



1-1-2011

Evaluating the Effects of Camera Perspective in Video Modelingfor Children with Autism: Point of View Versus Scene Modeling

Courtney Cotter
Western Michigan University

Follow this and additional works at: <https://scholarworks.wmich.edu/dissertations>



Part of the Applied Behavior Analysis Commons, Child Psychology Commons, and the Clinical Psychology Commons

Recommended Citation

Cotter, Courtney, "Evaluating the Effects of Camera Perspective in Video Modelingfor Children with Autism: Point of View Versus Scene Modeling" (2011). *Dissertations*. 360.
<https://scholarworks.wmich.edu/dissertations/360>

This Dissertation-Open Access is brought to you for free and open access by the Graduate College at ScholarWorks at WMU. It has been accepted for inclusion in Dissertations by an authorized administrator of ScholarWorks at WMU. For more information, please contact wmu-scholarworks@wmich.edu.



EVALUATING THE EFFECTS OF CAMERA PERSPECTIVE IN VIDEO
MODELING FOR CHILDREN WITH AUTISM:
POINT OF VIEW VERSUS
SCENE MODELING

by

Courtney Cotter

A Dissertation
Submitted to the
Faculty of The Graduate College
in partial fulfillment of the
requirements for the
Degree of Doctor of Philosophy
Department of Psychology
Advisor: Scott T. Gaynor, Ph.D.

Western Michigan University
Kalamazoo, Michigan
June 2010

EVALUATING THE EFFECTS OF CAMERA PERSPECTIVE IN VIDEO
MODELING FOR CHILDREN WITH AUTISM:
POINT OF VIEW VERSUS
SCENE MODELING

Courtney Cotter, Ph.D.

Western Michigan University, 2010

Video modeling has been used effectively to teach a variety of skills to children with autism. This body of literature is characterized by a variety of procedural variations including the characteristics of the video model (e.g., self vs. other, adult vs. peer). Traditionally, most video models have been filmed using third person perspective (i.e., scene models), where the viewer is watching the actor perform in a scene. Recently, studies have successfully incorporated the use of first person perspective into video models (i.e., point of view models), where the view is directly from the actor's point of view. Currently, no studies have directly compared the effects of camera angle on learning when video models are used as teaching tools. Six boys with autism ages 4-8 years learned yoked pairs of tasks, with one task assigned to each type of modeling condition. The effects were evaluated using an adapted alternating treatments design that allowed for a direct comparison between conditions with task difficulty held constant. Few differences in rate of acquisition and attention to the model were observed. Video modeling was not always successful as a teaching tool for targeted tasks. Supplemental teaching strategies (e.g., in vivo modeling with error correction)

were employed when video modeling was ineffective for one or both tasks. This study provides evidence that camera angle does not generally have an effect on video modeling effectiveness. It also provides further evidence that video modeling may not always be an effective teaching tool for all children with autism.

UMI Number: 3470401

All rights reserved

INFORMATION TO ALL USERS

The quality of this reproduction is dependent upon the quality of the copy submitted.

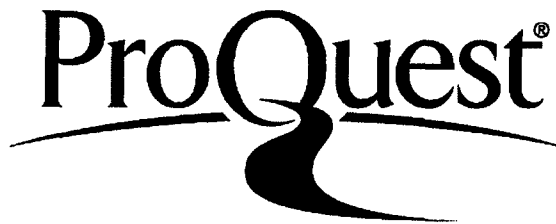
In the unlikely event that the author did not send a complete manuscript and there are missing pages, these will be noted. Also, if material had to be removed, a note will indicate the deletion.



UMI 3470401

Copyright 2010 by ProQuest LLC.

All rights reserved. This edition of the work is protected against unauthorized copying under Title 17, United States Code.



ProQuest LLC
789 East Eisenhower Parkway
P.O. Box 1346
Ann Arbor, MI 48106-1346

© 2010 Courtney Cotter

ACKNOWLEDGEMENTS

I would like to thank Dr. Linda LeBlanc for her help in designing and executing this study. Throughout my graduate career, Dr. LeBlanc has fostered my interest in working with children with autism. Additionally, she has provided me with the technical knowledge needed to ask research questions relevant to this population and design studies to answer these questions. Her guidance on this and many other projects has been much appreciated.

I would also like to thank the other members of my committee, Dr. Scott Gaynor, Dr. Galen Alessi, Dr. Wayne Fuqua, and Dr. Jamie Owen-DeSchryver. Their comments on earlier versions of this manuscript were helpful in refining the final version of this manuscript. Special thanks to Dr. Scott Gaynor for stepping into the role of my dissertation chair when needed.

Courtney Cotter

TABLE OF CONTENTS

ACKNOWLEDGEMENTS	ii
LIST OF TABLES	vi
LIST OF FIGURES.....	vii
INTRODUCTION.....	1
Video Modeling with Children with Autism.....	4
Benefits of Video Modeling	5
Procedural Variations in Instruction	6
Characteristics of the Model.....	9
Rationale for the Current Study.....	14
METHOD.....	15
Participants and Setting.....	15
Screening Procedures	16
Materials and Targets	18
Selection of Target Skills	18
Completed Target Skills	19
Discontinued Targets	23
Video Models	24
Randomization Procedure	26
Measurement.....	27

Table of Contents - Continued

Design.....	30
Procedures	31
Baseline	31
Video Modeling.....	31
Supplemental Instructional Procedures.....	32
Added Instruction	34
In Vivo Modeling with Error Correction.....	35
Edible Reinforcers	35
Prompting and Prompt Fading.....	36
RESULTS.....	36
Attention to the Video	45
DISCUSSION.....	46
REFERENCES	67
APPENDICES	73
A. Craft Tasks.....	73
B. Tangram Tasks	75
C. Drawing Tasks	77
D. Conversation Scripts.....	80
E. Sample Data Sheet for Child Data and Procedural Integrity - Baseline..	81

Table of Contents - Continued

F. Sample Data Sheet for Child Data and Procedural Integrity - Treatment	82
G. HSIRB Approval Letter	83

LIST OF TABLES

1. Descriptive Information About Video Models and Attending Behavior.....	64
2. Number of Trials of Instruction	65

LIST OF FIGURES

1. Nathan: This figure depicts the results of Nathan’s learning of six target behaviors (two craft tasks and one conditional discrimination task).	55
2. Brendan: This figure depicts the results of Brendan’s learning of four target behaviors (construction tasks using tangrams).....	56
3. Dave: This figure depicts the results of Dave’s learning of four target behaviors (two drawing tasks and two construction tasks).	57
4. Ethan: This figure depicts the results of Ethan’s learning of four target behaviors (four construction tasks using tangrams).	58
5. Jake: This figure shows the results of Jake’s learning three pairs of target skills.	59
6. Aidan: This figure shows the results of Aidan learning two pairs of target skills (two craft tasks and two drawing tasks).	60
7. Trials to Criterion: This figure depicts trials to criterion for all targets taught in both camera angle perspectives.....	61
8. Targeted Skills: This figure depicts the total number of skills and subcategories of those skills with respect to performance outcomes and interventions. ..	63

INTRODUCTION

Autism is a developmental disorder that was first identified by Leo Kanner (1943) based on his careful observations of 11 patients. Kanner described the patients as being socially aloof, generally having adequate language but not using it to communicate, and having an insistence on sameness or resistance to change. The definition of autism has been refined over time to allow for more precise diagnosis (Volkmar, Chawarska, & Klin, 2005). Currently, autism is classified in the Diagnostic and Statistical Manual of Mental Disorders (4th ed., text revision) under the class of Pervasive Developmental Disorders (American Psychiatric Association, 2000). It is one of three disorders in this category along with Asperger's Disorder and Pervasive Developmental Disorder – Not Otherwise Specified.

The core deficits of autism include qualitative impairments in communication and social interaction and excesses in restricted, repetitive and stereotyped patterns of behavior (American Psychiatric Association, 2000). Common communication problems include language delays, problems initiating or sustaining a conversation, and use of stereotyped, repetitive or idiosyncratic language. Social interaction problems include poor peer relationships, poor use of nonverbal behaviors (e.g., eye contact, facial expressions) to regulate interactions, and a lack of social or emotional reciprocity. Finally, ritualistic and repetitive behavior includes intense interests, strict adherence to nonfunctional routines, and stereotyped and repetitive movements (e.g., hand flapping).

Autism has become an increasingly diagnosed condition over the past two decades. Previous reports estimated approximately 3.4 in 1000 live births resulted in autism (Yeargin-Allsopp et al., 2003) while the current estimate has doubled to approximately 6.7 in 1000, or 1 in 150 children diagnosed with an autism spectrum disorder (Centers for Disease Control and Prevention, 2007). Males are four times more likely to be affected by autism than females (Yeargin-Allsopp et al.), though females are more likely to have comorbid mental retardation (Centers for Disease Control and Prevention). Children with autism are found in all countries and in all socioeconomic classes.

Of the many treatments for autism, Applied Behavior Analysis (ABA) offers the only intervention with empirical support for producing significant improvements in core deficits and overall intellectual and adaptive functioning (Green, 1996; Lovaas, 1987; Reichow & Wolery, 2009; Eikeseth, Smith, Jahr, & Eldevik, 2007). Early Intensive Behavioral Intervention (EIBI) generally consists of up to 40 hours per week of intensive one-to-one training with the child using common ABA instruction techniques, such as prompting, reinforcement, shaping, and modeling to teach new skills (Green). EIBI was originally identified as an effective treatment for children with autism in Lovaas' 1987 study. This study compared two groups of children with autism. The first group received 40 hours per week of intensive one-to-one treatment, while the second group received 10 hours per week of intensive one-to-one treatment. Treatment was provided for a minimum of 2 years. At the end of the study, 47% of the group receiving 40 hours per week intensive one-to-one treatment achieved average intellectual functioning and

educational functioning, while only 2% of the group receiving 10 hours per week of intervention achieved this same outcome. A number of recent studies have replicated this effect, providing further support for the notion that EIBI can increase the intellectual and adaptive repertoires of children with autism (Cohen, Amerine-Dickens, & Smith, 2006; Eldevik, Eikeseth, Jahr, & Smith, 2006; Howard, Sparkman, Cohen, Green, & Stanislaw, 2005; Sallows & Graupner, 2005; Smith, Eikeseth, Klevstrand, & Lovaas, 1997).

In their replication of Lovaas (1987), Cohen et al. (2006) found that after three years, the EIBI group had significantly higher IQ and adaptive behavior scores than the comparison group with a typical special education curriculum. Additionally, of the 21 participants in the EIBI group, 6 advanced to regular education without support and 11 advanced with support, compared to only 1 of 21 who advanced in the comparison group (Cohen et al.). Howard et al. (2005) found similar results when they compared EIBI with “eclectic” special education interventions and non-intensive special education curriculums. They found the EIBI group had higher mean standard scores than both groups in cognitive and adaptive functioning (Howard et al., 2005).

Recently, researchers have examined variations of the traditional Lovaas (1987) model and showed similar effectiveness (e.g., Eldevick et al., 2006; Sallows & Graupner, 2005). Sallows and Graupner (2005) compared the traditional EIBI intervention to a parent-directed group with equal hours of instruction, but less supervision. After four years of treatment, they found the groups to have similar improvements in intellectual and adaptive skills (Sallows & Graupner). Additionally, out of both groups, 48% of the

children had advanced to regular education classrooms, which was consistent with Lovaas' original effects. Alternatively, Eldevick et al. (2006) compared lower intensity EIBI (12 hours per week) with eclectic interventions. After two years of treatment, they found the EIBI group made larger improvements than the eclectic group (Eldevick et al.). However, the results were not as robust as previous research with more hours per week of the EIBI intervention, suggesting that length of instruction is a critical variable to effectiveness. While this research is a step in the right direction, there needs to be further examination of the critical variables that make EIBI effective, as well as instructional techniques that enhance the efficiency of instruction.

