

Students' Development of Mathematical Meanings While Participating
Vicariously in Conversations between Other Students in Instructional Videos

Joanne Lobato^{1, 3}, John Gruver², and Michael Foster³

- ¹ Department of Mathematics and Statistics, San Diego State University, San Diego, California, United States
- ² Department of Mathematical Sciences, Michigan Technological University, Houghton, Michigan, United States
- ³ Center for Research in Mathematics and Science Education, San Diego State University, San Diego, California, United States

Citation

Lobato, J., Gruver, J., & Foster, M. (2023). Students' development of mathematical meanings while participating vicariously in conversations between other students in instructional videos. *Journal of Mathematical Behavior*, 71. <https://doi.org/10.1016/j.jmathb.2023.101068>

Acknowledgment

The development of this article was supported by the National Science Foundation through Awards DRL-1416789 and 1907782. The views expressed do not necessarily reflect official positions of the Foundation.

Abstract

Although interest in using videos in educational settings has surged in recent years, researchers know little about what mathematical meanings students develop from watching these videos or how they do so. To contribute to this gap in the research, we examined how two students appropriated mathematical meanings from instructional videos. In contrast to typical instructional videos, which rely heavily on an expert's exposition, the videos in our study featured the unscripted conversation of two high school students as they engaged with novel mathematical problems. This allowed us to examine how other students watching the videos coordinated meanings expressed by both the video participants and each other, including meanings that were initially incorrect or incomplete. To analyze these data, we adopted a Bakhtinian-inspired lens, which allowed us to conceptualize meaning as emerging from the relationships among multiple voices. Additionally, the appropriation of meanings from the videos was not straightforward. Instead, we found evidence of the repetition (mimicry) of words and actions from the video participants, revision, resistance, and the invocation of previously-appropriated voices, before the students were able to make the meanings expressed in the videos their own.

Keywords: Instructional videos, appropriation, Bakhtin, parabolas

Students' Development of Mathematical Meanings While Participating Vicariously in Conversations between Other Students in Instructional Videos

1. Introduction

Interest in the use of online mathematics videos for student learning has surged in recent years. Mathematics educators have increasingly incorporated videos into their classes, as part of “flipped” classrooms, “blended” or “hybrid” instruction, individualized video-based instruction, and massive open online courses (MOOCs). For example, in flipped instruction, students are assigned a video to view at home before class, which typically consists of the expository presentation of the mathematical content and worked examples that would normally constitute a class lecture (de Araujo et al., 2017). As a result, class time that would be devoted to direct instruction is freed up for students to negotiate how to apply the knowledge from the videos through classroom activities (Lo et al., 2017). Research indicates that such an instructional reorganization can be fruitful. A recent meta-analysis of 37 experimental studies comparing flipped and traditional classrooms showed a statistically significant effect in favor of flipped classrooms on student achievement (Güler et al., 2023). However, this research has not investigated the ways in which the videos themselves may be contributing to these positive outcomes.

Weinberg and Thomas (2018) identified a gap in the research on the use of online mathematics videos for student learning, calling for research that transcends the current body of work on flipped classrooms. Specifically, there is a lack of investigation into how students interpret, make sense of, and develop mathematical meanings from the videos themselves. This lack could be linked to what Weinberg et al. (2018) call “an implicit empiricist epistemology,” which assumes that “students uniformly construct the meaning the instructor believes the video

to convey” (p. 1263). When research *has* examined students’ experiences with the mathematical video component of instruction, it has focused on the identification of students’ attitudes toward the videos and some usage patterns (Attard & Holmes, 2022; Muir, 2021; Muir & Geiger, 2016).

The primary aim of this paper is to investigate how students develop meaning from mathematics videos. We treat students’ mathematical meanings as relational and emerging from an interplay between multiple utterances and perspectives—an approach that is aligned with Bakhtin’s dialogical theory (which we elaborate in Section 3.3). Consequently, it would be too limiting to examine students’ meaning development from the dominant type of mathematics video, which focuses on the presentation of procedures without considering multiple perspectives or unpacking mathematical relationships (Klinger & Walter, 2022). Examining how students develop meaning exclusively from procedurally-oriented videos would not be reflective of the full range of meanings students need to grapple with as they engage with rich mathematics. Thus, in this study, we asked learners to view videos that focused on making sense of mathematical concepts and ideas.

When deciding which videos to use in this study we also considered the presentation style of the videos. Most videos use an expository approach, which essentially recreates a lecture experience for viewers (Bowers et al., 2012). However, there is ample evidence that in classroom settings, alternative pedagogical approaches are preferable to lecture (Freeman et al., 2014). Researchers and educators have advocated for approaches where students are actively engaged with the mathematical discourse of other students, meaning they are actively comparing and contrasting ways of reasoning and extending and critiquing the reasoning of others (Boston, 2012; Hufferd-Ackles et al., 2004; National Council of Teachers of Mathematics, 2014). This can help students reexamine their own claims, clarify their thinking, and build more

sophisticated ways of reasoning by incorporating the reasoning of others (Francisco, 2013; McCrone, 2005).

With these considerations in mind, we turned to an alternative type of mathematics videos that several researchers have begun to make, which feature students as they engage in conversations to resolve mathematical struggles and develop mathematical meanings (Kolikant & Broza, 2011; Lobato et al., 2019; Seethaler et al., 2020; Weinberg et al., 2022). These videos hold greater potential for inviting viewers to think critically about the ideas being discussed and engage in a process whereby they negotiate and construct mathematical meanings that go beyond the interpretation of a procedure. Such videos often contain students' initial misconceptions and partially correct explanations, which are subsequently explored through discussion and argumentation. Because viewers of these videos participate indirectly or *vicariously* in the conversations of the video participants, their engagement is referred to as *vicarious learning*. Consequently, we ground our study in the vicarious learning literature, which we turn to next.

2. Literature review

2.1 *Vicarious learning*

Bandura introduced the term vicarious learning in the 1960s, to capture the social phenomenon of children learning by observing the behavior of others (Bandura et al., 1961, 1963). About 40 years later, the Scottish Vicarious Learning Project (SVLP) introduced vicarious learning into educational research, in the new context of learning from video or audio recordings of other people learning (e.g., McKendree et al., 1998). The SVLP argued that the central role of education was not to teach students facts and procedures, but rather to teach a process of logical reasoning, which they called *derivation*, where learners make their assumptions explicit and participate in argumentation. Furthermore, the SVLP contended that the

derivation process thrives in dialogues among students and between teachers and students: “educational discourse is full of explicit derivations because there is a continual need to check that assumptions are indeed shared and that interpretations are aligned” (McKendree et al., 1998, p. 115). Thus, through technology, the SVLP believed that conversations containing a derivation could be captured and indirectly engaged with by the viewers. Subsequently, a field of interdisciplinary research on vicarious learning in STEM education emerged.

A major emphasis of the vicarious learning studies in the sciences has been on the comparative effectiveness of videos that feature conversations (e.g., between a tutor and a student) versus expository videos (Chi et al., 2017; Gholson & Craig, 2006; Muldner et al., 2014; Muller et al., 2007). For example, Muller (2008) randomly assigned undergraduate physics students to learn from one of two videos: (a) an expository video in which an instructor conveyed physics information and (b) a video of a conversation between a tutor and a student for the same topic. The undergraduates reported that the expository video was clear, concise, and easy to understand, while the video featuring a conversation was confusing. However, the latter group had significant pre-post learning gains (on a mechanics conceptual inventory), while the expository group did not.

Vicarious learning studies in mathematics have tended to explore the experience of learners during engagement with the videos, rather than framing the investigation in terms of a comparison with expository videos. For example, Lobato and Walker (2019) examined how *vicarious learners* (VLs)—the students engaging with the videos that feature conversations—orient toward the video participants. In contrast to previous conceptions of VLs as emotionally and cognitively detached spectators (e.g., McKendree et al., 1998), the VLs acted as if they were in a collaborative group with the video participants and made regular displays of emotion.

Weinberg et al. (2022) explored how viewing a video that features interactions between students, along with attempting the mathematics task discussed in the video, can lead VLs to experience an intellectual need for the targeted mathematics topic. Other research has shown that VLs can learn to decenter (distinguish between their own reasoning and the mathematical thinking of video participants; Walters, 2017), make mathematical connections to the ideas expressed in a video when scaffolded by a teacher (Kolikant & Broza, 2011), and improve their mathematical reasoning from pre- to post-interviews (Putnam, 2021).

Taken together, the research on vicarious learning in science and mathematics education strongly suggests that the processes involved in the indirect participation in conversations are beneficial to learners. However, research is needed to identify those particular processes and connect them to an underlying theory of learning. We position our study, with its focus on the particular meanings that VLs develop from engaging with unscripted videos, as an important step in the exploration of the processes that entail vicarious learning. Our hope is that this, and future work, leads to the development of a vicarious learning theory that captures the complexity and fruitfulness of vicarious learning experiences.

2.2 Types of videos and types of engagement with videos in vicarious learning studies

Videos used in vicarious learning studies vary along at least three dimensions. First, the nature of the conversation between the video participants can range from scripted (e.g., Kolikant & Broza, 2011) to unscripted (e.g., Chi et al., 2017), with a middle ground of being loosely scripted (i.e., main points are identified in advance but the rest of the conversation is improvised; Weinberg et al., 2022). The videos used in our research study are unscripted, to allow for the expression of meanings—misconceptions as well as mathematically correct conceptions—from the perspective of real learners rather than those anticipated by researchers. Second, the video

participants have typically been either a student-tutor pair (e.g., Muller et al., 2007) or a pair of students (e.g., Lobato et al., 2019). The videos used in our research study feature two high school students (see Figure 1), to increase the chances that VLs of the same age will personally identify with the video participants (Mayes, 2015). Finally, the length of the videos can vary. Most vicarious learning research has been conducted with participants viewing one short video (e.g., Muller, 2008). In contrast, the videos in our study were selected from a series of videos comprising a video unit that we developed on parabolas (available at www.mathtalk.org). We explored longer periods of engagement with videos, because the development of productive mathematical meanings can be complex and require more than 10-20 minutes.

Just as there has been variety in the types of videos used in vicarious learning studies, a variety of conditions under which vicarious learners engage with videos has been explored. For instance, Chi et al. (2008) found that collaborative pairs of VLs watching a videotape of tutor-student interactions were more effective than individuals. Gholson and Craig (2006), in a review of vicarious learning studies, concluded that creating explanations and posing deep questions is important, not just for the video participants, but for the VLs as well. Weinberg et al. (2022) explored the importance of the VLs working on the same or similar tasks as the video participants. Consequently, in our study the VLs participated in collaborative pairs and worked together on mathematics tasks during the research sessions.

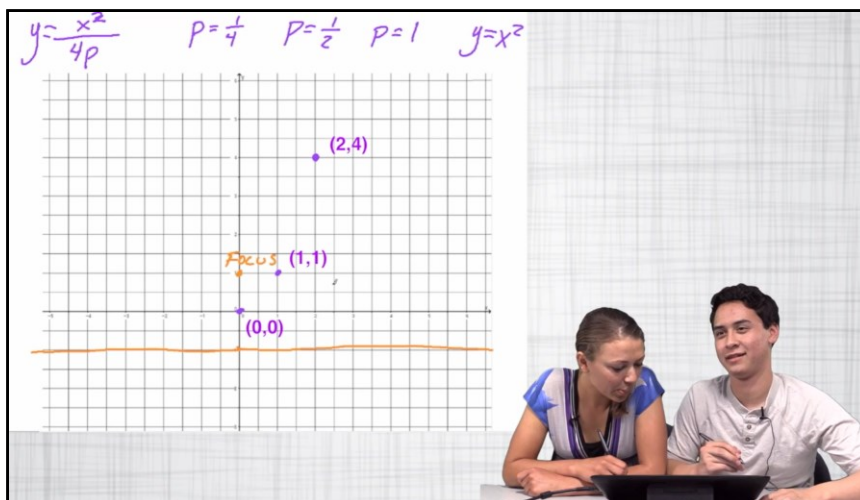


Figure 1. Screenshot from a video used in this study, which features the unscripted conversation of two high school students as they explore parabolas

3. Adopting a Bakhtinian-inspired lens

The purpose of this paper is to investigate how VLs develop meanings from mathematics videos that feature conversations between video participants. The VLs in our study had to manage both limited and productive mathematical meanings presented in the videos, while also coordinating ideas from each other and from their previous math experiences. Our interest in accounting for the complex coordination of discourse and ideas from multiple sources led us to the work of Mikhail Bakhtin, a Russian literary theorist and philosopher of language. Bakhtin (1986) adapted the term *polyphony* from music theory to emphasize the “multiplicity of independent and unmerged voices” (p. 208). In his theory of dialogism, Bakhtin (1981) posited that meaning only emerges through dialogic relationships among multiple voices. Furthermore from a Bakhtinian perspective, people cannot autonomously express themselves; rather the words and forms of expression that one uses come from others (Cooren & Sandler, 2014; Lemke, 2000). Bakhtin introduced the term *appropriation* within linguistics to account for process by which the words and discourse of others can be made one’s own. We were drawn to the

Bakhtinian construct of appropriation because of its emphasis on dialogic relationships, which is useful for our investigation of the connection between the meanings expressed by the video participants and the meanings that VLs develop. In Section 3.1 below, we elaborate the appropriation process as conceived by Bakhtin, followed by his constructs of voice and meaning development.

