñuelos y pandebonos? Pues, esa fila es como una lista en programación. Ahora, si hablamos de una lista enlazada (linked list), imagina que en la fila cada persona tiene un papelito con el nombre de la persona que está detrás de ellos. Cada persona sabe quién es el siguiente, y eso es lo que mantiene la fila conectada y organizada. Si queremos quitar a alguien del medio (por ejemplo, si una tía se da cuenta de que se le olvidó el dinero y se va), la persona que estaba antes simplemente actualiza su papelito para apuntar al siguiente, manteniendo la conexión. Así que una lista enlazada es como esa fila de la panadería, donde cada persona está conectada al siguiente, y podemos agregar o quitar elementos sin romper toda la estructura.	Abuela, you know how you make your famous ajiaco (a traditional Colombian soup)? Imagine each ingredient you add to the pot is like an item in a list. In a regular list, it's like having all the ingredients laid out on the table in a specific order. You can easily see and pick any ingredient you want, just like how you can access any item in a list directly. Now, a linked list is a bit different. Think of it like a chain of ingredients where each ingredient points to the next one. For example, you start with the chicken, and it points to the potatoes, then the potatoes point to the corn, and so on. You can't just jump to the corn; you have to follow the chain from the chicken to the potatoes to get to the corn. This way, each ingredient knows what comes next, just like each item in a linked list knows the next item in the sequence. It's a special way of organizing things that can be very useful in certain situations, like when you want to add or remove ingredients without messing up the whole recipe!

Table 5: Example outputs for Phase 2

6 Discussion

The responses from different LLMs and their approaches highlight several outcomes. In regard to RQ1, the prompt applied had a noticeable impact on the outputs generated. Our findings align with what scholars have described around prompt engineering for higher education, which suggests considering the value of including not just a verb and focus but also the context, condition, focus and condition, alignment, constraints, and limitations [39, 50]. The outputs had the highest cultural relevance across platforms when requesting responses, taking on a particular role or position. Rather than a generic “If you were...” prompt, adding specificity around being an instructor and developing culturally relevant examples was perceived as best suiting the intended goals. This lends support to the notion that iterative refinement, critical thinking, and content knowledge may be needed to guide responses and to convey particular types of information [51].

In terms of RQ2, examples centering on cuisine were the most commonly applied form of culture. This was followed by descriptions around family and social structures and then by events such as holidays and celebrations. The use of culinary examples is not entirely unprecedented, and researchers have previously mentioned that food can serve as an important way to communicate meaning and identity to others, in addition to finding solidarity [52]. Likewise, the focus on family and social structures fits with descriptions of the high value of family in Hispanic or Latine cultures [53]. These dynamics are shaped by a collectivist approach, which values putting the needs of the group before the individual, in direct contrast to the individualistic orientations within the U.S. [53, 54].

Yet, as indicated in our analysis, the responses given may continue to perpetuate stereotypes and/or neglect the nuances of individual backgrounds and lived experiences. Like many other social constructs, being Hispanic or Latine is multifaceted and influenced by factors such as nationality, immigration status, racial identification, and gender identification, making it hard to generalize what may be beneficial [55, 56]. In our outputs, frequent references made were centered on Mexican culture, which may not be reflective of the entire population of Hispanic or Latine students in the U.S. While seemingly problematic, authors have noted that U.S.-centered literature tends to focus on Mexican and Puerto Rican cultural tendencies “because of easier access for observation and research than other Latino groups” [53, p. 110]. Accordingly, it may be less an issue of bias from algorithms and more a result of the data available to generate examples.

Regarding RQ3, our nationality-specific analysis in Phase 2 revealed that specificity in prompting can mitigate the “Mexican default” bias observed in broader queries. When constrained to specific identities, models like ChatGPT and Meta AI successfully generated distinct regional markers—such as the *Turrón de Doña Pepa* for Peru—demonstrating that the models possess this cultural knowledge but require precise activation. However, this specificity varied by platform, with some models struggling to adapt their pedagogical metaphors beyond surface-level noun swapping.