Video Modeling with Children with Autism

One potential means to enhance early intensive behavioral intervention with children with autism is incorporation of technology in teaching situations (Goldsmith & LeBlanc, 2004). Video is a particularly popular technology enhancement due to ease of use, accessibility, and low cost (Goldsmith & LeBlanc). Video has typically been incorporated into instruction with individuals with autism as a means of providing an appropriate model for the child to imitate. Video modeling involves the learner observing a video of a model correctly performing the target behavior and then performing the target behavior himself (Delano, 2007). Video modeling has been successfully employed to teach a variety of skills to children with autism including social initiations (Nikopoulos & Keenan, 2004), perspective taking (Charlop-Christy & Daneshvar, 2003; LeBlanc et al., 2003), giving compliments (Apple, Billingsley, & Schwartz, 2005), and engaging in conversational speech (Charlop & Milstein, 1989).

Video modeling has become such a popular instructional technique for children with special needs that the *Journal of Positive Behavior Interventions* devoted a special issue to the topic (Sturmey, 2003).

Benefits of Video Modeling

There are many suggested benefits for using video modeling with children with autism. First, video modeling removes the social component of instruction, which could be aversive for children with autism, allowing the child to focus solely on the target skill (Bellini & Akullian, 2007). Additionally, it has been hypothesized that many children with autism respond best to visual stimuli, so instruction that depends heavily on visual observation, such as video modeling, may be suited to their particular needs (Sherer et al., 2001). Videos also allow consistency of presentation of the behavior across trials and can allow the therapist to isolate and enhance aspects of the behavior that are particularly salient to acquisition (LeBlanc et al., 2003). Video models can also be observed in the absence of a trained therapist, increasing the amount of exposure a client is likely to have to the modeled behavior.

One study suggests that video modeling may also be a particularly efficient instructional method compared to live or in vivo modeling. Charlop-Christy, Le, and Freeman (2000) conducted a comprehensive study comparing the effectiveness and efficiency of video modeling with in vivo modeling. In all cases, video modeling required fewer training sessions to skill mastery and skills taught via video modeling generalized across people, settings, and stimuli. Charlop-Christy et al. also recorded the time and cost efficiency of in vivo modeling versus video modeling. For four of five participants,

video modeling required less time to implement than in vivo modeling. For the fifth participant, equal amounts of time were required for video and in vivo modeling. In all cases video modeling was more cost effective than in vivo modeling. Although this remains the only published experimental comparison between video and live modeling, an extensive literature has documented the beneficial effects of video modeling using various procedural variations.

Procedural Variations in Instruction

The procedural variations employed in video modeling are numerous and have primarily focused on different characteristics of the video and different methods of implementing the video (i.e., supplemental instructional components, number of times the video is shown). Nikopoulous and Keenan (2004), LeBlanc et al. (2003), and Charlop and Milstein (1989) provide a small illustration of the variety of skills that have been taught through the use of video modeling, as well as the procedural variability that characterizes this body of literature. These studies are described in detail below as a sample of the typical range of procedures used in video modeling to teach children with autism.

Nikopoulous and Keenan (2004) used video modeling to teach children with autism to initiate social interactions. Following a video model, participants were placed in a room with similar toys and given 25s to initiate a social interaction. If such an interaction occurred, the child was moved to various conditions (i.e., condition similar to baseline, condition with different toys than those observed in the video to measure stimulus generalization) to determine how robust the learned response was. If the child

did not provide a response within 25s of being placed in a condition similar to baseline, the child was shown a video of a less complex interchange. One child showed an increase in social initiations after viewing a video model of a relatively complex exchange between two peers while the other two showed an increase in social initiations after viewing a video model of a less complex interaction.

LeBlanc et al. (2003) used video modeling to teach perspective-taking skills to children with autism. In their study, researchers showed children a video of an adult correctly completing a perspective-taking task. Interestingly, the video also provided the observer with rule statements about how to appropriately complete the perspective-taking task. Children were then asked perspective-taking questions related to the task that was completed by the adult. If children answered correctly, they received a variety of reinforcers. If children answered incorrectly, they were shown the video again until correct responding occurred. Next, children were presented with similar perspective-taking situations in which various stimuli used in the sample task were replaced with slightly different stimuli (e.g., pencil found in M&M's box replaced with pennies). All children were able to correctly respond to the presented perspective taking tasks following video modeling.

Charlop and Milstein (1989) used video modeling to teach children with autism to engage in appropriate conversational speech in the form of several scripted conversations. During baseline, the therapist held the item that was the topic of conversation and said the first line of the scripted conversation regarding that item. The therapist then waited 10s for a response and continued with the next two lines of the

scripted conversation if no response occurred. During video modeling, the entire conversation was modeled on a videotape three times. The therapist then said, "Let's do the same" and provided the first line of the scripted conversation. Generalization probes of untrained conversations were tested to determine whether conversational skills generalized across conversational topics. One participant was able to meet the criterion for both the modeled conversation and the generalization topic after viewing one videotaped conversation. A second participant met the criterion for generalization after viewing the two video models of conversations and a third participant met these goals after viewing three different conversation video models.

The three studies described above illustrate a few of the variations in procedures used in experimental evaluations of video modeling. First, different consequences were provided for incorrect or lack of responding. Nikopoulous and Keenan (2004) provided a less complex video model while LeBlanc et al. (2003) provided the same video for repeated trials and incorporated prompts. Still different, Charlop and Milstein (1989) did not provide additional access to the video, but instead provided lines of the scripted conversation until the child responded or the trial was completed. The videos also differed in the number of exemplars the video demonstrated. Charlop and Milstein showed the desired behavior occurring three times (e.g., complete conversation three times), while LeBlanc et al. and Nikopoulos and Keenan each showed the appropriate behavior once in the video, though the same video may have been shown multiple times. The time allowed for a response also varied, with Nikopoulous and Keenan waiting 25s while Charlop and Milstein waited only 10s. Finally, LeBlanc et al. was the

only study in which the narrator describes the model's behavior and why that behavior occurred during the video (e.g., "He looks in 1 because the footprints lead to 1 – it's a clue" (p. 255). All of the studies achieved positive outcomes, as have most published studies of video modeling, but the lack of direct experimental comparison of differences in procedural implementation prevent conclusions about whether these characteristics differentially impact effectiveness.

Characteristics of the Model

There are several characteristics that might impact the effectiveness of a video model. While many of these characteristics have been included in videos in research studies, very few studies have examined the differential impact of each characteristic and whether that characteristic makes a video model more or less effective. Generally, guidelines for creating models are based on the work of Bandura (1977), who suggests that a) when the model engages in the target behavior the consequences associated with the response (e.g., receipt of reinforcers) should also be depicted, b) the model should either be similar to the learner or should be someone of higher status, and c) the behavior should be modeled in the context in which the learner should engage in the behavior. Researchers have also suggested that video models should focus on the salient aspects of the behavior to be imitated (McCoy & Hermansen, 2007), but it is unclear to what extent videography renders a model more effective. Additionally, researchers have suggested that video models are beneficial because they are highly engaging for children with autism who may be more likely to attend to a television than a person; however, a small data set indicates that attention to video models may vary

across individuals with autism (Dillon & LeBlanc, in press.). Though these particular recommendations seem quite reasonable, they have not been experimentally demonstrated to enhance the effectiveness of video models with children with autism.

Sherer et al. (2001) provides one of the few direct experimental comparisons of different aspects of the video model by comparing video models of a peer actor or edited video depicting the target child performing the skill (i.e., self model). Sherer et al. attempted to teach children with autism conversation skills from a list of twenty questions provided by parents and caregivers. These questions were ones that the participants were unable to answer independently (i.e., without prompts) that the caregivers wished their children were able to answer. Two videotapes were created for each child. In one, a typically developing peer served as the model and in the other, footage from prompted interchanges with the participant was edited to remove the prompts and create an effective self-model. Of the five participants, two acquired the conversations quickly and one acquired them more slowly. Two were unable to reach the acquisition criterion of answering 100% of the questions correct. Of the two that reached acquisition quickly, one showed a preference for self-modeling and one showed a preference for peer modeling. These findings may indicate that more research is needed to determine whether self or other modeling is superior. These findings may also indicate that preferences are individual and that the child's learning preferences should be considered when choosing a treatment option.

More recently, researchers have begun to investigate alternative perspectives to the traditional scene filming. Point-of view video modeling depicts a skill being

performed as it would be seen from a first person point of view (i.e., by the person doing the action), rather than as a third person observer of a scene. Hine and Wolery (2006) investigated the effectiveness of point of view modeling in teaching children with autism to engage in appropriate play behaviors. Two preschool aged children with limited verbal abilities were taught to engage with “sensory materials” (i.e., potting soil and gardening tools, colored rice and cooking tools) appropriately, rather than in the stereotypic way they had previously engaged with the materials. The video models were shown on a laptop computer and depicted an adult’s hands interacting appropriately with the toys shot from first person perspective. Each video was less than two minutes in duration and provided three exemplars of the appropriate behavior. The video model was sandwiched between two short segments of the child’s favorite cartoon and included a verbal cue “Play with your toys,” before the hands entered the scene to model appropriate play. After viewing the video, the child was placed in the context observed in the video and told to play with their toys and to stay near the sensory bin. The intervention was effective for both children for the gardening task, but was only effective for one child for the cooking task.

Shipley-Benamou, Lutzker, and Taubman (2002) also investigated the use of point of view video modeling in teaching daily living skills (i.e., pet care, making orange juice, mailing a letter, setting the table) to young children with autism. Participants viewed a video of the appropriate chain of behaviors shot from the first person perspective, and then were immediately given the opportunity to complete the chain of behaviors. The point of view model was effective for two of three participants and their

new skills maintained when they were asked to complete the task without watching the video. The third participant also learned her target skills in the video modeling condition, but required the television to be lowered to eye level and an additional gestural prompt to attend to the video during the viewing of the video and to the appropriate materials after the video was completed. When these prompts were added, this participant was also successful in mastering each of her target skills, and these skills were maintained in the no video condition and partially maintained in the one-month follow up condition.

Norman, Collins, and Schuster (2001) incorporated point of view perspective in teaching three children with developmental disabilities to complete behavioral chains of self-help skills (i.e., cleaning sunglasses, putting on a wrist watch, and zipping a jacket). Participants viewed a video model of the entire target chain followed by re-presentation of the video footage of one step of the task prior to an opportunity to respond. During initial trials the single steps were presented with a zero second delay, and in later trials, the delay was increased to 5 s to allow participants the opportunity to respond independently. Correct responses resulted in praise, and training continued until the child completed all steps of the chain without re-presentation of individual steps. All of the participants successfully learned the chains from the point of view perspective videos. One learner required differential praise of independent responses and presentation of five massed trials of troublesome steps to facilitate acquisition.