Despite the relevancy of Bakhtin's ideas to our research, he was not an educational researcher and did not study meaning development in the realm of mathematics. In applying Bakhtinian ideas from the domain of literary criticism to mathematics education research, adaptations and extensions are necessary to meet the particular needs and goals of our field. Consequently, we draw on both Bakhtin's translated works, as well as on the writing of educational researchers in general, and mathematics educators in particular, who have applied Bakhtinian constructs to various educational domains.

3.1 Appropriation

As learners coordinate a polyphony of meanings and discourses, Bakhtin claimed they may go through a process called appropriation. Introducing this term, Bakhtin (1981) wrote:

[L]anguage, for the individual consciousness, lies on the borderline between oneself and the other. The word in language is half someone else's. It becomes, 'one's own' only when the speaker populates it with his own intention, his own accent, when he appropriates the word, adapting it to his own semantic and expressive intention. Prior to this moment of appropriation, the word does not exist in a neutral and impersonal language (it is not, after all, out of a dictionary that the speaker gets his words!), but rather it exists in other people's mouths, in other people's contexts, serving other people's intentions: it is from there that one must take the word, and make it one's own. (p. 293)

For Bakhtin, appropriation is a process by which a word that one person uses (i.e., that one person "owns") is taken up and is used by another (i.e., that another person takes ownership of). Fundamentally, this process is about taking a word and filling it with one's own intentions,

where one's intentions can lead to the use of the word in a novel context or even a shift in the semantic meaning of the word in the expression of some new idea.

It is important to note that while Bakhtin's translators, Emerson and Holquist, opted for the use of "word" in this and subsequent quotes, Bakhtin (1981) had a broader meaning in mind than an individual utterance. While defining *discourse* in their Bakhtinian glossary, Emerson and Holquist acknowledge that "word" and "discourse" are interchangeable saying, "The Russian word *slovo* covers much more territory than its English equivalent, signifying both an individual word and a method of using words" (Bakhtin, 1981, p. 427). This suggests that the appropriation process is not only relevant to individual words but also to discourses. For simplicity's sake, we will continue to make use of "word," but keep in mind the broader Russian meaning.

Appropriation is not always a straightforward process and often involves resistance and negotiation. According to Bakhtin (1981):

And not all words for just anyone submit equally easily to this appropriation, to this seizure and transformation into private property: many words stubbornly resist, others remain alien, sound foreign in the mouth of the one who appropriated them and who now speaks them and who speaks them; they cannot be assimilated into his context and fall out of it; it is as if they put themselves in quotation marks against the will of the speaker. Language is not a neutral medium that passes freely and easily into the private property of the speaker's intentions; it is populated—overpopulated—with the intentions of others. Expropriating it, forcing it to submit to one's own intentions and accents, is a difficult and complicated process. (p. 293-294)

Before speakers are able to appropriate a word from others, they may repeat or mimic the word to try it on but feel like they are speaking a foreign language (Amhag & Jakobson, 2009). They may also resist using a word (e.g., Taylor, 2003) or enter into a struggle with the word as they attempt to submit it to their will. Because words are "overpopulated with the intentions" of other users, speakers negotiate with the word and its respective history. Bakhtin (1981) referred to this negotiation process as a "struggle" between the new word and those that are already influential

to the speaker (p. 346). This suggests that as we encounter and attempt to appropriate a new word, the word is negotiated with those words that we have previously appropriated. Our history, along with the history of the word being appropriated, shape our ability to make use of and fill a word with our own intentions. It is important to note that a negotiation process is also the central distinguishing feature between what Bakhtin called *internally persuasive* versus *authoritative* words (Morson, 2004). However, a proper exploration of these two constructs is beyond the scope of this paper; instead, we focus more narrowly on the negotiation process germane to appropriation.

Central to our analysis is the question of how our students appropriated mathematical meanings from the videos. We examine the VLs' appropriation of words, actions, and the associated meanings of those words/actions (which we connect to the Bakhtin's notion of *voice* in Section 3.2). Consequently, the appropriation of a voice involves the making of a voice one's own, that is, the speaker expresses the voice and adapts it to "his own semantic and expressive intention" (Bakhtin, 1981, p. 293). Bakhtin did not specify the particular ways in which one fills a word or voice with one's own intention. In mathematics education research, Radford (2000) articulated one way, namely expressing a word from someone else but transforming its associated meaning. For example, in his study of a Grade 8 algebra class, a teacher informed the students that the word "rank" can be used to express the number of circles corresponding to a figure number. Students used the word "rank," but supplied an additional meaning to the teacher's word relative to a vertical visual pattern that was salient to them.

In this study, we will show how appropriation can also happen through the integration of a new voice with voices that have already been appropriated (which we refer to as *previously-appropriated voices*). The polyphony of voices that occurs when a new voice comes in contact

with previously-appropriated voices results in a web of meanings. Indeed, it is through this connection with previously-appropriated voices that struggle occurs and negotiation takes place (Bakhtin, 1981). In sum, a voice that has been transformed or integrated has been used for one's own intentions. It has become one's own; it has been appropriated.

3.2 *Voice*

Central to Bakhtin's writings is the idea of voice. Two meanings for *voice* are present in both Bakhtin's writings and those of his interpreters. According to Hirschkop (2021), Bakhtin describes both one's "distinctive speech manner" and one's "distinctive ideological position" as "voice" (p. 89). The former is evident in Bakhtin's preeminent translators' (Emerson and Holquist's) definition of voice as, "the speaking personality, the speaking consciousness. A voice always has a will or desire behind it, its own timbre and overtones (Bakhtin, 1981, p. 434)." Voice, here, is about one's discursive representation. Indeed, we see this meaning cited by many of his interpreters (e.g., Silseth, 2012; Wertsch, 1991). This meaning of voice allowed Bakhtin (1981) to differentiate poetry from prose, critically analyze the work of Dostoyevsky (Belova et al., 2008), and more generally differentiate the novel from other genres.

Alternatively, voice as an ideological position is also expressed within Bakhtin's and his interpreter's work. Emerson and Holquist comment that *ideology* in Russian is simply an idea system and is not to be confused with the more politically-oriented term in English (Bakhtin, 1981). This suggests that ideas, perspectives, meanings are part of one's voice (i.e., what one is speaking about). Indeed, in educational research, this association of ideas or meanings with voice has gained traction. For example, Kolikant and Pollack (2015) characterize voices as "ideas, viewpoints, knowledge, beliefs, concerns, etc." (p. 327), and Linell (2009) defines voice as "an expressed opinion, view or perspective" (p. 116). Bakhtin also associated voice with meaning,

e.g., while discussing the refraction of meaning present within multiple voices, Bakhtin (1981) stated, “[i]n such discourse there are two voices, two meanings and two expressions (p. 324).” Our focus will be on voices as words or actions and their associated meaning. However, we acknowledge that Bakhtin (1981) claimed that voice as speaking personality is nonetheless present at the level of words and utterances: “The speaking person in the novel is always, to one degree or another, an *ideologue*, and his words are always *ideologemes*” (p. 333). Thus, a lens into one’s ideological position will inevitably intersect with their discursive representation.

3.3 *Meaning development*

From a Bakhtinian perspective, meaning only emerges when two or more voices are brought together in such a way that allows them to illuminate each other (Bakhtin, 1981; Koschmann, 1999). Meaning is not derived from dictionary definitions or the rules of grammar; rather, it emerges through the relationships among voices (Barwell, 2018; Silseth, 2012). Indeed, Bakhtin (1986) said, “understanding itself is ... the layering of meaning upon meaning, voice upon voice” (p. 121). Kazak et al. (2015) argue that existing theories of conceptual development derived from both Piagetian and Vygotskian perspectives fail to take into account the important role of relationships among two or more voices in enabling moments of insight and meaning. Furthermore, meanings depend on the context and discourse in which verbalization or writing occurs (Amhag & Jakobson, 2009). As Bakhtin (1981) put it, “Each word tastes of the context and contexts in which it has lived its socially charged life; all words and forms are populated by intention” (p. 293).

4. **Research question**

We now restate the purpose of our study, by grounding it in Bakhtin’s dialogism and using his construct of appropriation. Specifically, our research question for this study is: *How do*

vicarious learners appropriate meanings from mathematics videos that feature unscripted student conversations? Bakhtin's conceptualization of appropriation differs from the variety of neo-Vygotskian formulations of appropriation, which include, among others: (a) the uptake and use of cultural tools (e.g., Carlsen, 2010); (b) changes in individuals because of their involvement in mediated activities (e.g., Moschkovich, 2004), and (c) changing patterns of cultural participation (e.g., Rogoff, 1995). Instead, a Bakhtinian perspective emphasizes dialogical relationships and a process that is often marked by struggle, resistance, and negotiation. We responded to this research question by applying codes developed from Bakhtinian constructs (which we describe in Section 5.4) to analyze a pair of VL's engagement with unscripted videos as they came to appropriate voices (i.e., words or actions and their associated meanings) from the videos. Responding to this research question addresses a shortcoming in the research—a lack of investigation into the meanings that students develop by engaging with mathematics videos.

5. Methods

5.1 Participants

We studied two VLs' appropriation of mathematical meanings from instructional videos that feature the unscripted mathematical conversations of two video participants. The VLs were paired because we see their communication with each other as important to meaning development, and because previous vicarious learning research found that students who view videos in pairs learn more (Chi et al., 2008). Thus, we wanted to account for the VLs' negotiation of mathematical meanings in response to those presented in the video. We focused on one pair of VLs because the analysis of meaning development and the VLs' management of multiple voices requires complex (and time-consuming) qualitative analysis.

The two VLs, Desiree and Belinda, were recruited from a previous study in which 26 students interacted with just one of our video-based lessons (Lobato et al., 2019). During that study, Desiree and Belinda were able to make mathematical progress in the lesson. As such, they were well positioned to continue with the remaining video lessons in the unit, which is the focus of the current study. Desiree and Belinda were in Grade 9 and were earning grades in the B-to-D range in their regular Algebra 1 class. The VLs attended an ethnically diverse, low-income high school in a southwestern city in the United States. Both were fluent in Spanish and English.

5.2 *Data collection*

The VLs participated in 9 research sessions, each lasting 75-90 minutes. The research sessions occurred in a classroom at the VLs' school, after school hours. During the research sessions, Belinda and Desiree viewed instructional videos and worked together on mathematics tasks. They had access to the videos on a laptop before, during, and after working on mathematics tasks.

Two researchers also participated in the research sessions. One researcher operated two camcorders. One camcorder captured the VLs' interactions with each other; the other captured the VLs' inscriptions as they worked on the mathematics tasks. The VLs' written artifacts were also collected at the end of each research session. The other researcher interacted with the VLs. To help ensure that the videos served as the primary source of instruction for the VLs, the researcher's actions were limited. She sat across the room while the VLs watched videos and worked on mathematics tasks. After the VLs had completed viewing a video, the researcher sat across the table from the VLs and asked what they noticed in the video. After the VLs completed work on a task, they explained their thinking to the researcher. Because the researcher's actions were purposely constrained, she left many areas of confusion for the VLs unresolved.

5.3 Mathematics videos and paired tasks

The videos the VLs viewed feature the conversations of two Grade 9-10 students named Sasha and Keoni (as shown in Figure 1). In the videos, Sasha and Keoni engaged with mathematics problems they found challenging. During the problem-solving process, the video participants built upon each other's ideas and negotiated meanings. Because the problems were challenging, Sasha and Keoni's initial ideas were often mathematically incomplete or incorrect. However, they actively worked together to resolve struggles.