As illustrated, while all LLMs and prompts may not be equal, such tools can offer opportunities to create culturally relevant resources. For educators looking to do so, we suggest enhancing specificity through careful prompting, such as adopting specific personas and referencing frameworks. It is vital to recognize that students are not a monolith and to avoid amplifying bias

by exploring individual nationalities or subgroups. Furthermore, resources should be verified for accuracy and pedagogical alignment; our findings indicate that models like DeepSeek and Claude prioritize visuals and code, while Gemini and ChatGPT favor narrative analogies. Finally, we recommend using these resources to promote critical thinking by asking students to critique the validity of AI-generated examples. This work serves as a starting point for using LLMs to enhance culturally relevant pedagogy across diverse populations. However, outputs must be scrutinized for potential harm, and we advocate for a mindful approach where instructors partner with students as co-constructors of knowledge to support learning.

As illustrated, while all LLMs and prompts may not be equal, such tools can offer opportunities to create culturally relevant resources to support additional resource development. For educators looking to do so, we suggest several guidelines and potential next steps:

1. *Consider prompts carefully to enhance the specificity of responses.* Using language like taking on the role of an instructor and directly referencing the intended framework can help to achieve more culturally appropriate responses.
2. *Students are not a monolith, and what may resonate with one individual may not be appropriate for all students.* It can be important to balance offering practical examples with not continuing to amplify bias. This may entail exploring individual nationalities or subgroups to assess how doing so may impact the cultural elements suggested.
3. *Resource creation may require multiple checks to ensure it achieves its intended goal.* It is important to confirm that the responses generated correctly describe the intended data structure and also that they make sense in the context of the intended population. Instructors could try to offer examples as resources in courses to explore how such examples may impact students' understanding and other outcomes. Alternatively, the outputs could be triangulated against those created by a human respondent, asking students or computing instructors how they would explain a particular data structure to a friend or family member. However, such an undertaking should consider that humans can also be fallible, and if they do not fully understand the concept or cannot explain it clearly, they could suggest an example that could perpetuate confusion further.
4. *Apply resources created not only as examples but also let them serve as opportunities to promote critical thinking and reflection among students.* Given that generative AI is increasingly pervasive, finding new ways to apply it and promote questioning can aid in students' development. One such approach could entail asking students to review the material generated and then explain why or why not a particular output may make sense. Doing so could require further research and actively engaging to develop their understanding.

Our work is intended as a starting point to encourage instructors to consider how LLMs could serve to aid in enhancing culturally relevant pedagogy. We also want to note that while our inquiry focused on a particular population, this could be extended to other groups to try to create a broader set of resources.

Yet, outputs should not be accepted unquestioningly, and it is important to be critical of what is produced and to recognize the inherent potential for harm. Through such measures, we hope to

encourage instructors to be mindful about their approach and to partner with students to draw upon their experience to become co-constructors of knowledge and aid in their efforts to understand data structures.

7 Limitations and Future Directions

We acknowledge that LLMs are rapidly evolving tools; the outputs analyzed here represent a snapshot of model capabilities in February 2025 and April 2025. Specific cultural biases or limitations observed in this study may shift as newer versions of these models are released. LLMs do not make meaning in the same way as a human, and generated text is based on syntactical rules and probabilities [28]. Accordingly, authorship of the content is not considered entirely “natural” since the outputs are predicated on prompts developed for the intention of analysis and source material applied, as well as its quality, is not transparent. In addition, contextual information like user preferences, browser histories, and the location where the queries were conducted could impact the outputs generated. The prompts and outputs here were highly specific to a particular context and cannot be generalized to speak about examples pertinent to all Hispanic or Latine students around the world. The tools applied were intentionally selected in recognition that instructors may be more apt to use public, commercial tools to generate course resources. However, in the future, this work could be expanded to explore stable, offline models.

In addition, it would be beneficial to validate the examples generated with Hispanic or Latine students, seeking their preferences and feedback on the examples through a survey, focus group, or interviews. While this was beyond the scope of the current analysis, gathering student and even faculty perspectives could yield valuable insight into how these examples are perceived, interpreted, and could be applied in educational settings. We aim to use this data to develop educational resources that leverage these tailored examples to improve engagement, belonging, and concept retention among Hispanic and Latine computing students. However, doing so requires validation and dissemination.

We also want to note that the outputs generated may not resonate with all Hispanic and Latine learners at every institution. While this approach is intended to provide students with materials that can aid in conceptualizing data structures, just because an individual is from a country does not mean the example is meaningful to their lived experiences. Moreover, we do not expect faculty to know students’ specific backgrounds nor to provide each to meet individual needs in larger courses. However, we do suggest these could be offered as resources for students to engage with and explore as they may try to internalize what these data structures may mean.