Alberto, Cihak, and Gama (2005) also examined the use of point of view video modeling as a teaching technique; however, the comparison was between point of view

video models and static picture prompts rather than live prompts. Alberto et al compared the use of these two different teaching techniques on community-based independence skills (i.e. taking \$20.00 out of an ATM, using a debit card to purchase groceries) in students aged 11-15 with developmental disabilities. The video modeling condition employed a point of view video model of the entire target chain while the teacher described the behaviors that were being completed at each step of the chain. The other condition employed a series of static pictures that depicted point of view representations of the steps of the target chain while the teacher described the behaviors being completed. Each student learned one target using the video and one with the pictures with counterbalancing across participants (i.e., one participant learned to withdraw money from an ATM with static pictures and one participant learned this skill with video modeling). Both point of view video models and static picture prompts were effective teaching strategies for the participants in this study. Results were mixed in terms of efficiency of the teaching strategy and number of errors made in each teaching strategy. Four of eight participants reached mastery criterion in a smaller number of sessions in the static picture prompting condition, and the same was true for the video modeling condition. Three participants made a greater number of errors in the static picture prompting condition, while five participants made a greater number of errors in the video modeling condition. While point of view video modeling was not found to be superior to static picture prompting, this study provides further evidence that point of view video modeling is an effective teaching tool for children with developmental disabilities.

Each of the studies described above provides support for the use of point of view video modeling as an effective teaching tool for children with developmental disabilities. However, it remains unclear whether this type of model provides additional benefit compared to traditional scene perspective video models. Researchers have not yet directly compared scene modeling with point of view modeling to determine whether one is differentially more effective as a teaching tool than the other. Videos using scene type modeling occur more frequently in this body of literature, with approximately 6 times more published studies using scene type video modeling than point of view video models (Dillon, LeBlanc, & Geiger, 2009). It is unclear whether the results of studies using point of view type video models can be directly compared to prior implementations of scene models because other characteristics of the video model or other extraneous instructional procedures besides the point of view manipulation could have impacted the results. For example, in the Hine and Wolery (2006) study, the researchers sandwiched the video model between clips of a child's favorite cartoon, which might impact the effectiveness of the video model, regardless of the viewpoint of the model.

Rationale for the Current Study

As stated previously, video models in the research literature have included a variety of characteristics. The impact of each of these characteristics on the effectiveness of a given video model is unclear because of a lack of direct comparison studies. Before we can create "best practice" standards for video modeling, we must first identify the critical characteristics of optimal video models and how to incorporate

these characteristics into the creation of video models for children with autism. The purpose of this study is to determine whether point of view video modeling is differentially more or less effective than scene video modeling for teaching children with autism. This study could be the first in a series of studies investigating each of the characteristics of video models to determine which characteristics create the most effective models. Taken together, this series of studies could provide empirically-based guidelines for the creation of “best practice” video models.

METHOD

Participants and Setting

Six children with a previous diagnosis of Autism or Pervasive Developmental Disorder, Not Otherwise Specified (PDD-NOS) provided by independent school districts or area professionals (e.g., psychologist, social worker) were included in this study. Participants ranged in age from 4 to 8 years old. Specifically, Nathan was 7 years 10 months old, Brendan was 4 years 5 months, Dave was 8 years 11 months, Ethan was 4 years 11 months, Jake was 7 years 11 months, and Aidan was 6 years 10 months at the time of participation in the study. All participants had a strong motor imitation repertoire and a strong echoic repertoire. Brendan and Ethan were the most verbally sophisticated participants. Both of these participants communicated in complete spontaneous sentences. Jake, Nathan, and Dave communicated in spontaneous phrased speech and some full sentences, though the sentences were mostly echolalic. These participants also often relied on caregiver prompting to correctly formulate sentences

that were appropriate for a given context. Aidan's speech was the most impaired. He spoke mostly in phrases, and many of these phrases were echolalia.

All sessions took place in a therapy room in a clinical laboratory area in 1504 Wood Hall. The rooms were approximately 8 feet wide by 10 feet long. A table and chair for the child and a television for the child to view the video model were present in the room. A chair for the adult running the trial was placed to the side of the table where the child worked and the data collector's chair was placed next to the television. Sessions were recorded using a camera that was mounted to the wall near the ceiling opposite from where the child was sitting.

Screening Procedures

The Autism Diagnostic Observation Schedule - Generic was administered with all participants to verify the child's diagnosis of autism or PDD-NOS (ADOS-G; Lord et al., 2000). The ADOS-G is a semi-structured, standardized assessment of communication, social interaction, and play or imaginative use of materials with interobserver reliability ranging from 82% to 91.5% (Lord et al., 2000; Hill et al. 2001) and acceptable validity (Lord et al., 2000). The ADOS-G consists of standard activities that allow the examiner to observe behaviors that have been identified as important to the diagnosis of autism spectrum disorders at different developmental levels and chronological ages. The ADOS-G Module 2 was conducted with all participants except Aidan who participated in ADOS-G Module 1 due to his limited speech repertoire. For both modules, the cutoff score indicating autism was 12 and all participants scored at or above that score. Nathan's score on the ADOS was a 22, Brendan's score was an 18, Dave's score was a 22, Ethan's

score was a 12, Jake's score was a 21, and Aidan's score was an 18. These scores indicate that all of the participants' behaviors are similar to other children with autism spectrum disorders.

In addition, parents of the participants completed the Gilliam Autism Rating Scale-2 (GARS-2). The GARS-2 is a rating scale that can be used with individuals between the ages of 3 and 22 to help in diagnosing autism spectrum disorders. The GARS-2 internal consistency reliability scores for each subscale range from .84 to .94, indicating good reliability. The GARS-2 has also been determined to be a valid instrument for screening individuals suspected of having an autism spectrum disorder (Gilliam, 2006). The GARS-2 provides an Autism Quotient score, and this score is the identified as indicating that it is unlikely, possibly likely, and very likely that a child's behavior is similar to other children with autism spectrum disorders. Autism Quotients between 70 and 84 indicate that it is possibly likely that a child's behavior is similar to other children with autism spectrum disorders, while an Autism Quotient above 85 indicates that it is very likely that a child's behavior is similar to other children with autism spectrum disorder. Nathan's score on the GARS-2 was a 79, Brendan's score was a 100, Dave's score was a 79, Ethan's score was an 89, Jake's score was a 103, and Aidan's score was a 98. This indicates that it is very likely that Brendan, Ethan, Jake, and Aidan's behaviors are similar to other children with autism spectrum disorders, and possibly likely that Nathan and Dave's behavior is similar to other children with autism spectrum disorders. Prior research has indicated that the GARS slightly underestimates characteristics of autism in comparison to the ADOS-G (Mazefsky & Oswald, 2006). The GARS-2 attempted

to correct this issue by lowering the cutoff scores for autism (Montgomery, Newton, & Smith, 2008). Interestingly, however, our findings are consistent with the pattern observed in studies comparing the GARS and ADOS, despite our use of the GARS-2, indicating that the GARS-2 may be similar to the GARS in underestimating characteristics of autism.

Materials and Targets

Selection of Target Skills. After it was confirmed that each participant met criteria for autism or PDD-NOS, an interview was conducted with the child's parent(s). The purpose of the interview was to determine which of several possible target skills might be valuable for each child to learn. Pairs of skills were yoked based on type and difficulty level and were rated by a group of 29 behavior analysts working in an early intensive behavioral intervention program for children with autism. The professionals' experience in behaviorally oriented autism education programs ranged from 2 to 25 years. In order for a skill pair to be included in the study, at least 75% of the professionals had to agree that the two skills were of equal performance difficulty. The parent(s) were provided a list of skill pairs and asked to nominate the pairs that their child could not perform independently that were high priorities for acquisition. At least four target skills, grouped into two yoked pairs, were selected from the skills identified by the parent(s) during the interview. The baseline phase provided confirmation that participants did not already have skills prior to treatment. A few targets that were identified by parent report resulted in problematic patterns (e.g., carryover due to stimulus materials that were too similar) or problematic behavior (e.g., pica of craft

materials) and were discontinued and replaced with new pairs of tasks that would allow a more efficient comparison of the two treatments of interest. The targets selected for experimental comparison and the discontinued targets are described in detail below.

Completed Target Skills. Nathan's target skills were two sets of craft tasks (see illustration in Appendix A) and a conditional discrimination task. The first pair of craft tasks involved learning to make a fish (point of view video model) and a butterfly (scene video model) by assembling geometric figures into the appropriate pattern. The fish consisted of an oval, a crescent, and a small white circle with a dot in it. The crescent placement was at one end of the oval so that it resembled a tail whereas the circle with a dot was placed at the other end of the oval to resemble an eye. The butterfly consisted of two hearts and an oblong shape. The hearts were placed together at their points and the oblong piece was placed in the middle vertically covering the conjoined points of the hearts. The second pair of craft targets involved construction of a clown (point of view) and an ice cream (scene). Both pairs consisted of a larger circle, a triangle, and a smaller circle. For the clown, the child learned to place the triangle above the circle with the smaller circle at the point of the triangle resembling a pom pom on the top of a hat. For the ice cream, the child learned to place the circle on top of the triangle with the smaller circle on top of the larger circle, resembling a cherry on an ice cream. In the conditional discrimination task, the experimenter placed three cards in front of Nathan, two depicting pictures from one category and the third depicting a picture from a different category. In the scene condition, Nathan was shown two pictures from the clothing category, and one picture from the food category. In the

point of view condition, Nathan was shown two pictures from the vehicles category, and one picture from the animal category. Nathan was asked to identify “which one is different” from each array. The locations of the stimuli were rotated across trials.

Brendan and Ethan had the same four target skills involving puzzle constructions with Tangram® shapes (see illustration in Appendix B). Brendan’s first pair of tasks used five shapes to create two patterns with nonsense names, Ping (point of view) and Lud (scene). Each pattern had a different “anchor piece” or initial piece provided for the child to build onto with other pieces. The second pair of five shape patterns was entitled a Fye (point of view) and a Bok (scene) and, again, each had a different anchor piece. Ethan learned the Fye (point of view) and Bok (scene) as the first pair and the Ping (point of view) and the Lud (scene) for the second pair. In order to be considered correct, the pieces had to be placed in the correct orientation to each other, but did not have to be placed in that orientation in any specific order. The children were given a magnetic board with the pieces needed to create the target shape lined up vertically on the left side of the board. The anchor piece was placed in the middle of the magnetic board before the board was handed to the participant.

Dave learned a pair of drawing tasks (see illustration in Appendix C), and a pair of block construction tasks. The drawing targets were a balloon (point of view) and a fish (scene). The balloon required Dave to draw a circle and a small triangle below it with the top point of the triangle touching the bottom of the circle. A wavy line was drawn from the bottom of the triangle for the string of the balloon. The fish required Dave to draw a circle and a small triangle to the right of it with the top of the triangle touching the side

of the circle. A small dot was drawn inside the circle on the opposite side from the triangle for the eye. The construction targets were a boat (point of view) and a truck (scene). These tasks each required using three blocks from the “Go, Diego, Go!” Dora Lego® set. The boat required the child to place the body of the boat on top of the piece that looked like the rudder of the boat. The child then had to place a propeller off the back of the body of the boat. The truck required the child to place a roll bar that was positioned across the top of the truck and a Lego® off the back of the truck that simulated an exhaust pipe.