The videos show a simultaneous view of the video participants and their work (see Figure 1). A teacher, who remained off-screen, in order to keep the focus on the students, can be heard guiding the problem solving of the video participants. The videos also contain two forms of post-production: voice-overs and annotations. Voice-overs appear at the beginning to introduce each video and at the end of each video to summarize what the video participants have just done. Annotations highlight specific features of the video participants' work (e.g., labels for the focus and directrix of a parabola, as shown in Figure 1). The videos shown to the VLs in this study are part of a video unit on parabolas (available at www.mathtalk.org). The overarching goal of the video unit is to derive and develop meaning for the vertex form of an equation for a parabola, $y = \frac{(x-h)^2}{4p} + k$, where (h, k) is the vertex and p is the distance between the vertex and the focus. The desired meaning making includes being able to interpret the parameters and variables in the equation as *quantities*. A quantity is one's conception of a measurable attribute of an object (Thompson, 2011). In this case, the goal was for the VLs to interpret the parameters and variables as distances.

In the previous study (from which the VLs were recruited), the VLs had developed meaning for the geometric definition of parabola, which is the set of points that are equal

distance from a point called the focus and a line called the directrix. This meaning development had occurred as the VLs created a parabola by measuring distances from a point and line using a ruler and string, in a similar manner to the video participants. In the current study, the VLs continued to follow the trajectory of the video participants as they watched the videos and completed mathematics tasks that mirrored those given to Sasha and Keoni. If the video participants' mathematical task was complex, then the VLs worked on the same task. Other times, paired tasks were used (e.g., numerical values were changed for some component of the task), to ensure a high level of problem solving for the VLs.

Over the course of the nine research sessions, the VLs were able to use the geometric definition of a parabola to develop equations for parabolas. They began by using the Pythagorean theorem to find the coordinates for specific points on particular parabolas. They then generalized their process, yielding equations for several parabolas that had the origin as their vertex but whose distance from the vertex to the focus varied. They then generalized this process, allowing them to write the general equation for a parabola whose vertex is at the origin, $y = \frac{x^2}{4p}$, where the parameter p represents the distance from the vertex of a parabola to its focus. Finally, they derived an equation for a parabola where the vertex, (h, k) , is not at the origin: $y = \frac{(x-h)^2}{4p} + k$.

5.4 *Data reduction and analysis*

Preliminary analysis of the research sessions started with the creation of descriptive accounts for all 12 hours of videotaped data (Miles & Huberman, 1994). To reduce the data to a tractable amount, we initially identified 9 candidate topics in which the VLs demonstrated the development of a mathematical meaning. From these candidates, we selected the development of the meaning of p as a distance because of its importance mathematically and its complexity for the VLs. From the introduction of p to the VLs' development of the meaning of p as a distance

was 2.5 hours (Sessions 5 and 6). While Bakhtin emphasized the importance of history in meaning development, we believe that we can sensibly put boundaries on Sessions 5 and 6, because *p* was not introduced until Session 5 and the meanings for *p* that emerged in the videos did so in Sessions 5 and 6.

Within these two sessions, we identified key episodes in which some meaning for *p* could be inferred from the VLs' talk, gestures, or written inscriptions. We then analyzed those episodes by applying codes that we initially developed from the Bakhtinian constructs presented in Section 3. When these codes were insufficient, we created new codes induced from our data (following Miles & Huberman, 1994). When coding, two principles guided our inferences. First, we applied the principle of *fit* from Glaser and Strauss (1967). This means that inferences are consistent with the data and do not distort the data. It does not imply that our inference is the only one that fits the data; rather the inference is plausible and can be supported with evidence from the data. Second, we kept codes semantically close to what they represent; that is, we used a low level of inference rather than a high level (Miles & Huberman, 1994).

We first coded and named *voices*. As presented in Section 3.2, two meanings for voice are present in the Bakhtinian literature. For the purposes of our analysis, we focused on voice as words or actions and their associated meanings. However, because of the relationship between this meaning for voice and voice as speaking personality, we note that the meanings associated with words/actions are a *speaker's* meanings (more accurately, our inference of a speaker's meanings). Consequently, we named the speaker when coding voices (whenever we knew the origin of the voice). We inferred the speaker's meaning according to the two principles that we described previously. We identified several voices present in the videos, as well as a number of

the voices that the VLs seemed to bring from their previous experiences (e.g., previous research sessions or their regular Algebra 1 class).

We then coded the episodes using the following five codes: (a) repetition, (b) resistance, (c) revision, (d) appropriation, and (e) invoking a previously-appropriated voice. Four of these codes were inspired by the Bakhtinian literature (repetition, resistance, appropriation, and invoking a previously-appropriated voice). The other code (revision) was induced from our data. We briefly describe how we used ideas from the literature to arrive at these codes. The code definitions, and the type of evidence we looked for when applying the codes, are summarized in Table 1.

The codes of *repetition* and *resistance* are inspired by Bakhtin's (1981) passage about how difficult appropriation can be, namely that "many words stubbornly resist, others remain alien, sound foreign in the mouth [of the speaker]...it is as if they put themselves in quotation marks against the will of the speaker" (p. 293). We interpret this passage, in the context of our study, to mean that VLs might reproduce a word from the video participants or even recite a meaning they hear from the video (such as "*p* is a distance") but be mimicking the voice without providing evidence that they have made it their own. This is what we take to be *repetition*. We define *resistance* to entail rejecting a voice or simply setting a voice aside (for a time). For example, Taylor (2003) illustrated one way that resistance can be coded from data collected in a mathematics class for preservice elementary teachers. One student's sarcastic demeanor and facial expressions indicated resistance to the practice and meaning of "proving." The student spoke the teacher's word of "proving" but the "word sounded so foreign in her [the student's] mouth, she appeared to roll it around and spit it back at me" (Taylor, 2003, p. 341). This suggests

an initial rejection of proving. In the data for our study, resistance appeared in the form of finding a voice from a video too confusing and temporarily setting it aside.

The code of *revision* emerged inductively from our data to capture when the VLs altered actions from the video participants (e.g., revising their location of a parabola's focus) or used somewhat different words than the video participants but without evidence from which we could infer their associated meaning. As a result, for the code of revision, there is insufficient evidence of the VLs making a voice their own.

In contrast to repetition, revision and resistance, *appropriation* occurs when a voice becomes one's own, i.e., "when the speaker populates it with his own intention, his own accent...adapting it to his own semantic and expressive intention" (Bakhtin, 1981, p. 293). However, Bakhtin did not specify the particular ways in which one fills a word or voice with one's own intention (see Section 3.1). Consequently, we first looked for evidence of the VLs using words from the video participants but transforming their associated meanings (following Radford, 2000). However, in our data, appropriation instantiated itself through the process of the VLs' integrating a voice with other previously-appropriated voices. The specific words used by the VLs may be identical to those from the original utterance of the voice, or they could be slightly different (i.e., the words have been adapted to the VLs' expressive intention). Finally invoking a previously-appropriated voice means that one brings to bear a voice that one has already appropriated to reason about a new situation. These are voices the VLs appear to be comfortable with already. For example, the smooth execution of the actions associated with a voice is an indicator that a voice had been previously-appropriated.

Code	Definition	Nature of Evidence Sought
------	------------	---------------------------

voice	Words or actions and the speaker's associated meanings	A speaker expresses words or performs an action from which we can infer a meaning.
repetition	Reproduction of words or actions, or reproduction of meanings	The VLs re-utter a voice that has been expressed in a video but do not provide evidence of making it their own.
resistance	Rejection or setting aside of a voice	An important voice is presented within the video, but the VLs do not substantively engage with it at that time. The VLs verbally express rejection of the voice or confusion.
revision	Alteration to the words or actions of a voice but apparently not a revision to the associated meaning	The VLs use somewhat different words or actions than those from the original voice, but do not provide enough evidence from which we can infer the associated meaning. This could be the start of appropriation, but the VLs have not yet verbalized the associated meaning or provided sufficient evidence of making the voice their own.
appropriation	Expressing a voice and adapting it to one's own expressive or semantic intention (i.e., they make the voice their own)	The VLs express the same words or actions from a voice but with a transformed meaning, or they express the same meaning but integrate the voice with other previously-appropriated voices. In the latter case, the specific words used by the VLs may be identical to those from the original utterance of the voice, or they could be slightly different (i.e., the words have been adapted to the VLs' expressive intention).
invoking a previously-appropriated voice	Bringing to bear a voice that has already been appropriated to reason about a new situation	The VLs express a voice, or smoothly execute the actions associated with a voice, from a previous experience (e.g., from a previous research session or a school mathematics experience).

Table 1. Analytic codes

6. Background: Relevant voices from the previous research sessions

The results (presented in Section 7) will provide evidence that the VLs appropriated the voice “ p is a distance” from the video during Session 6 of the research study. However, during Session 6, the VLs also drew upon several voices that they had previously appropriated. Of particular relevance are voices from Session 5, where p was first explored. Presenting an evidentiary case for the appropriation of all these voices is beyond the scope of this paper.

Consequently, we simply identify the voices here that will be instrumental in the VLs' eventual appropriation of the voice of p as a distance during Session 6.

In Session 5, the VLs derived the general equation of a parabola with vertex at the origin (i.e., $y = \frac{x^2}{4p}$) after generalizing a pattern when working with a set of parabolas, all with vertices at the origin but with different foci. Although we had intended for p to be conceived as the distance between the vertex and the focus, this is not the meaning the VLs developed in Session 5. Instead, they managed various voices representing different meanings for p , such as: (a) the number on the y -axis next to the focus in the graph of the parabola, (b) the number you multiply by 4 in the denominator of the equation for the parabola, and (c) the number you add to or subtract from y when deriving the equation for a parabola. In each case, p was a number rather than a distance for the VLs. In Session 5, the VLs also drew on two voices that had emerged in previous sessions: (a) a point on the parabola is equidistant from the parabola's focus and directrix (from the geometric definition of a parabola), and (b) the Pythagorean theorem relates the horizontal distance, the vertical distance, and the actual distance between the focus and a point. The VLs also spoke about mathematical ideas they had learned in school. One such voice that is particularly relevant for their work in Session 6 is that the graph of a parabola can be created by using point substitution with an equation.

7. Results

In what follows, we present two instances of the VLs' development of mathematical meanings by appropriating voices from the videos. Section 7.1 is devoted to the VLs' appropriation of the video participants' voice of *placing the focus in a location that looks right*. The use of this voice resulted in both the video participants and the VLs making incorrect placements of the focus of a parabola. Section 7.2 presents the VLs' appropriation of the voice

from the video of *p as a distance*—the meaning we hoped students would develop from engaging with our videos. We include both occurrences of appropriation, rather than focusing only on the one that resulted in the development of the desired mathematical meaning, for three reasons. First, we wish to honor the complexity and challenge that this mathematics presented for the VLs. Second, the nature of the videos used in this study is such that the conceptual struggles of the video participants are visible. Consequently, we want to understand how the VLs coordinated various voices from the videos, rather than just examine the appropriation of correct mathematical meanings. Finally, in the two instances of appropriation, the ways in which the VLs developed their mathematical meanings differed, illustrating different possible pathways to appropriation.

7.1 Appropriation of Keoni's voice of placing the focus in a location that looks right

In this section, we present evidence to support the claim that the VLs appropriated a voice from Keoni, namely *placing the focus in a location that looks right*. This section is split into six subsections. First, we briefly describe the paired mathematics tasks given to the video participants and the VLs. Second, we turn to the video episode in which Keoni expressed the voice. The remaining four subsections are devoted to data from the VLs, along with our analysis. These four subsections were coded as containing (a) repetition, (b) revision, (c) invoking previously-appropriated voices, and (d) appropriation. In each subsection, we first present an episode without interpretation. Then we present our analysis in a separate paragraph or two. We follow these conventions: (a) new voices are italicized, (b) the VLs' previously-appropriated voices are identified but not italicized, and (c) codes are bolded.

7.1.1 Mathematical task and initial work

The VLs and the video participants worked on paired tasks in which a p -value was given, and they were asked to label the focus, directrix, and p on the graph. They were also asked to state the equation and graph the parabola. The video participants were given $p = \frac{1}{4}$ and $p = \frac{1}{2}$. The VLs were given $p = 1.5$ and $p = 2.5$. If students used the meaning of p as the distance from the vertex to the focus, then the focus would be $(0, \frac{1}{4})$ for the video participants' first parabola and $(0, 1.5)$ for the VL's first parabola, because the vertex is given as $(0, 0)$.