Going forward, an expansion to other cultural subgroups could further help to understand how well LLMs may generate culturally relevant content for a broader audience. It could also be important to focus on prompting for nationality-specific examples to understand how that may impact the stereotypes that may apply. Yet, this will also require further validation and assessment.

Ongoing evaluation of LLMs is important for ensuring quality and relevance. Given that many LLMs are opaque in how they generate responses, additional transparency around output generation and model behavior may be necessary. This also highlights the need for continued assessment of how effectively cultural content is produced.

8 Conclusions

We approached the creation of contextually pertinent examples through a comparison of different LLMs, prompts, and data structures. Analysis of the responses demonstrated that such tools can offer a helpful starting point and that they may incorporate multiple aspects of culture. Most often, references given for our population referred to culinary examples. While the responses skewed towards Mexico/Central America and did not fully represent the breadth of Hispanic or Latine students, they could offer a useful starting place to expand. Crucially, our findings also highlight that this skew can be mitigated through precise, nationality-specific prompting, which successfully yielded distinct regional markers. Although such examples may necessitate further validation, LLMs can provide a way to develop resources that may help more students to connect with such content. Going forward, we hope the findings from this investigation encourage faculty to continue to work towards seeking additional opportunities to develop pertinent content rather than applying a one-size-fits-all approach. However, we also want to remind others to be aware of the potential concerns and to make sure to question content generated with AI tools to mitigate potential harm.

References

- [1] “Computing curricula 2020: Curriculum guidelines for undergraduate degree programs in computing,” 2020, association for Computing Machinery.
- [2] G. L. McDowell, *Cracking the Coding Interview: 189 Programming Questions and Solutions*, 6th ed. CareerCup, 2015.
- [3] D. Zingaro, C. Taylor, L. Porter, M. Clancy, C. Lee, S. Nam Liao, and K. C. Webb, “Identifying student difficulties with basic data structures,” in *Proceedings of the 2018 ACM Conference on International Computing Education Research*, 2018, pp. 169–177.
- [4] K. S. K. K. Syed, “Conceptualization of data structures course using visualization techniques-a case study,” *Journal of Engineering Education Transformations*, vol. 30, no. 4, 2017.
- [5] H. S. Narman, C. Berry, A. Canfield, L. Carpenter, J. Giese, N. Loftus, and I. Schrader, “Augmented reality for teaching data structures in computer science,” in *2020 IEEE Global Humanitarian Technology Conference (GHTC)*. IEEE, 2020, pp. 1–7.
- [6] C. Pu, “Integrating generative ai with data structures and algorithm analysis course homework,” in *2024 IEEE Frontiers in Education Conference (FIE)*. IEEE, 2024, pp. 1–9.
- [7] S. A. Smarr, N. S. Morrow, and J. E. Gilbert, “Hbcu students’ experiences with data structures and the need for equity pedagogies,” in *Proceedings of the 2024 on RESPECT Annual Conference*, 2024, pp. 190–195.
- [8] T. C. Howard, “Culturally responsive pedagogy,” *Transforming multicultural education policy and practice: Expanding educational opportunity*, pp. 137–163, 2021.
- [9] N. Raihan, M. L. Siddiq, J. C. Santos, and M. Zampieri, “Large language models in computer science education: A systematic literature review,” in *Proceedings of the 56th ACM Technical Symposium on Computer Science Education V. 1*, 2025, pp. 938–944.
- [10] A. Attard and A. Dingli, “Empowering educators: Leveraging large language models to streamline content creation in education,” in *ICERI2024 Proceedings*. IATED, 2024, pp. 1312–1321.