Jake learned three pairs of tasks. First, he learned to draw a spaceship (point of view) and a boat (scene) (see Appendix C). The spaceship consisted of an oblong shape with a flat bottom and a crescent drawn below the oblong shape with the curved part of the crescent touching the flat part of the oblong shape. A small circle was drawn inside the oblong shape, resembling a window. The boat was drawn with a half circle with the flat side facing the top of the page. A line extended from the middle of the top of the half circle upwards and a small triangle extended to the left of line (i.e., a flag). His second pair of tasks involved conversation scripts about preferred activities for the weekend (point of view) and at recess (scene) (see Appendix D). Each script consisted of two pairs of exchanges between the researcher and Jake. His third pair of targets was a conditional discrimination task of spatial relations. This task involved placing colored blocks in relation to a container or to each other. Specifically, for inside, next to, and on top of, the blocks were placed in relation to a clear plastic container with a lid. The container was round and approximately 3 inches in diameter and 3 inches tall. The

blocks were placed in relation to each other for the middle, left, and right task. We chose to teach the middle, left, and right target in the point of view condition because if this target was taught in the scene condition, the answers given in the video would appear to be backwards (e.g., left being right, and right being left) to the observer. Each video showed the model answering each of three questions. In the inside, next to, and on top of condition, the model answered the questions: “Which one is inside the container?”, “Which one is on top of the container?”, and “Which one is next to the container?”. In the middle, left, and right condition, the model answered the questions: “Which one is on the right?”, “Which one is on the left?”, and “Which one is in the middle?”. Each video had multiple exemplars of these questions being answered for different stimulus arrangements. Three exemplars were shown on every trial so that each trial showed each block color in each location. The questions were also asked in a different order on each trial, to ensure that the child was not engaging in rote responding and was listening to the auditory stimulus provided at the start of the trial.

At the start of the task, the blocks were arranged in the appropriate locations. The researcher asked the relevant question (e.g., “Which one is on the right?” or “Which one is inside the container?”) and waited 5s for an answer. If Jake said the appropriate color name of the block in the queried position, he was praised (e.g., “That’s right! Green is in the middle!”). If Jake said an incorrect color name, the researcher said “No” in a monotone voice and moved onto the next question.

Aidan learned a pair of craft tasks and a pair of drawing tasks. Aidan was taught to create the same butterfly and fish crafts as Nathan (shown in Appendix A). Also the

same as Nathan, the fish was taught in the POV condition and the butterfly was taught in the scene condition. For the drawing tasks, Aidan learned to draw an apple (point of view) and a kite (scene) (see Appendix C). The apple consisted of drawing a circle with a small line extending upward from the middle of the top of the circle for a stem. Next, an oval was drawn to the right side of the line in the corner between the line and the circle on the bottom for the leaf on the side of the apple's stem. The kite consisted of drawing a diamond with a wavy line extending from the bottom point of the kite. A small triangle was drawn to the right of the wavy line, approximately halfway between the bottom of the diamond and the end of the line, resembling a bow on the kite.

Discontinued Targets. Targets were discontinued with three children as problems arose that made it dangerous for the children to complete the tasks or difficult to ensure an appropriate comparison between the two video conditions. Craft tasks involving construction paper were discontinued with Dave because he began to eat pieces of the construction paper and the non-toxic glue stick. For Aidan, the clown and ice cream crafts were discontinued because the similarity of the craft materials for the two tasks resulted in significant interference (i.e., partial acquisition in one condition interfered with the other condition) and a dysfunctional response pattern. Though the colors differed, the shapes were identical but needed to be placed in different configurations for the two conditions. Three errors occurred repeatedly across trials resulting in a hybrid figure with errors being produced in sessions of each treatment condition. First, he placed the small circle at the pointed tip of the "cone" rather than the mouth of the cone in the ice cream condition. This placement was correct for the

clown's hat (i.e., the other condition) but was inaccurate for the ice cream. Second, he placed the triangle upside down and below the large circle in the clown condition, positioning that simulated the cone in the ice cream condition. Third, he consistently failed to place the small dot (step 3) that would have completed the creation of either the clown or the ice cream. Aidan was only able to engage in accurate responding on 3 of 11 POV trials (i.e., 27.2%) and 1 of 13 scene trials (7.7%). Other trials resulted in erroneous responding that demonstrated some type of carryover difficulty. Thus, these targets were replaced with another pair of targets with stimuli that differed more in stimulus properties (i.e., shape or color) to promote discrimination between the two tasks.

Finally, we attempted to teach Nathan the same conditional discrimination task involving spatial positions taught to Jake (i.e., middle, left, right; inside, next to, on top of). Nathan demonstrated a positional bias only in the inside, next to, and on top of condition, answering each question with the color block in the "on top of" position. Because of this bias, we attempted to teach a different conditional discrimination task, which one is different, which is described in detail above.

Video Models. Video models were created for each task. Each video depicted an adult model and videos ranged in duration from 12.9 s to 49.5 s (see Table 1, columns 5 and 8 for individual video durations). The models were all filmed in the same room as later experimental sessions were conducted (i.e., in the therapy room in Wood Hall). Most videos showed one exemplar of the required response; however, the conditional discrimination videos showed three exemplars of each required response. Each video

contained the spoken instruction (e.g., “Make a butterfly”) that was later repeated to the child after viewing the video.

The point of view video models were filmed in the first person perspective (i.e., point of view of the person completing the task). The only items visible in the video were the model’s hands and the objects relevant to the task. For example, the craft video showed adult hands holding and then placing the construction paper pieces in the proper orientation. The scene models were filmed from the third person perspective (i.e., as though the viewer was watching an actor on stage). The model’s head, torso and arms as well as the objects needed to complete the task were shown in the video. The hands in the point of view video model and the head, torso, and arms of the model in the scene type video models took up a comparable amount of the television screen. In the point of view video models, the model was shot from approximately 1 foot from the materials and hands modeling the task. In the scene models, the model was shot from approximately 3-4 feet from the materials and the person modeling the task. This difference accommodated the incorporation of context in the scene models (i.e., a portion of the room shown as the background for the video model).

For drawing and construction tasks, scene videos would typically require the participant to engage in some perspective taking in order to complete the task accurately. That is, the video viewer sees the task being completed upside down or at a one hundred eighty degree rotation and would have to alter the overall positions to complete the task. Because children with autism notoriously struggle with perspective taking (Shabani et al., 2002, Spradlin & Brady, 1999), we attempted minimize this

difficulty in two ways. In the drawing tasks, the model drew the items on a white board rather than on paper with the perspective directly over the drawer's shoulder (i.e., person is visible in scene but placement does not require rotation). In the other tasks resulting in a permanent product, the model held the item up at the end of the model to show the correct placement at the end of the video.

Randomization Procedure. Trials of each target skill within a yoked pair were rapidly alternated for the first four participants in this study (i.e., Nathan, Dave, Aidan, and Jake). Within each yoked skill pair, one skill was assigned to even numbers and the other to odd numbers and the trial order was determined using a list of numbers generated by a random numbers generator at www.random.org. That is, each time an even number appeared in the list, the even number assigned skill was targeted and each time an odd number occurred, the odd number assigned skill was targeted. Due to the randomization procedure, in some instances there are groups of up to 5 trials in a row for these participants. However, most trials of each target occur less than 3 times consecutively, resulting in a very rapid alternation between the two targets.

Some difficulty with carryover was observed in Aidan's discontinued craft task (i.e., clown and ice cream). Aidan also demonstrated difficulty related to the acquisition of his drawing task, with his error involving the omission of one part of the drawing that was similar in both conditions (small oval on the side of the stem for the apple, small triangle on the side of the string of the kite). When other methods of enhancing the video were not effective in allowing Aidan to demonstrate accurate responding in either condition, Aidan was presented with massed trials of each target. Despite massed trial

presentation, Aidan still did not completely master these targets in video modeling. To avoid problems with carryover in subsequent participants (i.e., Brendan and Ethan), alternation of trials between the two targets occurred less rapidly such that at least two trials of one target would occur before switching to the other target. To do this, a list of numbers was generated from random.org.; however, the number list was modified slightly resulting in a quasi-random sequence. In any instance in which one even number was followed by one odd number (i.e., would result in rapid alternation), a second even number was inserted following the first even number (for a total of 2 trials), and then a second odd number was added following the first odd number (for a total of 2 trials). One task was then assigned to even numbers and the other to odd numbers. In this way, at least two trials of a given task occurred before alternation to the other skill occurred.

Measurement

The primary dependent measure was the percentage of skill steps completed accurately and independently. Primary scorers collected paper and pencil data during the session in the room with the researcher and participant. Sample data sheets can be found in Appendices E and F. A correct independent step for the craft tasks was scored when a particular piece of the shape was placed in the correct orientation to all other shapes and attached using the glue stick. A correct independent step for the Tangram tasks or construction tasks was scored when a piece was placed in the correct location and orientation to all other shapes in the item. A correct independent step for the drawing tasks was scored when the child drew a shape, or a very close approximation of the shape, in the correct orientation to all other shapes. A correct independent step for

the conversation script was scored when the child said the entire correct statement in response to the previous line of the script. A correct independent step for the conditional discrimination task was scored when the child correctly named the colored block in the corresponding location (i.e., right, left, inside) or correctly identified the item that was different. Partial responses were not scored as correct for any of the target skills (e.g., partially drawn shape, correctly drawn shape in the wrong orientation, partial statement of the scripted statement). However, all tasks had 2-5 steps, and each independently completed step was scored as correct. As a result a child could receive a score between 0 and 100% on any given trial.

The number of trials to the mastery criterion (i.e., 100% accuracy on 4 consecutive trials) was computed for each skill for multiple comparisons across conditions. The percentage duration of attending to the video models was scored from video footage. To score these data, observers started a stopwatch when the child was oriented to the television screen to start watching the videotape (i.e., "Watch this!" while researcher pointed to the television screen). The stopwatch was stopped any time the child shifted his gaze from the television screen. The resulting data is the duration of eye orientation toward the video, which is discussed as total duration of attendance. The number of seconds of attention was divided by the total duration of video length (calculated by timing from the moment the child was oriented to the television screen to the end of the video for each trial and then averaged across trials) and then multiplied by 100% to obtain percent of attendance to the video.

A second trained observer collected interobserver agreement (IOA) data for at least 25% of trials across all phases of the study. Total agreement was calculated as the number of specific steps scored identically by the two independent observers divided by the total number of scored steps and converted to a percentage. All sessions were videotaped using a wall-mounted camera positioned in the corner of the room and IOA data was collected from videotape. IOA was collected on 93.1% of baseline trials and was 99.3% (range 91.7%-100%). IOA was collected on 79.1% of treatment trials and was 96.1% (range 88.0%-100%). Overall, IOA was collected on 80.9% of trials across all participants and all phases of the study and was 96.6% (range 88.0% to 100%). IOA for total duration of attending was calculated by summing duration of attention for each target for the primary and secondary observer (i.e., sum of primary observer and sum of secondary observer). The smaller sum (i.e., smaller total duration) was divided by the larger sum (i.e., larger total duration) and this result was multiplied by 100%. IOA for total duration of attending was collected on 27.8% of trials across participants and was 91.3% (range: 64.4% to 100%).

Procedural integrity was scored from video footage for at least 25% of trials across all phases of the study. Because each condition of the study required the researcher to engage in several potentially different steps, treatment integrity was calculated as the percentage of steps correctly completed by the researcher. A sample data sheet for procedural integrity can be found in Appendices E and F. An average percentage of correctly completed steps was calculated for each condition. Procedural integrity for baseline was 100% and was collected on 88.8% of trials. Procedural integrity for

treatment was 99.9% (range: 93.1%-100%) and was collected on 73.4% of trials. Overall, procedural integrity was collected on 75.6% of trials across all phases of the study and all participants and was 99.9%

A second observer scored 38.3% of trials across all phases for IOA of the procedural integrity measure. Total agreement was calculated such that each step of each trial was an agreement (i.e., scored identically as accurate or inaccurate) or non-agreement (i.e., one accurate, one inaccurate). The number of agreements was divided by the total number of steps (i.e., agreements plus disagreements) and converted to a percentage. Agreement was 99.9% across all participants and phases.