Determining the focus proved challenging, as the data presented below will demonstrate. However, finding the equations and plotting points on the graph was more straightforward, which is where both the VLs and the video participants began. Specifically, the VLs substituted the given p -values into the general equation $y = \frac{x^2}{4p}$ to obtain the equations for the video participants' version of the task ($y = \frac{x^2}{1}$ and $y = \frac{x^2}{2}$), as well as their own ($y = \frac{x^2}{6}$ and $y = \frac{x^2}{10}$). The video participants used the same approach to arrive at the equation $y = x^2$, when $p = \frac{1}{4}$. The video participants also substituted values for x into their equation $y = x^2$, to obtain two points, $(1, 1)$ and $(2, 4)$, which they plotted on the graph (as shown in Figure 2). The VLs plotted additional points for the video participants' equation $y = x^2$. When the VLs resumed the video to check their work, they viewed a conversation between Sasha and Keoni about where to place the focus.

7.1.2 Keoni places the focus by “look” in the video

After the video participants had placed $(1, 1)$ and $(2, 4)$ using the equation, the teacher (who can be heard but not seen in the video) asked if they could also find a point using what they know about the geometry of parabolas. Keoni quickly plotted a point at $(0, 0)$, reasoning, “It’s equal distance from our focus and our directrix” (apparently recalling the geometric definition of

a parabola from a previous lesson). The teacher responded by asking, “Where is the focus?” Keoni paused, quietly laughed, and exclaimed, “Yikes!” Eventually, Keoni asked Sasha if they wanted their focus to be 1, and she agreed. Then Keoni (incorrectly) plotted the focus at $(0, 1)$ and began to draw the directrix at $y = -1$ (see Figure 2). The teacher asked why the focus was located at $(0, 1)$. Keoni shared his reason that “It’s just like a general place to put it.”

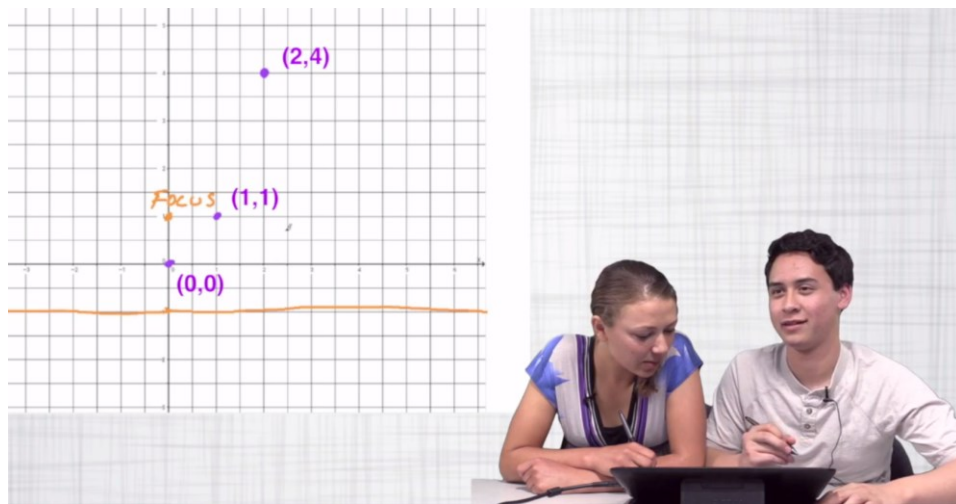


Figure 2. The video participants plot three points on the parabola, along with an incorrect placement of its focus

In this video, Keoni expressed the voice of what we are calling, *placing the focus in a location that looks right*. He engaged in the action of placing the focus at $(0, 1)$. From this action and his words, “It’s just like a general place to put it,” we inferred the associated meaning that a focus is a point that can be determined by intuition (i.e., it can be placed anywhere that looks right). We acknowledge that other reasonable inferences could be made. For example, perhaps his action of placing the focus above the x -axis suggests that, for Keoni, any point above the x -axis is a legitimate location for the focus. However, because Keoni did not mention spatial placement, we inferred a meaning using a lower level of inference. Regardless of the specific

criteria that makes a location look right to Keoni, it seems his meaning is disconnected from the meaning of p as the distance between the vertex and the focus.

7.1.3 Vicarious learners repeat video participants' placement of the focus

While the VLs watched the video, and just after Keoni expressed his reason for placing the focus at $(0,1)$ ["it's just like a general place to put it"], Belinda picked up a pen and drew a point at $(0,1)$ on the graph that she and Desiree had made of the video participants' parabola (Figure 3). Belinda then stopped the video before the VLs could see the video participants resolve their struggle.

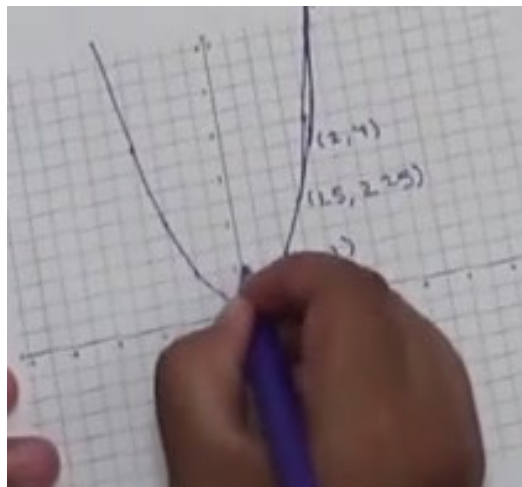


Figure 3. Belinda repeats Keoni's (incorrect) placement of the focus at $(0, 1)$.

We coded this brief episode as the **repetition** of Keoni's action of placing the focus at $(0, 1)$. From a Bakhtinian perspective, students may try on the words of others before fully appropriating them by adapting the words to their own intention. In a similar fashion, Belinda appeared to be trying on the action of placing the focus in the same location as the video participants.

7.1.4 Vicarious learners revise the focus after invoking previously-appropriated voices

Immediately after placing the focus at $(0, 1)$, Belinda drew a horizontal line segment from the focus to the point $(1, 1)$ on the parabola. Then she drew a vertical line segment from $(1, 1)$ to the x -axis (as shown in Figure 4) and exclaimed, “Whoa!”

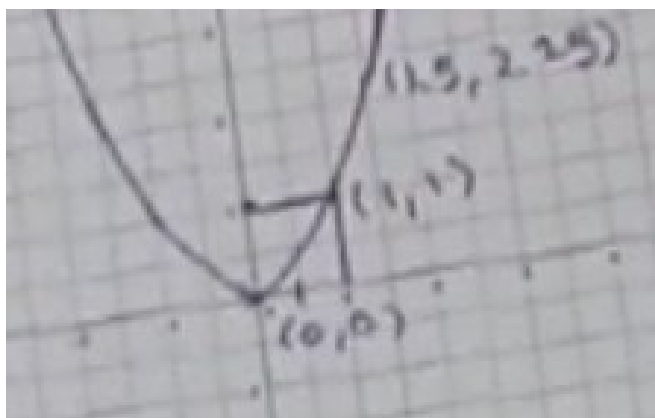


Figure 4. Belinda draws segments from a point on the parabola to the focus and the x -axis.

Desiree responded by arguing that the focus is too close to the point $(1, 1)$:

- Desiree: His [Keoni's] focus is too close to the point that he picked. The directrix is like two units away...
- Belinda: So... screw that [draws squiggles over the focus at $(0, 1)$ and her two line segments, apparently to erase them, as shown in Figure 5].
- Desiree: He [Keoni] put his directrix down here [sweeps pen horizontally from left to right over where the line $y = -1$ would be], and from here [places a finger at $(1, -1)$] to here [sweeps vertically up to $(1, 1)$] would be two.

Desiree then suggested they try the focus at 0.5 on the y -axis:

- Desiree: So I'm thinking the focus would be like right here, point five [plots a point at $(0, 0.5)$ as shown in Figure 5].

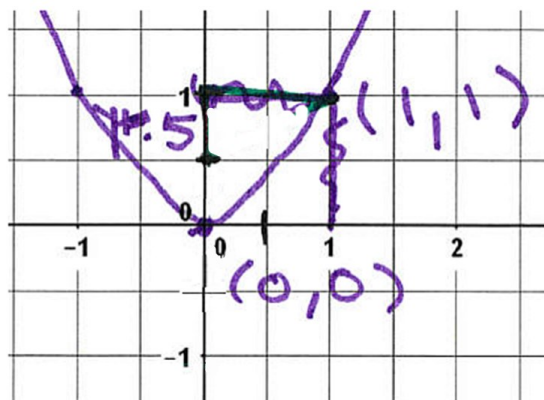


Figure 5. The VLs revise the placement of the focus to $(0, 0.5)$.

We coded this episode as containing **revision**, since the VLs' action of moving the focus to $(0, 0.5)$ seemed to be an alteration of Keoni's initial action of placing the focus at $(0, 1)$. This revised placement came after the VLs worked together to negotiate that $(0, 1)$ could not be the correct focus. Specifically, the VLs **invoked their previously-appropriated voice** of the geometric definition of a parabola, with the goal of testing whether or not $(0, 1)$ could work as the focus. Belinda drew two line segments of equal length. The first segment was from the focus to $(1, 1)$, a point on the parabola. The second segment started at $(1, 1)$ and dropped vertically down, mirroring their work from previous research sessions when they would determine the distance from a point on the parabola to the directrix. Desiree seemed to pick up on the fact that Belinda's vertical line segment was not long enough, indicating gesturally that Keoni put the directrix at $y = -1$, and verbalizing that the distance from $(1, 1)$ to the directrix should be 2. In other words, when the focus is at $(0, 1)$, then the distance from a point on the parabola is not the same distance to the focus as it is to the directrix, which means that the conditions of the geometric definition of a parabola are not met. Consequently, Desiree suggested moving the focus to 0.5, saying, "I'm thinking the focus would be like right here, point five."

In sum, Desiree and Belinda's use of Keoni's voice has undergone a revision through their changing of the location of the focus from $(0,1)$ to the point $(0, 0.5)$. However, there is insufficient evidence of the VLs' appropriation of Keoni's voice. While the VLs seem to be testing and negotiating with Keoni's focus of $(0, 1)$ by invoking a previously-appropriated voice, they have not verbalized the meaning associated with their revised placement of the focus. As a result, their semantic intention is unclear, and we cannot be certain appropriation has occurred from a Bakhtinian perspective.

7.1.5 Vicarious learners test the focus using previously-appropriated voices

The VLs immediately tested their focus at $(0, 0.5)$:

- Desiree: Then we can do the little triangle [draws lines segments from $(0, 0.5)$ to $(1, 1)$, $(0, 0.5)$ to $(0, 1)$, and $(0, 1)$ to $(1, 1)$, forming a right triangle, as shown in Figure 6].
- Belinda: But the directrix, where's the directrix?
- Desiree: It would be like right here [draws a mark at $(0, -0.5)$, as shown in Figure 6], wouldn't it?
- Belinda: Do you think it's the same distance? Here [places finger at $(1, 1)$] down [sweeps finger vertically down to the point $(1, -0.5)$] as here [places finger at $(1, 1)$] here [sweeps finger to the point $(0, 0.5)$]?]
- Desiree: I mean we can test it out.

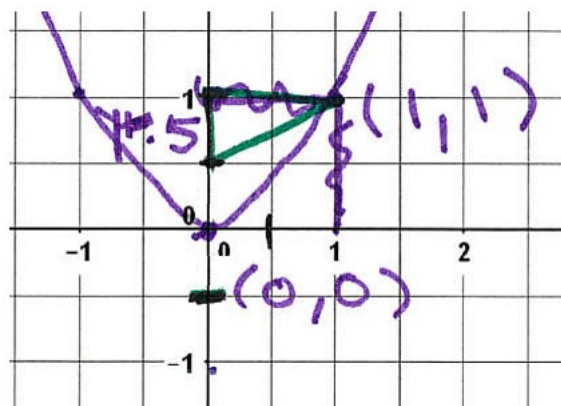


Figure 6. The VLs form a right triangle connecting the focus, a point on the parabola, and $(0, 1)$.

They used the Pythagorean theorem to determine that the distance from the focus to the point (1, 1) is approximately 1.1, as shown in Figure 7. With this information, they returned to their graph. Belinda argued that “this distance [places one finger at the focus and another finger at (1, 1)] is not the same as this distance [leaves one finger at (1, 1) but moves the other finger to their imagined directrix].” As a result, they seemed to reject (0, 0.5) as the focus. Desiree put her finger on the y -axis above (0, 1) and wondered, “The focus can’t be all the way up here, can it?” They decided to return to the video and “see what they [the video participants] do.”