- [11] K. L. Heitner and M. Jennings, "Culturally responsive teaching knowledge and practices of online faculty," *Online Learning*, vol. 20, no. 4, pp. 54–78, 2016.
- [12] C. Mathis, A. R. Daane, B. Rodriguez, J. Hernandez, and T. Huynh, "How instructors can view knowledge to implement culturally relevant pedagogy," *Physical Review Physics Education Research*, vol. 19, no. 1, p. 010105, 2023.
- [13] D. C. Freitas and H. L. Cardoso, "A nightmare on llms street: On the importance of cultural awareness in text adaptation for lrls," in *Proceedings of the 2nd LUHME Workshop*, 2025, pp. 40–48.
- [14] G. A. Garcia and O. Okhidoi, "Culturally relevant practices that "serve" students at a hispanic serving institution," *Innovative Higher Education*, vol. 40, pp. 345–357, 2015.
- [15] G. A. Garcia, A.-M. Núñez, and V. A. Sansone, "Toward a multidimensional conceptual framework for understanding "servingness" in hispanic-serving institutions: A synthesis of the research," *Review of Educational Research*, vol. 89, no. 5, pp. 745–784, 2019.
- [16] G. Gay, "Preparing for culturally responsive teaching," *Journal of teacher education*, vol. 53, no. 2, pp. 106–116, 2002.
- [17] G. Gay, *Culturally responsive teaching: Theory, research, and practice*. teachers college press, 2018.
- [18] G. Ladson-Billings, *The dreamkeepers: Successful teaching for African-American students*. San Francisco: Jossey-Bass, 1994.
- [19] F. López, D. Rivas-Drake, E. Serrano, and G. Delcid, "The role of asset-based pedagogy in promoting belonging and ethnic-racial identity among latine students," *Educational Psychology Review*, vol. 37, no. 1, pp. 1–38, 2025.
- [20] D. Coddling, B. Alkhateeb, C. Mouza, and L. Pollock, "From professional development to pedagogy: Examining how computer science teachers conceptualize and apply culturally responsive pedagogy," *Journal of Technology and Teacher Education*, vol. 29, no. 4, pp. 497–532, 2021.
- [21] J. J. Ryoo and T. Blunt, "“show them the playbook that these companies are using”: Youth voices about why computer science education must center discussions of power, ethics, and culturally responsive computing," *ACM Trans. Comput. Educ.*, vol. 24, no. 3, Aug. 2024. [Online]. Available: <https://doi.org/10.1145/3660645>
- [22] L. Malafouris, *How things shape the mind*. MIT press, 2013.
- [23] A. K. Dey, "Understanding and using context," *Personal and ubiquitous computing*, vol. 5, pp. 4–7, 2001.
- [24] J. Piaget, *The language and thought of the child*. Harcourt, Brace, 1926.
- [25] J. Piaget, *The moral judgment of the child*. London: Routledge & Kegan Paul, 1932.
- [26] J. Piaget and M. Cook, *The origins of intelligence in children*. New York: International University Press, 1952.
- [27] B. J. Wadsworth, *Piaget's theory of cognitive and affective development: Foundations of constructivism*, 5th ed. Longman Publishing, 1996.
- [28] M. Gillings, T. Kohn, and G. Mautner, "The rise of large language models: challenges for critical discourse studies," *Critical Discourse Studies*, pp. 1–17, 2024.
- [29] D. Huang, J. M. Zhang, Q. Bu, X. Xie, J. Chen, and H. Cui, "Bias testing and mitigation in llm-based code generation," *ACM Trans. Softw. Eng. Methodol.*, Mar. 2025, just Accepted. [Online]. Available: <https://doi.org/10.1145/3724117>