Design

An adapted alternating treatments design was used to evaluate the effectiveness of each type of video modeling (Sindelar, Rosenberg, & Wilson, 1985). This design combines aspects of a multiple baseline design and an alternating treatment design allowing two individual instructional interventions to be compared against the baseline condition. The staggering of the baselines across skills pairs illustrates experimental control to rule out potential validity threats of maturation and familiarity with the preparation with replication of treatment effects or non-effects across skills (within subject) and across participants. The two video modeling conditions (i.e., point of view, scene) were compared within each panel of the adapted alternating treatments design for skills that were yoked on level of difficulty.

Procedures

Baseline. The child entered the room, sat at the table and was given the materials needed to complete the task and the relevant instruction (e.g., “Make a Butterfly”). If a child consistently completed the majority of the steps of the behavior independently and accurately, a different target behavior was selected. Baseline continued until a stable baseline was obtained with a phase length consistent with a multiple baseline stagger before the intervention phase was implemented (i.e., three data points per target in short baseline, at least one additional data point per target skill in the longer baseline). The baselines were kept as brief as possible to avoid evoking problem behavior as a result of continually presenting a task without instruction or reinforcement.

Video Modeling. The child entered the room, sat at the table in front of the television and was shown a video of an adult model completing the target behavior. Following the completion of the video, the child was given the appropriate materials to complete the behavior. For all participants except Jake, the instruction to “Do what you saw in the video!” or “Now it’s your turn!” and the instruction relevant to the task (e.g., “Make a boat.”) were provided. During the conditional discrimination tasks for Jake and Nathan, the phrase “Do what you saw in the video” was not provided and instead just the discriminative stimulus for the conditional discrimination trial (e.g., “Which one is left?”; “Which one is different?”) was presented. For the script task, Jake was told, “Let’s say what they said” instead of “Do what you saw in the video.” If the child completed all steps of the behavior accurately and independently, praise and a short break were

provided before the next video presentation and task trial. If the child did not complete all steps accurately and independently, the child was told, “No,” in a monotone voice, was praised for working, and a short break occurred before the next video model viewing and trial presentation. The child was provided with mild praise or mild negative statement on each step of the task to provide the child with some feedback on their performance. That is, for each correct step, the child was given the feedback “Good,” or “Yes,” in a monotone voice for a correct step and the feedback “No,” in a monotone voice for an incorrect step. This feedback was given only once per step, that is, if the child responded to this feedback by changing his or her response on a step and then looking for additional feedback, no additional feedback was given. If all steps were completed correctly by the end of the trial, the child was praised at the end of the trial. This praise occurred even if the child had completed a step incorrectly, but self-corrected this step either before or after hearing the monotone “No.”

The mastery criterion was set at 4 consecutive trials with 100% accuracy for one target skill. When one skill reached the mastery criterion, instruction continued with the other target for at least 2 additional trials of the other skill and until there was no increasing trend in acquisition (i.e., if progress continued with the other skill, trials would have continued until mastery).

Supplemental Instructional Procedures. One form of supplemental instruction was provided as part of the direct experimental comparison of video type. If participants learned one target in the initial comparison, attended to the videos consistently during the initial comparison, and targets were conducive to being taught in the alternative

video modeling condition (e.g., scene target taught in POV), the researchers attempted to teach the unlearned target using the alternative video modeling condition. The procedures in this supplemental instruction condition were the same as described above in video modeling except that only one skill was targeted (i.e., the unmastered skill) and the perspective of the model changed in the video. This occurred for Ethan's first target and Jake's second target.

There were four cases in which at least one target was learned in the initial comparison, but the unlearned target was not taught in the alternative video modeling modality. First, Dave's attention to the video was poor during the initial comparison, which led researchers to use other supplemental instructional procedures that did not rely on videoed instruction for his supplemental teaching procedure. Secondly, for Aidan's first target, the error he was making in the POV condition was a very small detail that would not have been highlighted by scene type video modeling. Thus, it was determined that this target was not conducive to being taught in scene video modeling, and instead instructions were added to the POV video. For Jake's third target, the unlearned target (Left, Right, and Middle) was not conducive to being taught in scene type video modeling. Finally, Nathan's third target, the conditional discrimination task to identify the item that was different, presented the concern that Nathan might not have learned a generalized "different" selection but may have simply learned to select a particular picture that was modeled in the video. Teaching Nathan a generalized "different" repertoire would have required extensive, multiple exemplar training which might have proven lengthy and complicated and detracted from our primary purpose of

comparing the two types of models. Because of this, we used a simple prompting procedure to finish teaching the unmastered target to Nathan once he mastered selection of the appropriate picture in one condition.

Additional types of supplemental teaching procedures were used based on analysis of performance errors during video modeling. These supplemental teaching procedures were not considered part of the experimental analysis, but were put in place simply to ensure that participants learned their identified targets, and benefited from participation in the study.

Added Instruction. One type of supplemental instruction involved adding verbal instructions to the videos used in the original comparison of point of view and scene type video models. These instructions consisted of the model giving a verbal description of the action as the task was completed. For example, in the tangram task, the model would name the color of the relevant shape and identify where it should go in relation to the other shapes (e.g., red, red goes on the blue). This instruction was given while the appropriate action was completed in the video. Additional instructions were added if a child was not acquiring either target skill in the initial comparison and had good attendance to the videos in the original comparison. For one participant (i.e., Aidan), this instructional technique was further enhanced by repeating at the end of the video the verbal instructions for the portion of the task that was being omitted. That is, a narrator stated the instruction while completing the task in the video (e.g., circle, line, oval on the side) and then repeated the instruction for the portion of the task Aidan was omitting (e.g., remember, oval on the side) at the end of the video.

In Vivo Modeling with Error Correction. Another type of supplemental instruction was in vivo modeling with error correction. This type of supplemental instruction was used if a child had poor attendance to the video models in the initial comparison (e.g., Dave), or if a particular skill did not lend itself to being taught in the other video modality (e.g., small detail being a repeated error with the correction too small to be highlighted in scene modeling). In this condition, in vivo modeling consisted of the examiner providing a live model of the task being completed correctly and the immediately giving the child the opportunity to complete the task. Error correction consisted of a verbal statement identifying the error (e.g., reminder to include a missing part of an item, instruction to move a given part of a constructed target). For the scene target in pair 1 for Dave, a descriptive statement (i.e., labeling the shape being drawn) was also given during the in vivo model. This was the only participant for whom this verbalization was added to the in vivo model.

Edible Reinforcers. Another type of instruction involved providing children with edible reinforcers for correct responding, and this type of supplemental instruction was used in conjunction with other types of supplemental instruction (e.g., in vivo modeling) when decreased motivation to respond seemed to be interfering with acquisition. Decreased motivation was defined as children being able to respond accurately or at above chance levels on a large percentage of trials, but not responding accurately on enough trials consecutively to meet mastery criterion. For example, Jake was able to identify left, right, and middle at above chance responding (66% or higher) approximately 40% of the time, indicating that his lack of correct responding may have

been related to boredom or a lack of motivation to attend to the stimuli. Decreased motivation was also defined as a child beginning to engage in mild verbal protest when asked to complete a given task. For example, Dave began to make a whining noise when asked to draw the fish, and so it was decided that he would be offered edible reinforcers for correct responding.

Prompting and Prompt Fading. A final type of supplemental instruction used was prompting and prompt fading, specifically a proximity prompt conducted either by moving the child's hand close to the correct response or by moving the correct response closer to the child. This was used only with Nathan and Jake for their conditional discrimination tasks.

RESULTS

The individual results for each participant allow a detailed examination of response patterns and the effects of supplemental teaching conditions. Table 2 provides information on the number of trials conducted with each child per target in video modeling, and, if necessary, supplemental instruction conditions. Trial by trial data are depicted graphically in Figures 1-6. Summary bar graphs of trials to criterion can be found in Figure 7. At times, the total number of trials conducted with participants (depicted in Table 2 and Figures 1-6) differs from the number of trials to criterion (depicted in Figure 7). This occurred if it was not immediately apparent to the researcher that the mastery criterion was met, and as a result, additional trials were conducted past the point of mastery. For one participant, Jake (Figure 5, panel 2) 10 additional trials were conducted past the point of mastery. The task for which this

occurred was a drawing task, and while Jake accurately drew the target figure, he enhanced the drawing by adding scenery. Trials continued to be conducted with this participant until consultation with the research team could be completed to determine whether adding additional scenery to the target figure should be considered accurate responding.

For two participants, video modeling was generally an effective intervention with no clear difference between the point of view condition and the scene condition. Nathan (Figure 1) readily acquired both targets for two different skills (top and middle panels). In the point of view condition for each of these two skills, he demonstrated accurate responding from the first trial and met the mastery criteria in four trials. His performance was slightly more variable in the scene condition, where across the initial three trials of each task he typically completed only one of two steps accurately before mastery occurred on trial seven. For Nathan's third skill (bottom panel), he acquired the skill quickly in the point of view condition, meeting mastery criterion by trial six. Following the completion of the point of view target, seven additional trials were implemented in scene modeling without meeting the mastery criterion before prompting and prompt fading were implemented as a supplemental intervention. Nathan met mastery criteria after 21 trials in the prompt and prompt fading condition. Brendan (Figure 2) mastered the targets in the same number of trials in both conditions for each task pair that was evaluated (pair one: 20 trials; pair two: five trials; represented in top and bottom panels). Though the number of trials to mastery varied across the two task pairs (i.e., acquisition was less variable and more rapid for the

second set of tangram tasks than the first) the between condition comparison results were replicated. That is, the different camera perspectives of the video models did not result in differential effectiveness. The faster acquisition for the second pair of targets may have been due to increased familiarity with the tangram preparation, fatigue occurring in the session containing trials 28-38 of the first task pair, or because the randomization procedure happened to result in several consecutive trials of the point of view condition in a row before the scene condition trials occurred for the second pair of targets.

Dave's graph (Figure 3) illustrates that video modeling was generally ineffective and that supplemental interventions were required for most targets. He mastered the point of view skill of the first pair of targets (top panel) in 22 trials, but never mastered the scene modeling target with stable low responding occurring in trials 1-15 and after trial 17. More accurate responding was observed in trials 15 and 17, but this accurate responding did not maintain in subsequent trials. Because Dave's attention to the video was poor (POV: 52.4%, Scene: 58.9%), and he was engaging in mild verbal protests when asked to draw the scene target, the supplemental procedure of edible reinforcers with in vivo modeling and error correction was used for the first targets resulting in increased variability in responding and mastery in 35 additional trials. For the next pair of targets (bottom panel), Dave was unable to master either target in the video modeling phase. Responding for one target (point of view condition) was consistent at 50% (same 2 of 4 steps) while all steps were consistently inaccurate for the other target (scene condition). Again, because Dave's attention to the video was poor (POV: 59.0%, Scene: 58.9%), in

vivo modeling with error correction was added though no edible reinforcers were added because no noncompliance or problem behavior was evident. The in vivo modeling with error correction produced increased accuracy with eventual mastery of the target originally taught in the point of view condition in 21 trials and the target originally taught in the scene modeling condition in six trials.