Handwritten calculations in green ink:

$$\begin{aligned}
 a &= .5 \\
 b &= 1 \\
 c &= ? \\
 .5^2 + 1^2 &= c^2 \\
 .25 + 1 & \\
 \sqrt{1.25} &= c \\
 c &= 1.1180339
 \end{aligned}$$

Figure 7. The VLs’ calculations to compute the distance from the focus to (1, 1) using the Pythagorean theorem.

We analyzed this episode as another instance of **invoking previously-appropriated voices**. In this episode, the VLs worked together with the apparent goal of testing their revised focus of (0, 0.5). They drew upon the voice of the geometric definition of a parabola, which told them that the distance from the point (1, 1) to the focus needs to be the same as the distance from that point to the directrix. The VLs also appealed to the previously-appropriated voice of the Pythagorean theorem to help them identify the distance from (0, 0.5) to (1, 1). After (0, 0.5) failed their test, Desiree wondered if the focus could be “up here,” which suggests that she was considering another revision of Keoni’s focus.

7.1.6 Vicarious learners provide evidence of the appropriation of Keoni's voice

When the VLs resumed the video, they watched a segment in which the video participants corrected the placement of the focus from $(0, 1)$ to $(0, \frac{1}{4})$ and began to connect the focus with the given p -value [details about the video to be presented later in Section 7.2.1]. The VLs groaned, stopped the video and exclaimed, “We were way off!” Then they placed two red dots on the graph to indicate the focus and the directrix (as shown in Figure 8), and said they think that $\frac{1}{4}$ is correct. When the researcher asked what they noticed in the video, Desiree said it was “too confusing.” She recounted how she and Belinda had used 0.5 as the focus and directrix. When the researcher asked why they used 0.5, Desiree said that it felt right to her:

- Desiree: We plugged in point five because we were thinking it's probably point five of the focus and directrix and from this...
- Researcher: Ahh, why did you think the focus was point five away from the origin?
- Desiree: It felt right to me! [laughter]

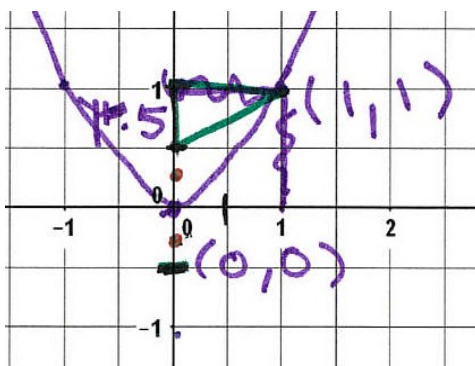


Figure 8. The VLs place a dot at $(0, \frac{1}{4})$ for the focus and a dot at $(0, -\frac{1}{4})$, which is where the directrix intersects the y -axis.

We coded this episode as evidence that the VLs had **appropriated** Keoni's voice of *placing the focus by look*. A voice has been appropriated when speakers express the voice and adapt it to their own semantic and expressive intention. One way appropriation can occur is when a speaker expresses the same meaning as another person's voice and integrates the voice with their own previously-appropriated voices. For the first time, one of the VLs provided a reason for

the placement of their focus at (0, 0.5). Desiree used different words from Keoni (“It felt right to me!” instead of “It’s a general place to put it”), suggesting an adaptation of the voice to her expressive intention. However, Desiree’s words seemed to carry the same meaning as Keoni’s. Furthermore, the VLs seemed to have imposed their own semantic intention on the voice, as evidenced by their two cycles of testing, negotiating with, and revising the focus. During this process the VLs appeared to integrate Keoni’s voice with their own previously-appropriated voices. Specifically, Keoni’s voice was expressed by Desiree in response to a query by the researcher about why the VLs had placed the focus at (0, 0.5). Their revised focus of (0, 0.5) had been the result of a negotiation process of testing and revising Keoni’s focus of (0, 1) that entailed invoking two previously-appropriated voices—the geometric definition of a parabola and the Pythagorean theorem. Consequently, the voice of *placing the focus by look* appears to be connected now with voices that the VLs had previously appropriated. However, we don’t know the precise moment of appropriation. Because the VLs didn’t verbally state a reason for their focus placements until later, we don’t know if appropriation occurred while they were testing and revising the focus, or if that happened later upon reflection.

A critical reader may wonder if the VL’s words and associated meaning for the focus could be an original construction rather than an instance of appropriation. However, for Bakhtin (1981), words are always “half someone else’s” (p. 293). As Cooren and Sandler (2014) put it, “A person cannot autonomously express himself or herself, because any form of expression that he or she uses (forms of language, genres of speaking and writing) comes from others” (p. 227). That said, the words and meanings one uses from others may have occurred historically rather than contemporaneously. We can never know for sure. However, the VLs were clearly engaged with the videos, as evidenced by: (a) Belinda picking up her pen to mimic the placement of the

focus by Keoni, (b) Desiree resuming the video after exchanges of negotiation with Belinda and then verbally responding to the ideas of the video participants, and (c) Desiree using similar phrasing as Keoni in her utterance of the voice. Consequently, our inference of the appropriation of Keoni's voice from the video is one plausible interpretation of the events.

7.1.7 Summary

The four episodes from the VLs, taken together, describe the process by which the VLs developed a meaning for the location of the focus of a parabola (albeit an incorrect mathematical meaning), as they appropriated Keoni's voice. The VLs initially watched Sasha and Keoni perform an action—placing the focus at $(0, 1)$ —and heard Keoni's justification that $(0, 1)$ is a “general place to put it.” They then repeated Keoni's action of placing the focus at $(0, 1)$ and began negotiating with his voice by testing his location against the previously-appropriated voice of the geometric definition of a parabola, revising the location, and testing again. Finally, the VLs provided evidence they had appropriated Keoni's voice of *placing the focus in a location that looks right*, with Desiree's justification of their placement of the focus, “It felt right to me!” Desiree's justification provides evidence that the VLs have developed the same meaning for the location of a focus of a parabola as expressed by Keoni—that the focus lies in a location that intuitively makes sense. However, the VLs indicated that they went beyond simply repeating Keoni's voice with two forms of evidence that they made his voice their own. First, while Desiree's words seem to carry the same meaning as Keoni's voice, she went beyond the reproduction of his words by adapting them to her expressive intention. Second, the VLs integrated Keoni's voice with their own previously-appropriated voices.

7.2 *Appropriation of the video participants' voice of p as a distance*

In this section, we present evidence to support the claim that the VLs developed meaning for the p -value of a parabola as a distance, through appropriation. This section is split into four subsections. First, we turn to the video episode in which Sasha and Keoni's voice of *p as a distance* was expressed. The VLs viewed this video twice—once as recounted previously in Section 7.1.6 and again right after the VLs expressed their reason for placing the focus by what looks right. The remaining subsections are devoted to data from the VLs, along with our analysis. These subsections were coded as containing: (a) resistance, (b) repetition, (c) invoking previously-appropriated voices, and (d) appropriation. We continue with the convention of italicizing new voices, identifying the VLs' previously-appropriated voices but not italicizing them, and used bolded font for the codes.

7.2.1 *Video participants express the voice of p as a distance*

As mentioned previously, the video participants were able to correct their placement of the focus from $(0, 1)$ to $(0, \frac{1}{4})$ and connect the focus with the given p -value. We provide details in this section regarding how this occurred. Right after Keoni had justified the placement of the focus at $(0, 1)$ as being a “general place to put it,” the teacher asked the video participants what the p -value would be if the focus were at $(0, 1)$, which set in motion a conversation that ended with Sasha and Keoni suggesting a revised focus at $(0, \frac{1}{4})$:

- Teacher: If the focus were there [referring to the point at $(0, 1)$], what would the p -value be?
- Keoni: One? One. [Annotation “ $p = 1$ ” with brackets indicating the distance between the origin and the focus, as shown in Figure 9, appears for future VLs; the video participants cannot see the annotation.]
- Teacher: And what p -value were we working with for this particular parabola?
- Participants: [In unison] One-fourth.
- Teacher: Can you adjust your focus then?
- Keoni: Yes... so like right here [points to $(0, \frac{1}{4})$], right? No? No.
- Sasha: Well our p is one-fourth, so yeah. I agree with you buddy!

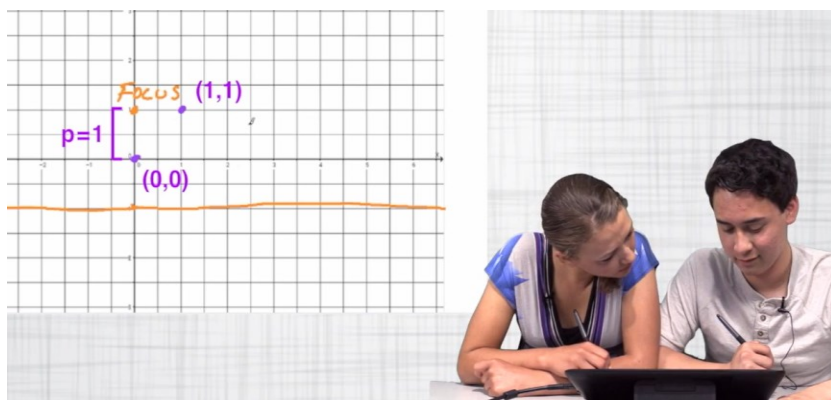


Figure 9. Annotation of “ $p = 1$ ” appears on the video to indicate the distance between the vertex and the focus.

Continuing, Sasha and Keoni plotted and labeled a new focus at $(0, \frac{1}{4})$. They drew a directrix by first placing a dot at $(0, -\frac{1}{4})$ and then drawing a horizontal line that passed through the dot. When the teacher asked how they knew where to place the directrix, the video participants responded with an explanation based on the distances from the origin to the focus and to the directrix:

Teacher:	Why is that?
Keoni:	Because our point has to be equal distance from our focus and our directrix, and if this is one of our points [points to origin] it has to be equal distance from both of them.

The video episode concluded with a voice-over describing what Sasha and Keoni discovered:

Sasha and Keoni figured out something important! When p equals one-fourth, then the focus is one-fourth of one unit from the vertex [annotation “ $\frac{1}{4}$ unit” appears with brackets indicating the distance from the vertex to the focus, as shown in Figure 10], and the directrix is one-fourth of one unit from the vertex [a second annotation of “ $\frac{1}{4}$ unit” appears indicating the distance from the vertex to the directrix].

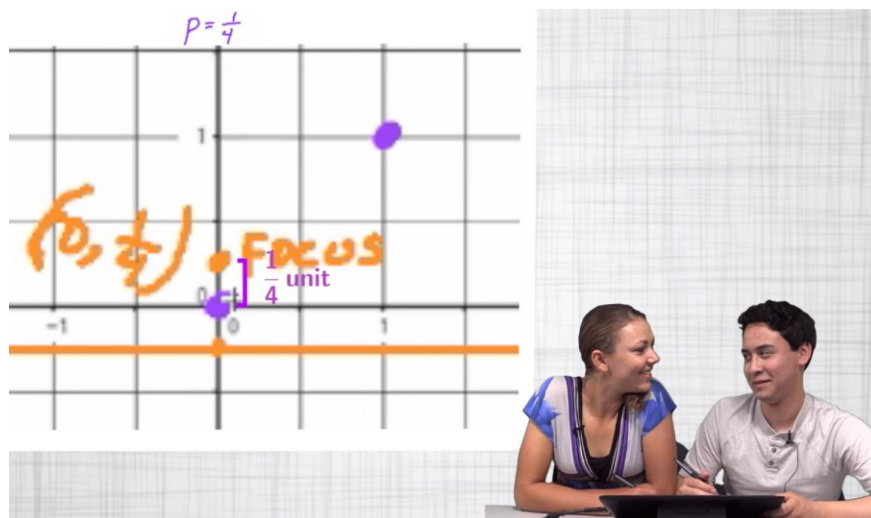


Figure 10. Annotation of “1/4 unit” appears on the video to indicate the distance between the vertex and the focus.