- [30] A. Jungherr, "Artificial intelligence and democracy: A conceptual framework," *Social media+ society*, vol. 9, no. 3, p. 20563051231186353, 2023.
- [31] N. R. Sahoo, A. Saxena, K. Maharaj, A. A. Ahmad, A. Mishra, and P. Bhattacharyya, "Addressing bias and hallucination in large language models," in *Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024): Tutorial Summaries*, 2024, pp. 73–79.
- [32] H. Park and D. Ahn, "The promise and peril of chatgpt in higher education: opportunities, challenges, and design implications," in *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, 2024, pp. 1–21.
- [33] M. Giannakos, R. Azevedo, P. Brusilovsky, M. Cukurova, Y. Dimitriadis, D. Hernandez-Leo, S. Järvelä, M. Mavrikis, and B. Rienties, "The promise and challenges of generative ai in education," *Behaviour & Information Technology*, pp. 1–27, 2024.
- [34] B. Paltridge, *Discourse analysis*. Springer, 2021.
- [35] B. Johnstone and J. Andrus, *Discourse analysis*. John Wiley & Sons, 2024.
- [36] Q. Jiang, Z. Gao, and G. E. Karniadakis, "Deepseek vs. chatgpt vs. claude: A comparative study for scientific computing and scientific machine learning tasks," *Theoretical and Applied Mechanics Letters*, p. 100583, 2025.
- [37] Q. Lang, M. Wang, M. Yin, S. Liang, and W. Song, "Transforming education with generative ai (gai): Key insights and future prospects," *IEEE Transactions on Learning Technologies*, 2025.
- [38] R. Manvi, S. Khanna, M. Burke, D. B. Lobell, and S. Ermon, "Large language models are geographically biased," in *International Conference on Machine Learning*. PMLR, 2024, pp. 34 654–34 669.
- [39] N. Knoth, A. Tolzin, A. Janson, and J. M. Leimeister, "Ai literacy and its implications for prompt engineering strategies," *Computers and Education: Artificial Intelligence*, vol. 6, p. 100225, 2024. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S2666920X24000262>
- [40] T. Huber-Warring and D. F. Warring, "Assessing culturally responsible pedagogy in student work: Reflections, rubrics, and writing," *Journal of Thought*, vol. 40, no. 3, pp. 63–90, 2005.
- [41] E. A. Sagastume García, "El hambre guatemalteco," Centro de Estudios de las Culturas en Guatemala, Tech. Rep., 2023, universidad de San Carlos de Guatemala.
- [42] J. E. Arellano, Ed., *Acahualinca: Revista Nicaragüense de Cultura*. Managua, Nicaragua: Academia de Geografía e Historia de Nicaragua, nov 2016, vol. 2.
- [43] J. E. Estévez Portillo, "Efecto potencial de la implementación de maíz y frijol biofortificados en la nutrición de la comunidad el jicarito, san antonio de oriente, francisco morazán, honduras," Master's thesis, Escuela Agrícola Panamericana, Zamorano, Honduras, 2016.
- [44] M. C. Martínez Solís and L. Rayo Trujillo, "Colombia diversa," 2019.
- [45] M. Seña, A. Giraldo, and D. Arango, "Mujeres ganadoras de la categoría: Canción inédita en el festival de la leyenda vallenata (valledupar – colombia)," *Anagramas Rumbos y Sentidos de la Comunicación*, vol. 24, pp. 1–26, 01 2026.
- [46] T. Gretton, "Posada and the 'popular': Commodities and social constructs in mexico before the revolution," *Oxford Art Journal*, vol. 17, no. 2, pp. 32–47, 1994. [Online]. Available: <http://www.jstor.org/stable/1360573>
- [47] N. E. Cantú and Guadalupe Cultural Arts Center, *La Quinceañera: Towards an Ethnographic Analysis of a Life-cycle Ritual*. San Antonio, Tex.: Guadalupe Cultural Arts Center, 2000. [Online]. Available: <https://books.google.com/books?id=P4ANAAAYAAJ>

- [48] C. Rodríguez-Nieto and Y. Escobar-Ramírez, “Conexiones etnomatemáticas en la elaboración del sancocho de guandú y su comercialización en sibarco, colombia,” *Bolema: Boletim de Educação Matemática*, vol. 36, pp. 971–1002, 12 2022.
- [49] T. E. Arias Ortiz, “Entre santos y cocodrilos. acercamiento a dos festividades en tabasco y guatemala,” *Península*, vol. 4, no. 1, mar 2010. [Online]. Available: <https://www.revistas.unam.mx/index.php/peninsula/article/view/44389>
- [50] B. Eager and R. Brunton, “Prompting higher education towards ai-augmented teaching and learning practice,” *Journal of University Teaching and Learning Practice*, vol. 20, no. 5, pp. 1–19, 2023.
- [51] W. Cain, “Prompting change: Exploring prompt engineering in large language model ai and its potential to transform education,” *TechTrends*, vol. 68, no. 1, pp. 47–57, 2024.
- [52] N. Stajcic, “Understanding culture: Food as a means of communication,” *Hemispheres*, no. 28, p. 77, 2013.
- [53] D. E. Battle, *Communication Disorders in Multicultural Populations: Communication Disorders in Multicultural Populations*. Elsevier Health Sciences, 2011.
- [54] I. Arevalo, D. So, and M. McNaughton-Cassill, “The role of collectivism among latino american college students,” *Journal of Latinos and Education*, vol. 15, no. 1, pp. 3–11, 2016.
- [55] G. A. Garcia and M. G. Cuellar, “Advancing “intersectional servingness” in research, practice, and policy with hispanic-serving institutions,” *AERA Open*, vol. 9, p. 23328584221148421, 2023.
- [56] A.-M. Núñez, “Employing multilevel intersectionality in educational research: Latino identities, contexts, and college access,” *Educational Researcher*, vol. 43, no. 2, pp. 85–92, 2014.