Ethan's graph (Figure 4) also shows that supplemental interventions were required for the majority of targets. For the first pair of targets (top panel), he acquired the point of view target in 5 trials, but was unable to master the scene target in 15 trials. Ethan then viewed the scene target in a point of view perspective video, but remained unsuccessful in learning this target after 7 additional trials. Ethan identified the changed perspective when the scene target was shown using point of view modality, making comments such as "But they're making it upside down!" He continued to construct the target using the orientation from the scene video. The graph of Ethan's second set of skills (bottom panel) illustrates highly variable responding and no mastery in either of the video modeling condition (POV: 21 trials, Scene: 23 trials). In an attempt to increase the effectiveness of video modeling, verbal instructions were added to the video naming each step as the model completed it. Though Ethan echoed these directions while completing the task, he was still unable to master the skills in 21 point of view trials and 20 scene trials. In vivo modeling with error correction was then implemented and proved effective with mastery in five trials for the scene target. Three of six trials of in vivo modeling with error correction for the point of view target were accurate, independent responses. Independent responses are denoted within this phase by gray

shaded data points. Open data points indicate that error correction was provided (i.e., an open data point at 100% would indicate that the participant made the correct response after correction).

The pattern illustrated in Jake's responding (Figure 5) is quite mixed with good effects of video modeling with some targets and one completely unsuccessful target. For the first pair of targets (top panel), Jake mastered the point of view and scene targets in an equal number of trials (16 trials). On the second pair of targets (middle panel), Jake mastered the point of view target in 9 trials but did not meet the mastery criterion for the scene target within 20 trials. Jake was also unable to learn the scene target when it was taught using a point of view video (6 trials). Correct performance occurred only when error correction was added, such that accurate independent performance was never achieved. Additional trials with error correction were not completed because Jake's approved time for participation (i.e., 10 sessions) expired. In the final pair of targets, Jake was able to learn the scene modeling target in 41 trials but was unable to master the point of view target in 43 trials. Because this particular point of view target (i.e., left, right, middle) was not conducive to scene perspective video modeling, prompting and prompt fading and edible reinforcers were implemented. Edible reinforcers were implemented because Jake exhibited pattern of perfect completion trials alternated with 0% to 33% (chance responding) accuracy trials on which he seemed uninterested in the task, suggesting that part of the performance difficulty may have been related to motivation. However, this teaching strategy, of prompting and prompt fading and edible reinforcers for correct responses was also an

ineffective teaching strategy for Jake with inconsistent responding persisting after 48 trials.

Aidan's graph (Figure 6) shows that he mastered the scene target in 13 trials, but was unable to master the point of view target after 17 trials for his first pair of targets (top panel). Additional spoken directions were then embedded in the point of view video. This strategy was selected because Aidan's error involved placing the eye of the fish in the center of the fish instead of on the side of the fish opposite the tail and the placement detail might not have been very salient in the scene model. The additional instruction was "All the way over!" and occurred while the model slid the eye from the middle of the fish all the way to the edge of the fish. However, point of view with added direction was also not effective for teaching Aidan with performance maintaining at 50% of steps accurate for 6 trials.

In his second pair of targets (bottom panel), Aidan was unable to master either target skill using video modeling. Because Aidan's attention to the videos was good and neither target was mastered, instructions were embedded in the video to see if Aidan was able to learn the targeted skill. The instructions added were verbal descriptions of the model's actions as she completed each step. When no change in performance was observed for 8 point of view trials and 10 scene trials, additional instructions were added in the form of a reminder to add the forgotten item while the model retraced this item (i.e., "circle, line, oval on the side, remember, oval on the side"). No change in performance was observed over 20 point of view trials and 23 scene trials so the next strategy was implemented. As a next strategy, the alternation between targets was

discontinued and trials of one target (i.e., scene) were massed. Massed trials were presented for both the scene and the point of view target, but neither target was acquired as a result of massed trial presentation. Finally, following massed trial presentation of the scene target, in vivo modeling with error correction was implemented for the scene modeling target (6 trials). Aidan required error correction on all trials in order to draw the target correctly. He never independently completed the target accurately, as denoted by the open data points within this phase. Because each participant was only available for 10 sessions, Aidan's participation was completed before in vivo modeling was able to be implemented with the point of view target.

Overall, only one participant (Brendan) mastered all of the target skills during the video modeling conditions without any supplemental strategies. One participant (Nathan) mastered five of six skills during the video modeling conditions, with the target in the scene condition unmastered until prompt and prompt fading was used to teach this skill to mastery. One participant (Jake) mastered four of six targets, with one unmastered target being in the scene condition and one in the point of view condition. Neither of these targets was trained to accurate, independent responding using supplemental instruction (i.e., video in the alternative orientation for unmastered scene target; prompt and prompt fading for unmastered point of view target). Three participants (Dave, Ethan, and Aidan) mastered only one of their four targets using video modeling alone. Video modeling in the modality that was effective for one target was not effective as a remediation strategy for the unlearned target in any case in which it was used suggesting that video modeling, per se, rather than the particular perspective

of the video was the problematic variable. Additionally, video modeling was not made more effective when instructions or repetition of instructions given in the video (i.e., additional instructions) were added to the video, providing further evidence that video modeling was a problematic teaching strategy for most participants. In vivo modeling with error correction was an effective teaching methodology for all three of Dave's unmastered targets and two of Ethan's unmastered targets. It is possible that this teaching methodology would have been effective for Ethan's other target and Aidan's targets, but alternative video methodologies were used to teach these targets and, as stated previously, continued to be ineffective as a teaching methodology. Also, in Aidan's case, time for participation expired and all selected remediation strategies were not able to be implemented for all targets. These data are interesting given the large number of studies available that indicate that video modeling is an effective teaching tool for children with autism (Bellini & Akullian, 2007).

As a general comparison of the rate of acquisition during video modeling, Figure 7 depicts the number of trials to criterion in the original video modeling comparison of camera perspective for every pair of skill targets. Unmastered targets are identified by an asterisk at the top of the relevant bar, and for these targets, the total number of trials presented to the participant is depicted in Figures 7. This figure also depicts the additional number of trials presented when an unlearned target was presented in the alternative video modeling modality as a lighter shaded portion of a stacked bar. Supplemental teaching procedures (i.e., video modeling with instruction, in vivo modeling with error correction, edible reinforcers, and prompting with prompt fading)

are not included in this graph to allow for a comparison of video modeling procedures related to camera angle alone. For Nathan (first panel), point of view video modeling (4 trials per task) was slightly more effective than scene modeling (7 trials per task) for each of his first two tasks and point of view video modeling (6 trials) was also more effective than scene modeling (14 trials) for Nathan's third pair of tasks. For Brendan (second panel), the two conditions were equally effective across skill pairs with the second pair acquired more rapidly than the first in both conditions (i.e., 20 trials for each condition in the first pair, 5 trials for each condition in the second pair). Dave (third panel) and Ethan (fourth panel) acquired one of two targets in the point of view condition and neither target in the scene condition. Ethan did not acquire the target skill that was taught in the alternative video modeling orientation (scene target fye taught in point of view orientation that was successful for bok). Jake (fifth panel) acquired four of the six target skills with the same number of trials to criterion for the first pair (16, POV; 16, scene), mastery only in the point of view condition for the second pair and mastery only in the scene condition for the third pair suggesting no differential effectiveness. Jake did not acquire skill two of pair two when it was taught in the alternative video modeling orientation. Aidan (sixth panel) acquired one of the two skills in the scene modeling condition and neither in the point of view modeling condition.

With respect to the number of skills acquired, there appears to be only a slight benefit to teaching using point of view video modeling when compared with scene type modeling but other advantages for point of view modeling were noted. All together, ten skills were mastered in the point of view condition while six skills were mastered in the

scene condition (Figure 8). Also, if participants only mastered one of a pair of skills, it was more likely that the skill mastered was in the point of view condition (five instances) than in the scene condition (one instance). In addition, three participants struggled with the perspective taking skills that were required to perform in the scene modeling condition. For these three participants, the target skills in the scene modeling condition were created upside down, and then the children turned the completed product around before giving it to the examiner. This behavior indicates that the child was aware of the correct orientation of the completed product, but did not know how to create it in the correct orientation. Understanding the creation of the product in the correct orientation required perspective taking skills, skills that are notoriously difficult for children with autism (LeBlanc et al., 2003), and that only begins to develop in typically developing children around the age of four (Cox, 1978). Some research has shown that typically developing children as old as 7, and possibly even older, may continue to struggle with more complex perspective taking tasks (Cox).

Attention to the Video

Participant attending to the video models varied more across participants than between conditions for a given participant (see Table 1). When average duration of attending was averaged across all targets, attending was slightly higher in the POV condition (77.6%) than the scene condition (74.3%) but was still well below 100% for both conditions. When pair-based comparisons are made, the average duration of attending was essentially equal (i.e., less than 3% difference for the pair) for 7 of the target pairs across participants (Nathan-1; Brendan-1, 2; Dave-2; Jake-3; Aidan-1, 2).

Attending was higher in the POV condition than the scene condition for 4 target pairs (Nathan-2; Ethan-1, 2; Jake-1) and higher in the scene condition for the remaining three target pairs. A paired samples t-test was conducted on these data (i.e., attention in the POV condition and attention in the scene condition for each pair of skills), and showed that there was no significant difference in attention between the POV and scene condition ($t=1.300$, $p=.216$).

In most cases, attending was comparable across targets for target pairs for all participants. Two important exceptions, however, are Nathan's second target and Ethan's first target. Nathan's attending to the POV video of his second target was 87.2%, while his attending to the scene video of his second target was 69.5%. He mastered the POV target in this pair (as well as in his first pair, which showed no differential attendance) more quickly than the scene target. Ethan's attending to his first POV target was 97.6%, while his attending to his first scene target was 70.2%. Ethan learned his POV target in five trials, and failed to master his scene target. It is possible that the differential effectiveness in learning observed in Nathan and Ethan may be associated with differential levels of attending to the two types of videos for the specific targets mentioned here. However, a pattern of better attention being related to greater accuracy in responding was not observed in other target pairs.

DISCUSSION

Though many studies have documented the beneficial effects of video modeling for instruction with children with autism, only a few have incorporated point of view perspective rather than scene perspective in the model. Four studies have

demonstrated that point of view video modeling can be effective (Alberto, Cihak, & Gama, 2005; Hine & Wolery, 2006; Norman, Collins, & Schuster, 2001; Shipley-Benamou, Lutzker, & Taubman, 2002), however, no studies have directly compared the effects of point of view and scene video modeling with young children with autism. The purpose of this study was to directly compare the effects of the two interventions on rate of acquisition and attention to the model.

One interesting outcome of this study was the limited effectiveness of video modeling, regardless of the perspective, for several of the participants. Only 1 of the 6 participants mastered all of the target skills during the video modeling conditions. Three participants only mastered 1 of their 4 targets in the initial evaluation of video modeling. Interestingly, some other studies comparing video modeling with alternative teaching strategies have not demonstrated that video modeling is a superior teaching technique. For example, Alberto, Cihak, and Gama (2005) demonstrated that both video modeling and static picture prompting were effective in teaching daily living skills to individuals with developmental disabilities. However, results were mixed in terms of efficiency of the teaching strategy and number of errors made using each teaching strategy. As such, when comparing these teaching strategies, a superior strategy could not be identified. Additionally, researchers have found that supplemental interventions sometimes have to be added to increase the effectiveness of video modeling as a teaching tool. For example, Norman, Collins, and Schuster (2001) had to add massed trials and differential praise for independent responding to the video models for one participant to allow the participant to acquire accurate, independent responding.