During this video episode, several voices emerged. First, Keoni expressed the voice *you know the location of the focus, you know the p-value*. He stated that a given focus of 1 would have a corresponding p -value of 1. These words were in response to the teacher’s probe about the respective p -value for a focus at $(0, 1)$, which suggests that his voice has an associated meaning that given any focus in the form of $(0, y)$, y would represent the p -value. Second, Sasha and Keoni expressed a related voice of *once you know the p-value, you know the location of the focus*. This voice is particularly relevant to the task at hand and emerged when Keoni pointed to $(0, \frac{1}{4})$ and suggested it could be the location of the focus. The associated meaning is then evidenced by Sasha’s supporting statement, “our p is one-fourth.” Third, Keoni leveraged a familiar voice when he justified the placement of the directrix at $y = -\frac{1}{4}$, the voice of the geometric definition of the parabola. The final voice that will become relevant for the VLs is Sasha and Keoni’s voice of *p as a distance*. When Keoni uttered, “If this is one of our points [gestures to the vertex at the origin] it has to be equal distance from both of them [referring to the focus and directrix],” he was evidencing an associated meaning of p as a distance. Specifically,

Keoni attended to the space between the vertex and the focus and between the vertex and the directrix. This attention to distance was reinforced by the annotations presented during the voice-over (Figure 10), which explicitly connected these distances to their measurement ($\frac{1}{4}$ of one unit) and to the given p -value of $\frac{1}{4}$.

7.2.2 Vicarious learners resist the voice of p as a distance and invoke previously-appropriated voices

After the VLs had watched the video twice (in which p as a distance was expressed), the researcher asked what they noticed about p , the focus, and the directrix. Desiree immediately responded, “It’s getting harder, intense.” The researcher asked if they knew where to put the focus, directrix and p -value for one of the parabolas. Instead of answering the researcher’s question, they shifted their attention to graphing a parabola using point substitution. Working with $y = \frac{x^2}{6}$ (which they had previously identified as the equation for the parabola when $p = 1.5$), the VLs substituted different x -values into the equation, calculated the associated y -values (Figure 11a), and organized the values in a table (Figure 11b). Then they plotted the points to create a parabola (Figure 12). This systematic process was a somewhat lengthy detour from the task of locating p , the focus and the directrix, taking about 6 minutes to complete.

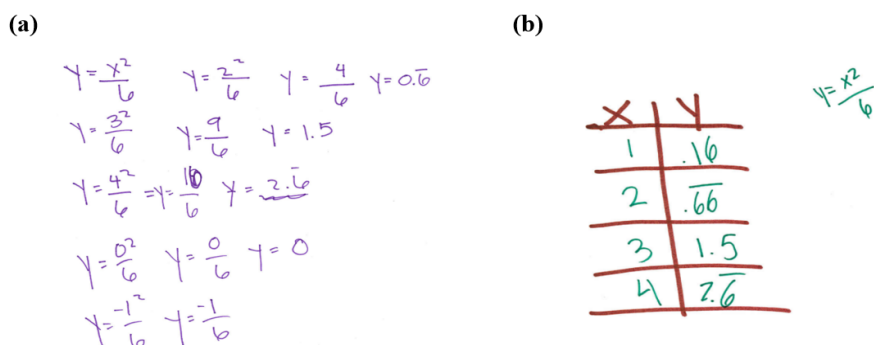


Figure 11. The VLs (a) calculate coordinate pairs using the equation $y = \frac{x^2}{6}$ and (b) record them in a table.

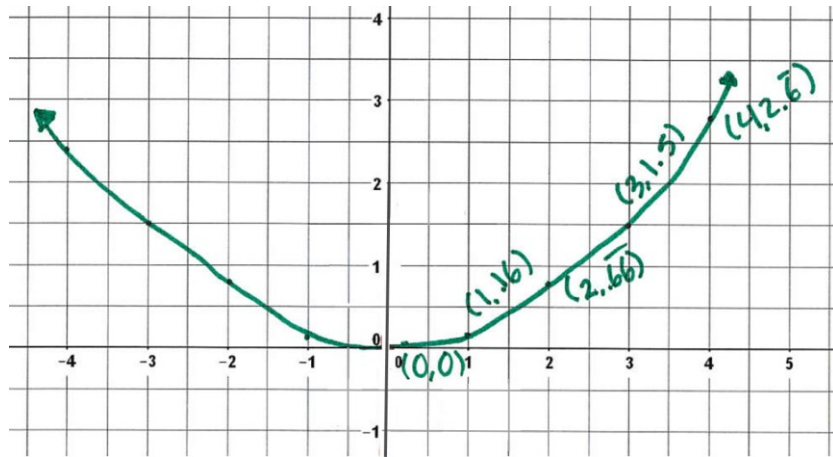


Figure 12. The VLs' graph of the parabola.

We coded the apparent setting aside of the video that the VLs had just viewed twice as a type of **resistance**. Despite being asked twice about p , the focus and the directrix, the VLs ignored these questions. Desiree's utterance about the video being "harder, intense" suggests that they found the new voices expressed in the video as too confusing or difficult to engage with at that moment. Instead, they **invoked a previously-appropriated voice**, namely the familiar voice of graphing a parabola using point substitution with an equation.

7.2.3 Vicarious learners repeat one of Sasha and Keoni's voices related to p

Having graphed the parabola using point substitution, the VLs returned to the original task statement and reread the prompt to find p , the focus, and the directrix:

- Desiree: [Reads the instructions aloud], "Where is p , the focus, and directrix?"
Wouldn't the directrix [misspeaking] be up here [places index finger from right hand at the point $(0, 1.5)$, pauses, and then places index finger from left hand at the point $(0, -1.5)$]?
Belinda: Oh. The focus?
Desiree: Yeah. This is the focus right here [plots a point at $(0, 1.5)$, as shown in Figure 13]. And directrix right here [draws line $y = -1.5$, as shown in Figure 13].
Belinda: Oh, because it's [p] one point five?
Desiree: Yep. So, focus [adds the label "focus" to the graph; Figure 13] and directrix [adds the label "directrix" to the graph; Figure 13].

When the researcher asked the VLs to tell her about their work, they started by describing their point-substitution method. The researcher acknowledged that they had graphed the parabola and then asked where p is on their graph:

- Researcher: Where is p on the graph?
 Belinda: This [sweeps finger from the origin down vertically to the directrix] and this [sweeps finger from the origin up vertically to the focus].
 Researcher: Can you label those [p -values]?
 Belinda: This is p [draws a line segment from the origin to the directrix and labels it “ p ”; see Figure 13], and this is p [draws a line segment from the origin to the focus and labels it “ p ”].
 Researcher: And what is the value of p for this parabola?
 Desiree: One point five.
 Belinda: [Adds “= 1.5” to the label of “ p ” for the segment from the origin to the directrix, as shown in Figure 13].

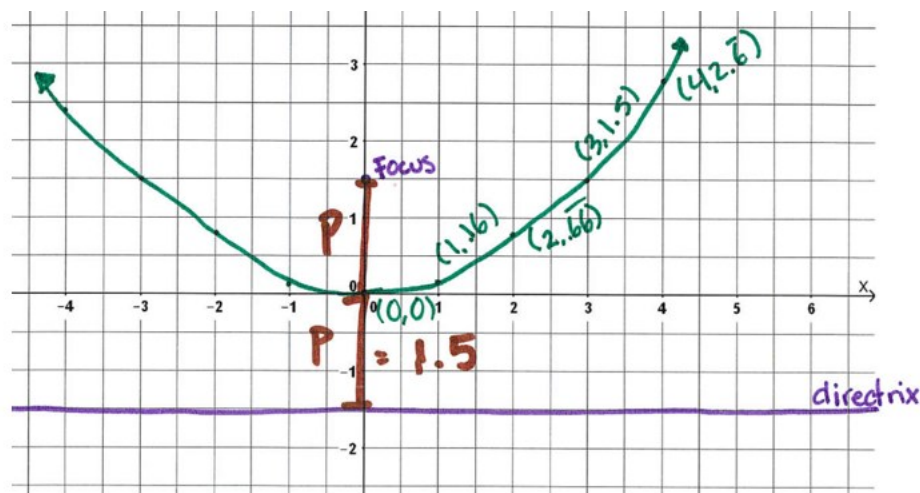


Figure 13. The VLs label p , the focus and the directrix.

Following the inclusion of p on their graph, the VLs were asked to reflect on connections between the p -value, focus, and directrix. In response, Belinda claimed, “When you’re given p , you know the focus and directrix.”

Within this episode, we coded Belinda and Desiree as having **repeated** one of the new voices related to p that was expressed in the video, namely *once you know the p -value, you know*

the location of the focus. Recall that in the video, Sasha had explicitly connected the p -value to the location when she stated, “Well our p is one fourth, so yeah. I agree with you buddy!” in response to Keoni’s suggested placement of the focus at $(0, \frac{1}{4})$. Similarly, Belinda asked if Desiree’s placement of the focus at $(0, 1.5)$ was “...because it’s [p] one point five?” Showing a repetition of both the words and the associated meaning of Sasha and Keoni’s voice. A bit later, Belinda expressed the voice explicitly, “When you’re given p , you know the focus and directrix.” This voice is significant, because it is very different from their previous voice of *placing the focus by look*. Applying their new voice allowed the VLs to use the given p -value. Once the VLs were able to use the given p -value, then the voice from the video that they had been setting aside— p as a distance—began to emerge.

Sasha and Keoni’s voice of p as a distance was initially indicated in Belinda’s sweeping gestures when asked where p is. Compared with Desiree’s deictic gestures for her proposed locations of the focus and directrix (i.e., pointing to $(0, 1.5)$ and $(0, -1.5)$), Belinda’s gestures are metaphoric (Solomon et al., 2018). These metaphoric gestures embody a view that p is the space between the focus or directrix and the vertex, which is suggestive of Sasha and Keoni’s voice of p as a distance. Furthermore, the voice of p as a distance is evidenced in the VLs’ inscriptions for p . Comparing the annotations from the video (Figures 9 and 10) to the VLs’ inscriptions (Figure 13) suggests that the VLs were seeing p in a similar way. Although the voice of p as a distance appeared to be in the process of being appropriated, the VLs were still struggling to articulate verbally that p is a distance.

7.2.4 Vicarious learners appropriate the voice of p as a distance

After completing the task where p is given as 1.5, the VLs went on to the task where $p = \frac{1}{2}$. In contrast to the previous episodes, Belinda immediately drew the directrix at $y = -0.5$,

created a line segment from the directrix to the vertex, and labeled the segment as p (as shown in Figure 14). She then drew a point at $(0, 0.5)$ and labeled the segment from focus to vertex as p (see Figure 14). Meanwhile, Desiree started using the substitution method with the equation $y = \frac{x^2}{2}$ to locate coordinate pairs, and then Belinda joined in. Together they created a table of x - and y -values, and Desiree plotted the points to create a graph of the parabola (Figure 14). After they were done, Belinda succinctly summarized their thinking to the researcher:

Belinda: Since we know that p is the distance between the focus [and vertex] and directrix [and vertex] we put point five [gestures to the “ p ” next to the line segment between the vertex and the focus], because one half equals point five [gestures to the “ p ” next to the line segment between the vertex and the directrix]. Then we solved for the points.

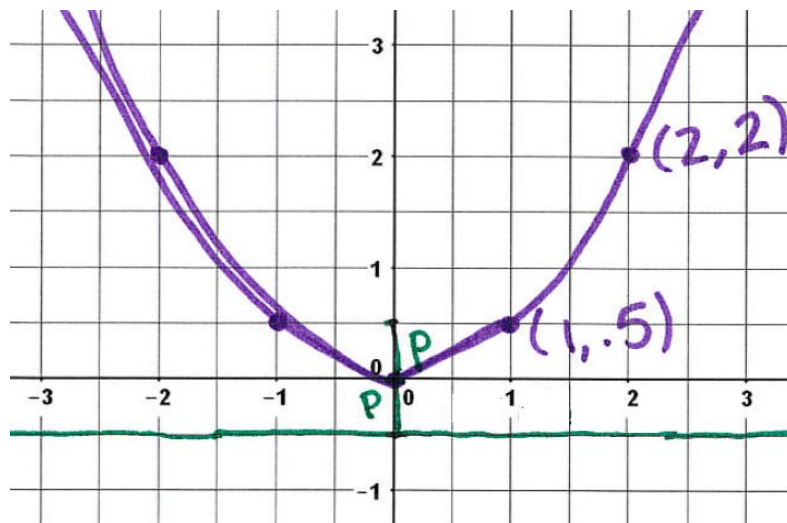


Figure 14. The VLs label the distances between the vertex and the focus, and between the vertex and the directrix, as p .