Shipley-Benamou, Lutzker, & Taubman (2002) had to add prompts for attention to the video to allow one participant to learn from the video model.

Supplemental teaching strategies were also required in this study to allow participants to acquire accurate independent responding. Supplemental instruction added to the video (e.g., added instruction) was not effective in allowing participants requiring supplemental instruction to achieve accurate independent responding. However, in vivo modeling with error correction was an effective teaching methodology for multiple targets, including all three of Dave's unmastered targets and two of Ethan's unmastered targets. This teaching strategy was effective for all targets it was used with, except Aidan's second pair of target skills, which achieved accurate, but not independent responding. It is possible that the error correction was the essential component in allowing participants to achieve accurate, independent responding, but a component analysis would have to be conducted to draw firm conclusions on this. Evaluating the impact of error correction both on video modeling and in vivo modeling may be an area for future research.

For those participants for whom video modeling was effective, it is interesting to consider the behavioral mechanism by which acquisition is achieved. Early research on imitation has hypothesized that imitation repertoires are eventually controlled by conditioned or automatic reinforcement (Bandura & Barab, 1971). That is, behavior that closely matches the behavior to be imitated is reinforced by the degree of sameness between the original behavior and the behavior emitted by the imitator which is referred to as parity. The more closely the behavior of the imitator matches the

behavior to be imitated, the greater the level of reinforcement (Bandura & Barab, 1971). Therefore, it is likely that the imitative behavior of those children for whom video modeling was effective came under the control of these reinforcers and their imitative behavior was acquired and maintained by the reinforcement obtained from the degree of sameness between their behavior and the behavior modeled in the video.

Very little difference in terms of acquisition and attention to the video was observed between the two types of video modeling investigated in this study. However, it is possible that a slight advantage of point of view video models, when compared to scene video models, exists. Specifically, point of view video models removed the perspective taking component of video modeling, a component participants were observed to struggle with during scene type video modeling. Specifically, Nathan, Brendan, and Ethan all completed the scene type tasks upside down and then turned them around before handing the item to the researcher as a completed product. This pattern of behavior was observed on some trials when Nathan was completing the ice cream, and on each trial when Brendan and Ethan were completing the tangram tasks. As stated previously, researchers attempted to adapt the scene type video models, while maintaining the integrity of this type of video modeling, to decrease the amount of perspective taking participants were required to engage in (e.g., showing the completed product at the end of the video), but even with these adaptations, difficulty with the perspective taking component of the scene type video models was still observed.

Additionally, in Ethan's first pair of tasks, researchers attempted to teach Ethan's scene task using POV videos following Ethan's mastery of his POV task. When the modality was changed, Ethan made several comments reflecting that he noticed the changed perspective. For example, Ethan stated, "But, they're making it upside down!" and "But, I can't really see it this way." These comments seem to demonstrate difficulty with perspective taking and also seem to have demonstrated Ethan's difficulty with a change in teaching modality related to his difficulty with perspective taking. Taking these occurrences together, it may be beneficial to use point of view teaching modalities for children who struggle with perspective taking. It may also be recommended that point of view video models be used to teach more complex tasks, even for children who do not struggle with perspective taking, to remove this additional complexity of stimulus presentation during the teaching process. Additionally, it may be recommended that changing teaching perspectives be avoided when teaching a specific task, as Ethan's experience provides limited evidence that this may result in confusion and, as seen in all participants for who this teaching strategy was used, may not result in an increase in accuracy of task completion.

There are, however, some limitations to the use of point of view video models. First, the use of point of view video models lends itself more readily to certain skills than other. For example, point of view video modeling lends itself well to teaching children to write letters. Because the video would be a close up of another child's hand drawing or writing a letter, it is intuitive that point of view models, that give detailed information about how to form a letter, would be superior to scene models, which would only show

a child creating a letter from a distance, and possibly upside down. Likewise, scene modeling is likely to be more effective in teaching conversational skills than point of view video modeling. A point of view video model of a social exchange would show the face of the other participant in the conversation while the individual hears the individual whom they are expected to imitate speaking. The scene model, however, would show two people conversing, providing information about how far away it is appropriate to stand from another person as well as other, subtle social nuances. While Jake was successful in learning scripts from both the point of view and scene modalities, his attendance data indicate that Jake was not readily attending to the visual stimuli of the video, and instead, was likely attending only to the auditory stimuli. As a result, the component that we speculated would decrease the effectiveness of point of view video modeling was effectively removed. We also did not require Jake to approach the examiner and engage in the scripted conversation, so it is unknown if Jake would have approached the examiner and stood at an appropriate distance to engage in the conversation. This may be an important variable to measure, in addition to the participant's ability to recite the conversational script, in order to determine the differential effectiveness of scene type and point of view type video models.

When the attention data are examined, very little difference is observed between the two video modeling modalities. When compared statistically, no significant difference in attendance data were observed.

One limitation of this research is that the targets chosen were not part of the child's curriculum. That is, these targets were ancillary to the child's school curriculum,

and as such, may have differed from targets the child is familiar with learning in terms of level of difficulty or genre of skill. This may have increased the length of time children required to learn these skills. Future research may consider obtaining permission to view a copy of participants' Individualized Education Plan and choosing skills that are in line with the skills that participants' need to learn for school.

An additional limitation is the fact that multiple different types of supplemental teaching strategies were used to teach participants unmastered skills. Because these teaching strategies were not implemented as part of the experimental evaluation, supplemental teaching strategies were selected based on child characteristics, and therefore the selection of a supplemental strategy or the order of implementation of types of supplemental instruction was not consistent across participants. As such, it is not possible to compare acquisition in supplemental teaching conditions across participants and targets. Future research may more systematically implement supplemental teaching strategies, and compare these strategies to video modeling, to allow for conclusions to be made about whether supplemental teaching strategies are more effective teaching techniques than video modeling.

The fact that not all children learned all of their targeted skills is an additional limitation to this study. In those pairs in which both skills were not mastered, it is not possible to conclude whether point of view or scene type video modeling is a more effective teaching technique. In these skill pairs, both types of video modeling presentations were ineffective, and does not allow for conclusions to be made about differential effectiveness.

It was also found to be difficult to teach conditional discrimination tasks using video modeling. For Jake, it took over 40 trials in each condition for mastery to occur, and mastery still did not occur for the left, right, middle discrimination. In addition, it was difficult to create a video for Nathan that could effectively teach him to discriminate between multiple stimuli to identify which stimulus was different. It was not clear whether Nathan learned to identify the “different” stimulus or simply the stimulus that was modeled in the video. It may be the case that the use of video models to teach conditional discriminations may require additional examination in order to create videos that can effectively teach these more complicated skills. However, it is also possible that video models may be a less effective modality of teaching conditional discriminations than the empirically supported methodologies currently used with children with autism (Green, 2001).

There are still many questions that remain to be answered about what characteristics lead to the most effective video models. For example, our participants were able to complete their targeted tasks with more accuracy and more quickly when error correction was added as a component of training to in vivo modeling. A component analysis of the importance of error correction for acquisition, both used in conjunction with video modeling and in vivo modeling, may be an important area for future research. Also, as stated previously, the existing literature on video modeling incorporates a variety of model characteristics and teaching procedures. It would be beneficial to investigate each of these variations to determine whether these variations differentially impact the effectiveness of video modeling as a teaching tool. It is hoped

that this study will be the first in a series of studies that can serve to inform clinicians about best practices related to the use of video modeling to instruct young children with autism. Because video modeling is such a popular mode of teaching young children with autism, studies that will provide information on “best practices” are imperative.

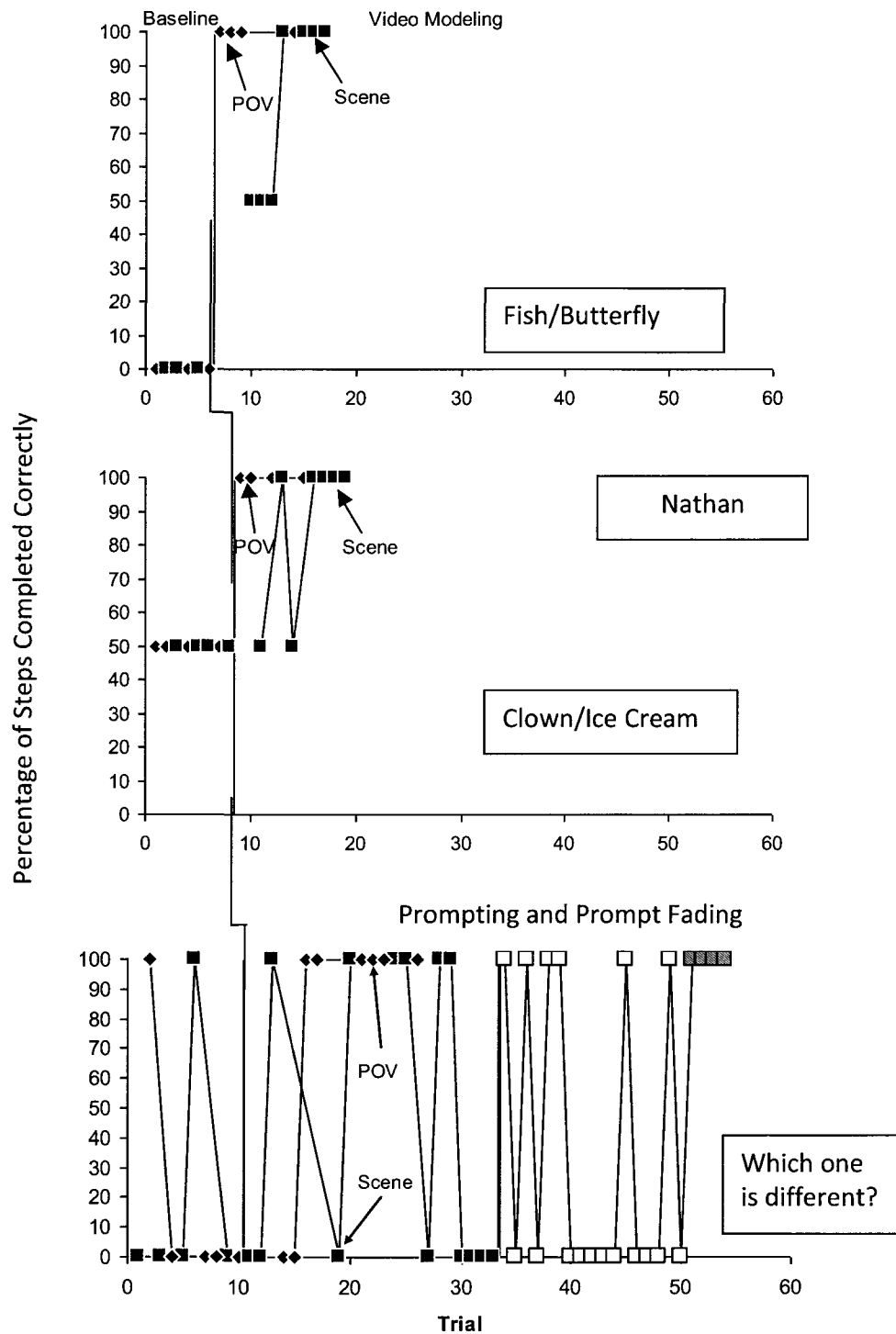


Figure 1. Nathan: This figure depicts the results of Nathan's learning of six target behaviors (two craft tasks and one conditional discrimination task). Shaded open data points in the supplemental teaching condition indicate independent responding.