We coded this episode as evidence that the VLs had **appropriated** Sasha and Keoni’s voice of p as a distance. A voice has been appropriated when speakers express the voice and adapt it to their own semantic intention. One way appropriation can occur is when a speaker expresses the same words and meaning, and they integrate the voice with their previously-

appropriated voices. Belinda expressed the same words and meaning of Sasha and Keoni's voice with her inscriptions of two line segments on the graph labeled as " p " and by verbally articulating that " p is a distance." Integration of this voice with the VLs' previously-appropriated voice is evidenced by the temporal relationship between previously-appropriated voices and the emergence of Sasha and Keoni's voice. For the previous task with $p = 1.5$, the VLs relied on their previously-appropriated voices to graph the parabola using point substitution. Only at the end of their work, after being asked by the researcher where p , the focus and directrix were on the graph, did p begin to emerge as a distance, as evidenced non-verbally by the VLs' gestures and inscriptions. In contrast, for the task with $p = \frac{1}{2}$, Sasha and Keoni's voice of p as a distance had primacy. Specifically, the VLs immediately engaged with p , marking it on the graph and using it to locate the focus and directrix. However, the new voice of p as a distance did not displace the previously-appropriated voice of *graphing a parabola using point substitution*. Rather, these voices were integrated together as the VLs used both to engage in the task, suggesting the voice of p as a distance had been filled with the VLs' own semantic intentions.

7.2.5 Summary

The episodes presented in Section 7.2, taken together, reveal a process of meaning development through appropriation. Fittingly, the appropriation of p as a distance started with resistance. As Bakhtin (1981) states, "not all words for just anyone submit equally to this appropriation ... many words stubbornly resist (p. 293)." Initially when the VLs viewed a video twice in which p as a distance was expressed, they set it aside as too confusing and hard. Furthermore, they ignored repeated queries to place p on the graph. Instead, they turned to their previously-appropriated voices to graph the parabola by using point substitution and creating a table of values. Reflecting on their graph, the VLs were able to connect another voice from the

video, namely *once you know the p-value, you know the location of the focus*, which resulted in correct placements of the focus and directrix. Once these different voices were able to meet, a web of connections formed. What surfaced for the VLs was the emergence of Sasha and Keoni's voice of *p as a distance*—first expressed nonverbally through gestures and inscriptions and later expressed verbally and integrated with a set of their previously-appropriated voices.

8.0 Discussion

8.1 Contributions

This paper offers a response to a gap in mathematics education research—a need to investigate how students develop meanings by engaging with mathematics videos. The findings presented from this research make three contributions. First, we identified particular meanings that VLs developed and showed how those meanings could have developed through the appropriation of voices from the videos. Initially the VLs appropriated the video participant's (Keoni's) meaning that a focus of a parabola is a point that can be determined by intuition (i.e., it can be placed anywhere that looks right). Correspondingly, the VLs made two incorrect placements of the focus. Eventually, the VLs appropriated the more productive meaning from the video of *p as a distance*. The appropriation of these two voices occurred in different ways. Specifically, the VL's appropriation of Keoni's voice of *placing the focus by what looks right* began with the VLs immediately repeating Keoni's action; in contrast, the VL's appropriation of *p as a distance* began with resistance (setting aside the video for a time). Additionally, the first instance of appropriation included revision (e.g., the VLs revised Keoni's focus), whereas the second instance of appropriation preserved the words and associated meaning of *p as a distance* expressed in the video. However, a throughline across both instances was the VLs' invocation of

previously-appropriated voices and the integration of voices from the videos with those previously-appropriated voices.

Second, this paper points to the important role that previously-appropriated voices can play. The VLs' moment of insight, when *p as a distance* emerged through the VLs' gestures and inscriptions (and later verbally), came after a lengthy episode in which the VLs invoked previously-appropriated voices to create a table using substitution and to graph the parabola. Once the parabola was realized in graphical form, the VLs were able to use a voice from the video participants expressing a numeric connection between the given *p*-value and the location of the focus (i.e., for a *p*-value of $\frac{1}{4}$, look for $\frac{1}{4}$ on the *y*-axis, and place the focus there). Once the focus was placed, it created one endpoint, and the vertex from their graphing exercise created the other endpoint. This made it possible for the VLs to gesture over the space between the vertex and focus and make an inscription representing the distance between the two points. Soon after, one of the VLs verbally expressed *p* as the distance between vertex and focus. In other words, the VLs were able to make meaning of the new idea of *p* as a distance by hearing the voice in relation to their previously-appropriated voices.

Finally, this paper offers one account of how VLs can engage with and coordinate the misconceptions expressed by video participants in unscripted conversations. Previous research indicated that the inclusion of common misconceptions in videos that feature conversations seems to be productive for VLs, as evidenced by results in improvements in learning outcomes (Boesdorfer et al., 2011; Muller et al., 2008). However, such research had not explored the particular ways in which VLs made sense of and managed such misconceptions. In our study, we were interested in how the VLs dealt with Keoni's misconception that it is okay to place the focus by what looks right (and not attend to the value of *p*). Initially, the VLs repeated Keoni's

initial incorrect placement of the focus at $(0, 1)$. However, the VLs immediately entered into a negotiation process by leveraging previously-appropriated voices (e.g., the geometric definition of a parabola). Keoni's idea of where to place the focus was adhered to by the VLs but was also treated as being open for negotiation. This suggests that the video participants were seen as providing crucial information, but such information did not go unquestioned. Having experienced this negotiation process may have given the VLs a set of tools (namely, invoking previously-appropriated voices and drawing on those voices to negotiate with the voices in the video) that they could use to come to a more productive meaning. Indeed, after initially resisting the voice from a video of *p as a distance* as too confusing, the VLs turned to their previously-appropriated voices, in a manner that echoed their negotiation process with Keoni's misconception. Bringing the previously-appropriated voice of graphing via point substitution to the foreground gave the voice of *p as a distance* something to interact with. Through this interaction and negotiation, the VLs came to develop the meaning of *p as a distance*.

8.2 Value of a Bakhtinian-Inspired Lens

We found a Bakhtinian-inspired lens useful for the analysis of our data for two main reasons. First, Bakhtin's emphasis on the words of others as a source for our own words provided a productive orientation for investigating meaning making from mathematics videos. According to Solomon (2012), Bakhtin resisted the two extremes of an individual creating meaning alone and of individuals being determined by social context. Instead, words are partly one's own and partly from other speakers, and each word or expression comes with a history of its use by other speakers or in other texts (Lemke, 2000). Furthermore, from a Bakhtinian perspective, meaning only emerges in the relationship of two or more voices (Kazak et al., 2015). According to Bakhtin (1981):

The word, breaking through its own meaning and its own expression across an environment full of alien words ... harmonizing with some of the elements of this environment and striking a dissonance with others, is able in the dialogized process, to shape its own stylistic profile and tone. (p. 277)

Nikitina's (2005) interpretation of this passage is that in an educational context, one's own voice is not lost but rather joins in with other voices (some in harmony, others in dissonance) to metaphorically create a musical chord, where different voices are still heard but something new is created. Consequently, our analysis focused on how new meanings for the VLs emerged from the interplay of multiple voices.

Second, Bakhtin's conceptualization of appropriation shines a light on the struggle and complexity that we saw in our data. According to Bakhtin (1981), "Language is not a neutral medium that passes freely and easily" (p. 294). Appropriation is not just the reproduction or uptake of the words or discourses of others but involves the making of words/discourses one's own. A person may try on words from others but fail to transform them into one's own voice. For Bakhtin, resistance and struggle are an essential part of meaning construction, rather than something to be avoided. This helped us better understand the episode in which the VLs viewed a video twice but set aside the new voice as too confusing. One might think the VLs did not engage with the video, but their resistance resulted in an episode in which they invoked previously-appropriated voices in order to tackle the part of the task that they could make progress on. After creating a table of values and graphing the parabola, the VLs were able to make use of a voice from the video: *if you know the p-value, you know the location of the focus*. This suggests that the VLs had, in fact, engaged with the video and recalled the video participants' voices later, after which point they were able to appropriate the voice of *p as a distance*.

8.3 Implications

Our discovery of the importance of the VLs invoking previously-appropriated voices and our analysis of how the VLs engaged with misconceptions has implications for classroom or online use. Basic models for the use of mathematics videos have typically envisioned individuals viewing videos independently, whether that be as part of flipped instruction, online courses, or working through individualized video-based online mathematics programs (Evergreen Education Group, 2016; Murphy et al., 2014). However, as Lohr (2009) advises, we need to imagine models of video use that go beyond current structures, including the future development of online student learning communities that could support collaborative viewing and discussion of alternative videos that feature student conversations.

The results from this study suggest that VLs need opportunities to invoke previously-appropriated voices and explore their relationships with voices from the videos. One support for the VLs in this study seemed to be the opportunity to explain one's thinking to a partner. For example, suppose that videos featuring student conversations are assigned to be viewed before class time (as an alternative to viewing procedural, expository videos) in flipped instruction. Then a discussion of both the misconceptions and productive meanings developed in the videos could be an important part of the classroom activities. We suggest that students be given opportunities to discuss with each other their reasoning on tasks related to those from the videos and to explore connections with the ideas presented in the videos.

If videos that feature conversations will be used by students in non-school settings, then online environments are needed that support discussion and collaborative problem solving. Technology to support such collaborative viewing of dialogic online videos is well under way. Existing video annotation platforms (e.g., Edthena.com or Thinglink.com) can be used to add notes to videos that highlight salient utterances or points of confusion in the video participants'

work. There are also tools currently available to make mathematics videos more interactive (e.g., PlayPosit.com or EdPuzzle.com) by having teachers tag places where the video should stop and an assessment question appear (in the form of multiple choice, fill in the blank, or free response). While these tools seem to work best with expository, procedural mathematics videos, they also have the capacity for threaded discussions, which could be used with conceptually-oriented mathematics videos that feature student-student interactions. Further development of technological tools to support mathematical problem solving (e.g., affording users to make multiple representations such as tables, graphs, and equations) could unlock the power of VLs to be in conversation with other users of the videos and support the invocation of their previously-appropriated voices.

References

- Amhag, L., & Jakobsson, A. (2009). Collaborative learning as a collective competence when students use the potential of meaning in asynchronous dialogues. *Computers & Education*, 52(3), 656-667. <https://doi.org/10.1016/j.compedu.2008.11.012>
- Attard, C., & Holmes, K. (2022). An exploration of teacher and student perceptions of blended learning in four secondary mathematics classrooms. *Mathematics Education Research Journal*, 34(4), 719-740. <https://doi.org/10.1007/s13394-020-00359-2>
- Bakhtin, M. M. (1981). *The dialogic imagination: Four essays* (C. Emerson & M. Holquist, Trans.; M. Holquist, Ed.). University of Texas Press. <https://doi.org/10.7560/715271>
- Bakhtin, M. M. (1986). *Speech genres and other late essays* (V. W. McGee, Trans.; C. Emerson & M. Hoquist, Eds.). University of Texas Press. <https://doi.org/10.7560/720466>

- Bandura, A., Ross, D., & Ross, S. A. (1961). Transmission of aggression through imitation of aggressive models. *Journal of Abnormal and Social Psychology*, 63(3), 575-582.
<https://doi.org/10.1037/h0045925>
- Bandura, A., Ross, D., & Ross, S. A. (1963). Imitation of film-mediated aggressive models. *Journal of Abnormal and Social Psychology*, 66(1), 3-11.
<https://doi.org/10.1037/h0048687>
- Barwell, R. (2018). From language as a resource to sources of meaning in multilingual mathematics classrooms. *Journal of Mathematical Behavior*, 50(1), 155-168.
<https://doi.org/10.1016/j.jmathb.2018.02.007>
- Belova, O., King, I., & Sliwa, M. (2008). Introduction: Polyphony and organization studies: Mikhail Bakhtin and beyond. *Organization Studies*, 29(4), 493-500.
<https://doi.org/10.1177%2F0170840608088696>
- Boesdorfer, S., Lorschach, A., & Morey, M. (2011). Using a vicarious learning event to create a conceptual change in preservice teachers' understandings of the seasons. *The Electronic Journal for Research in Science & Mathematics Education*, 15(1).
- Boston, M., (2012). Assessing instructional quality in mathematics. *The Elementary School Journal*, 113(1), 76-104. <https://doi.org/10.1086/666387>
- Bowers, J., Passentino, G., & Connors, C. (2012). What is the complement to a procedural video? *Journal of Computers in Mathematics and Science Teaching*, 31(3), 213-248.
- Carlsen, M. (2010). Appropriating geometric series as a cultural tool: a study of student collaborative learning. *Educational Studies in Mathematics*, 74(2), 95-116.
<https://doi.org/10.1007/s10649-010-9230-0>

- Chi, M. T., Roy, M., & Hausmann, R. G. (2008). Observing tutorial dialogues collaboratively: Insights about human tutoring effectiveness from vicarious learning. *Cognitive Science*, 32(2), 301-341. <https://doi.org/10.1080/03640210701863396>
- Chi, M. T., Kang, S., & Yaghmourian, D. L. (2017). Why students learn more from dialogue- than monologue-videos: Analyses of peer interactions. *Journal of the Learning Sciences*, 26(1), 10-50. <https://doi.org/10.1080/10508406.2016.1204546>
- Cooren, F., & Sandler, S. (2014). Polyphony, ventriloquism, and constitution: In dialogue with Bakhtin. *Communication Theory*, 24(3), 225-244. <https://doi.org/10.1111/comt.12041>
- de Araujo, Z., Otten, S., & Birisci, S. (2017). Teacher-created videos in a flipped mathematics class: Digital curriculum materials or lesson enactments? *ZDM—Mathematics Education*, 49(5), 687-699. <https://doi.org/10.1007/s11858-017-0872-6>
- Evergreen Education Group. (2016). *Keeping pace with K-12 online learning*. Author. https://static1.squarespace.com/static/59381b9a17bffc68bf625df4/t/593efc779f745684e6ccf4d8/1497300100709/EEG_KP2016-web.pdf
- Gholson, B., & Craig, S. D. (2006). Promoting constructive activities that support vicarious learning during computer-based instruction. *Educational Psychology Review*, 18(2), 119-139. <https://doi.org/10.1007/s10648-006-9006-3>
- Glaser, B. G., & Strauss, A., (1967). *The discovery of grounded theory: Strategies for qualitative research*. Aldine Press. <https://doi.org/10.4324/9780203793206-1>
- Güler, M., Kokoç, M., & Önder Bütüner, S. (2023). Does a flipped classroom model work in mathematics education? A meta-analysis. *Education and Information Technologies*, 28(1), 57-79. <https://doi.org/10.1007/s10639-022-11143-z>

- Hirschkop, K. (2021). *The Cambridge introduction to Mikhail Bakhtin*. Cambridge University Press. <https://doi.org/10.1017/9781316266236>
- Hufferd-Ackles, K., Fuson, K. C., & Sherin, M. G. (2004). Describing levels and components of a math-talk learning community. *Journal for Research in Mathematics Education*, 35(2), 81-116. <https://doi.org/10.2307/30034933>
- Franciso, J. M. (2013). Learning in collaborative settings: Students building on each other's ideas to promote their mathematical understanding. *Educational Studies in Mathematics*, 82, 417-438. <https://doi.org/10.1007/s10649-012-9437-3>
- Freeman, S., Eddy, S. L., McDonough, M., Smith, M. K., Okoroafor, N., Jordt, H., & Wenderoth, M. P. (2014). Active learning increases student performance in science, engineering, and mathematics. *Proceedings for the National Academy of Sciences of the United States of America*, 111(23), 8410–8415. <https://doi.org/10.1073/pnas.1319030111>
- Kazak, S., Wegerif, R., & Fujita, T. (2015). The importance of dialogic processes to conceptual development in mathematics. *Educational Studies in Mathematics*, 90(2), 105-120. <https://doi.org/10.1007/s10649-015-9618-y>
- Klinger, M., & Walter, D. (2022). How users review frequently used apps and videos containing mathematics. *International Journal for Technology in Mathematics Education*, 29(1), 25-35.
- Kolikant, Y. B. D., & Broza, O. (2011). The effect of using a video clip presenting a contextual story on low-achieving students' mathematical discourse. *Educational Studies in Mathematics*, 76(1), 23-47. <https://doi.org/10.1007/s10649-010-9262-5>
- Kolikant, Y. B. D., & Pollack, S. (2015). The dynamics of non-convergent learning with a conflicting other: Internally persuasive discourse as a framework for articulating

- successful collaborative learning. *Cognition and Instruction*, 33(4), 322-356.
<https://doi.org/10.1080/07370008.2015.1092972>
- Koschmann, T. D. (1999). Toward a dialogic theory of learning: Bakhtin's contribution to understanding learning in settings of collaboration. In C. M. Hoadley & J. Roschelle (Eds.), *Proceedings of the computer support for collaborative learning conference*. International Society of the Learning Sciences.
- Lemke, J. L. (2000). Material sign processes and emergent ecosocial organization. In P. B. Andersen, C. Emmeche, N. O. Finnemann, & P. V. Christiansen (Eds.), *Downward causation: Minds, bodies and matter* (pp. 181-213). Aarhus University Press.
- Linell, P. (2009). *Rethinking language, mind, and world dialogically: Interactional and contextual theories of human sense-making*. Information Age Publishing.
- Lo, C. K., Hew, K. F., & Chen, G. (2017). Toward a set of design principles for mathematics flipped classrooms: A synthesis of research in mathematics education. *Educational Research Review*, 22, 50-73. <https://doi.org/10.1016/j.edurev.2017.08.002>
- Lobato, J., Walters, C. D., Walker, C., Voigt, M. (2019). How do learners approach dialogic, on-line mathematics videos? *Digital Experiences in Mathematics Education*, 5(1), 1-35.
<https://doi.org/10.1007/s40751-018-0043-6>
- Lobato, J., & Walker, C. (2019). How viewers orient toward student dialogue in online math videos. *Journal of Computers in Mathematics and Science Teaching*, 38(2), 177-200.
- Lohr, S. (2009, August 19). Study finds that online education beats the classroom. *The New York Times*. <https://bits.blogs.nytimes.com/2009/08/19/study-finds-that-online-education-beats-the-classroom/>

- Mayes, J. T. (2015). Still to learn from vicarious learning. *E-learning and digital media*, 12(3-4), 361-371. <https://doi.org/10.1177/2042753015571839>
- McKendree, J., Stenning, K., Mayes, T., Lee, J., & Cox, R. (1998). Why observing a dialogue may benefit learning. *Journal of Computer Assisted Learning*, 14(2), 110-119. <https://doi.org/10.1046/j.1365-2729.1998.1420110.x>
- McCrone, S. S. (2005). The development of mathematical discussion: An investigation in a fifth-grade classroom. *Mathematical Thinking and Learning*, 7(2), 111–133. https://doi.org/10.1207/s15327833mtl0702_2
- Miles, M. B., & Huberman, A. M. (1994). *Qualitative data analysis* (2nd ed.). Sage Publications.
- Morson, G. S. (2004). The process of ideological becoming. In A. F. Ball & S. W. Freedman (Eds.), *Bakhtinian perspectives on language, literacy, and learning* (pp. 317-331). Cambridge University Press. <https://doi.org/10.1017/cbo9780511755002.016>
- Moschkovich, J. (2004). Appropriating mathematical practices: A case study of learning to use and explore functions through interaction with a tutor. *Educational Studies in Mathematics*, 55(1), 49-80. <https://doi.org/10.1023/b:educ.0000017691.13428.b9>
- Muir, T. (2021). Self-determination theory and the flipped classroom: a case study of a senior secondary mathematics class. *Mathematics Education Research Journal*, 33(3), 569-587. <https://doi.org/10.1007/s13394-020-00320-3>
- Muir, T., & Geiger, V. (2016). The affordances of using a flipped classroom approach in the teaching of mathematics: A case study of a grade 10 mathematics class. *Mathematics Education Research Journal*, 28(1), 149-171. <https://doi.org/10.1007/s13394-015-0165-8>

- Muldner, K., Lam, R., & Chi, M. T. (2014). Comparing learning from observing and from human tutoring. *Journal of Educational Psychology*, 106(1), 69-85.
<https://doi.org/10.1037/a0034448>
- Muller, D. (2008). *Designing effective multimedia for physics education* [Unpublished doctoral dissertation]. University of Sydney.
[http://www.physics.usyd.edu.au/super/theses/PhD\(Muller\).pdf](http://www.physics.usyd.edu.au/super/theses/PhD(Muller).pdf)
- Muller, D. A., Bewes, J., Sharma, M. D., & Reimann, P. (2008). Saying the wrong thing: improving learning with multimedia by including misconceptions. *Journal of Computer Assisted Learning*, 24(2), 144–155. <https://doi.org/10.1111/j.1365-2729.2007.00248.x>
- Muller, D. A., Sharma, M. D., Eklund, J., & Reimann, P. (2007). Conceptual change through vicarious learning in an authentic physics setting. *Instructional Science*, 35(6), 519-533.
<https://doi.org/10.1007/s11251-007-9017-6>
- Murphy, R., Gallagher, L., Krumm, A.E., Mislevy, J. & Hafter, A. (2014). *Research on the use of Khan academy in schools: Implementation report*. SRI International.
https://www.sri.com/wp-content/uploads/2021/12/2014-03-07_implementation_briefing.pdf
- National Council of Teachers of Mathematics (2014). *Principles to actions: Ensuring mathematical success for all*. Author.
- Nikitina, S. (2005). Pathways of interdisciplinary cognition. *Cognition and Instruction*, 23(3), 389-425. https://doi.org/10.1207/s1532690xc2303_3
- Putnam, R. (2021). *Vicarious learning in a prospective secondary mathematics teacher course* [Unpublished master's thesis]. California State University, Channel Islands.

- Radford, L. (2000). Signs and meanings in students' emergent algebraic thinking: A semiotic analysis. *Educational Studies in Mathematics*, 42(3), 237-268.
<https://doi.org/10.1023/A:1017530828058>
- Rogoff, B. (1995). Observing sociocultural activity on three planes: Participatory appropriation, guided participation, and apprenticeship. In J. V. Wertsch, P. del Rio, & A. Alvarez (Eds.), *Sociocultural studies of mind*. Cambridge University Press.
<https://doi.org/10.1017/cbo9781139174299.008>
- Seethaler, S., Burgasser, A. J., Bussey, T. J., Eggers, J., Lo, S. M., Rabin, J. M., Stevens, L., & Weizman, H. (2020). A research-based checklist for development and critique of STEM instructional videos. *Journal of College Science Teaching*, 50(1), 21-27.
- Silseth, K. (2012). The multivoicedness of game play: Exploring the unfolding of a student's learning trajectory in a gaming context at school. *International Journal of Computer-Supported Collaborative Learning*, 7(1), 63-84. <https://doi.org/10.1007/s11412-011-9132-x>
- Solomon, Y. (2012). Finding a voice? Narrating the female self in mathematics. *Educational Studies in Mathematics*, 80(1), 171-183. <https://doi.org/10.1007/s10649-012-9384-z>
- Solomon, A., Guzdial, M., DiSalvo, B., & Shapiro, B. R. (2018). Applying a gesture taxonomy to introductory computing concepts. In *Proceedings of the 2018 ACM Conference on International Computing Education Research* (pp. 250-257).
<https://doi.org/10.1145/3230977.3231001>
- Taylor, A. R. (2003). Transforming pre-service teachers' understandings of mathematics: Dialogue, Bakhtin and open-mindedness. *Teaching in Higher Education*, 8(3), 333-344.
<https://doi.org/10.1080/13562510309402>

- Thompson, P. W. (2011). Quantitative reasoning and mathematical modeling. In L. L. Hatfield, S. Chamberlain, & S. Belbase (Eds.), *New perspectives and directions for collaborative research in mathematics education: WISDOMe Mongraphs* (Vol. 1, pp. 33- 57). University of Wyoming.
- Walters, C. D. (2017). *The development of mathematical knowledge for teaching for quantitative reasoning using video-based instruction* [Unpublished doctoral dissertation]. University of California, San Diego and San Diego State University.
- Wertsch, J. V. (1991). *Voices of the mind: A sociocultural approach to mediated action*. Harvard University Press.
- Weinberg, A., & Thomas, M. (2018). Student learning and sense-making from video lectures. *International Journal of Mathematical Education in Science and Technology*, 49(6), 922-943. <https://doi.org/10.1080/0020739x.2018.1426794>
- Weinberg, A., Corey, D. L., Tallman, M., Jones, S. R., & Martin, J. (2022). Observing intellectual need and its relationship with undergraduate students' learning of calculus. *International Journal of Research in Undergraduate Mathematics Education*. <https://doi.org/10.1007/s40753-022-00192-x>
- Weinberg, A., Thomas, M., Martin, J., & Tallman, M. (2018). Failing to rewind: Students' learning from instructional videos. In T. T. Hodges, G. J. Roy, & A. M. Tyminski (Eds.). *Proceedings of the 40th annual meeting of the North American Chapter of the International Group for the Psychology of Mathematics Education* (pp. 1263-1266), University of South Carolina and Clemson University.