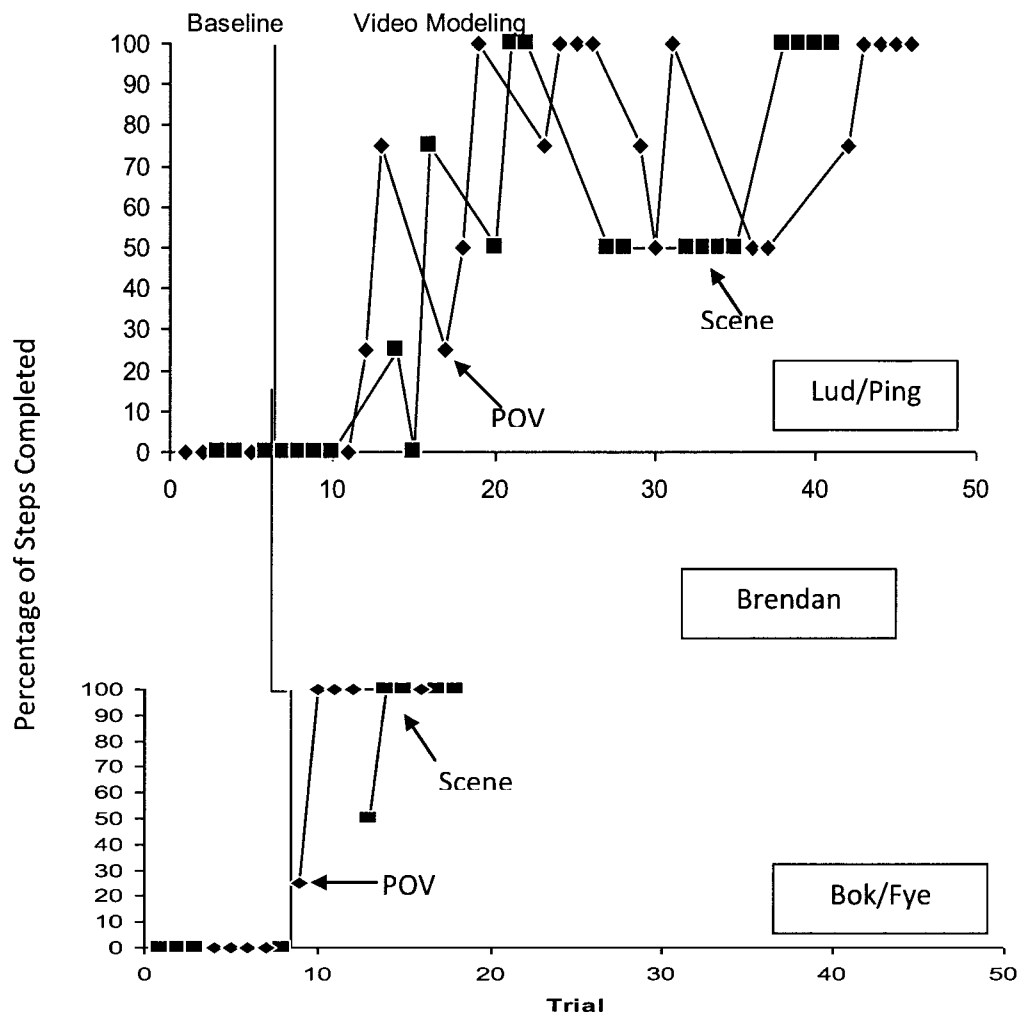


Figure 2. Brendan: This figure depicts the results of Brendan's learning of four target behaviors (construction tasks using tangrams).

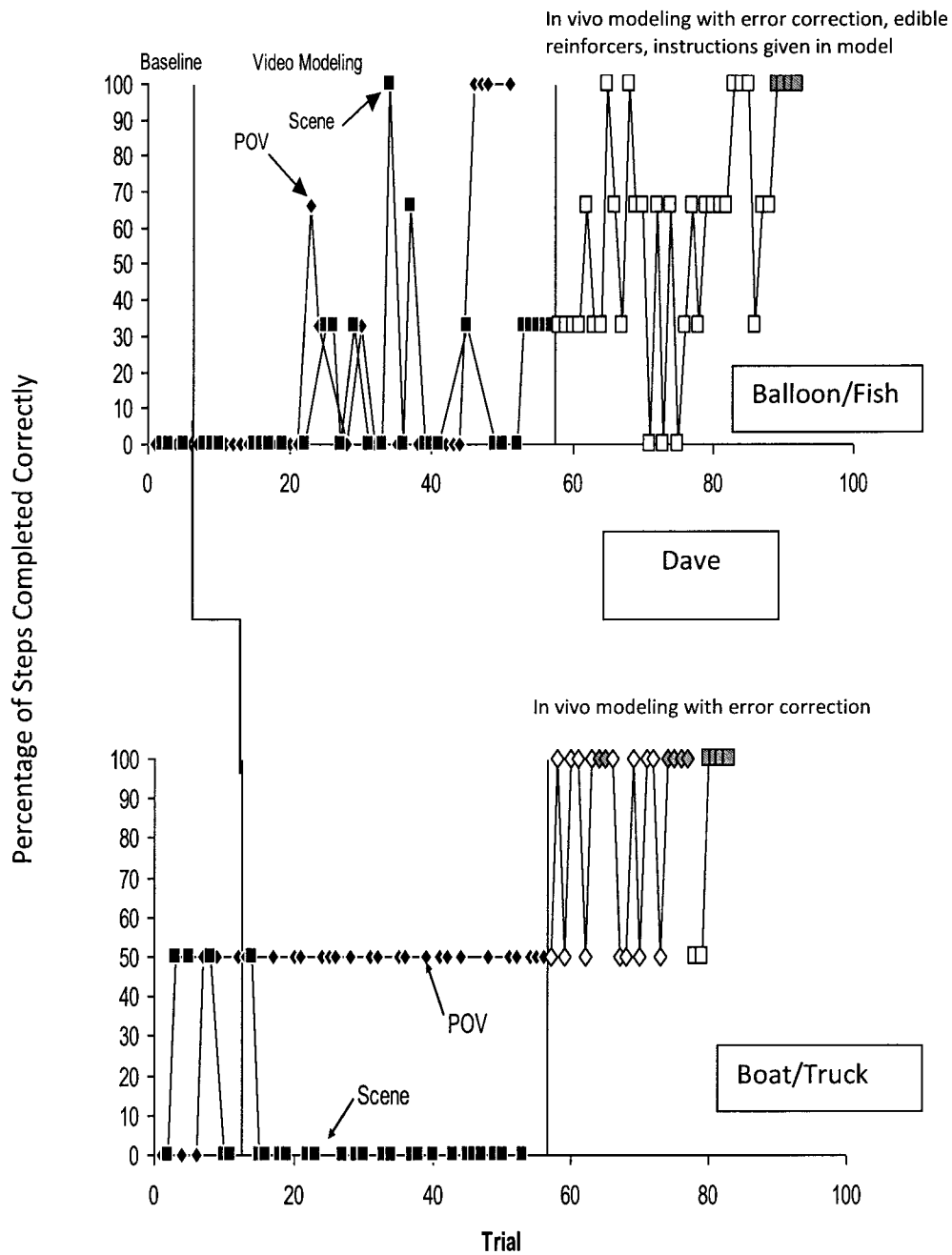


Figure 3. Dave: This figure depicts the results of Dave's learning of four target behaviors (two drawing tasks and two construction tasks). The shaded data points in the in vivo modeling of targets condition indicate independent responses.

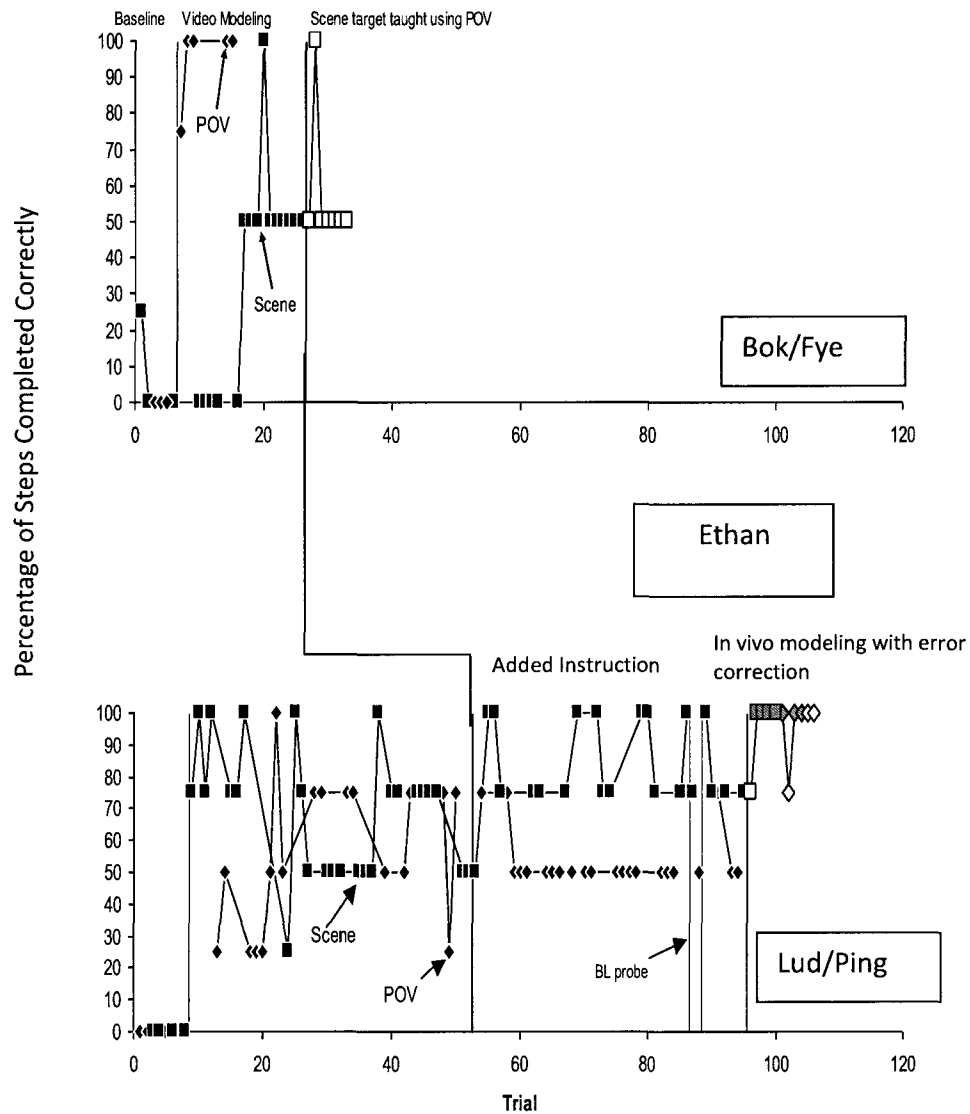


Figure 4. Ethan: This figure depicts the results of Ethan's learning of four target behaviors (four construction tasks using tangrams). Shaded data points in the in vivo modeling condition indicate independent responses.

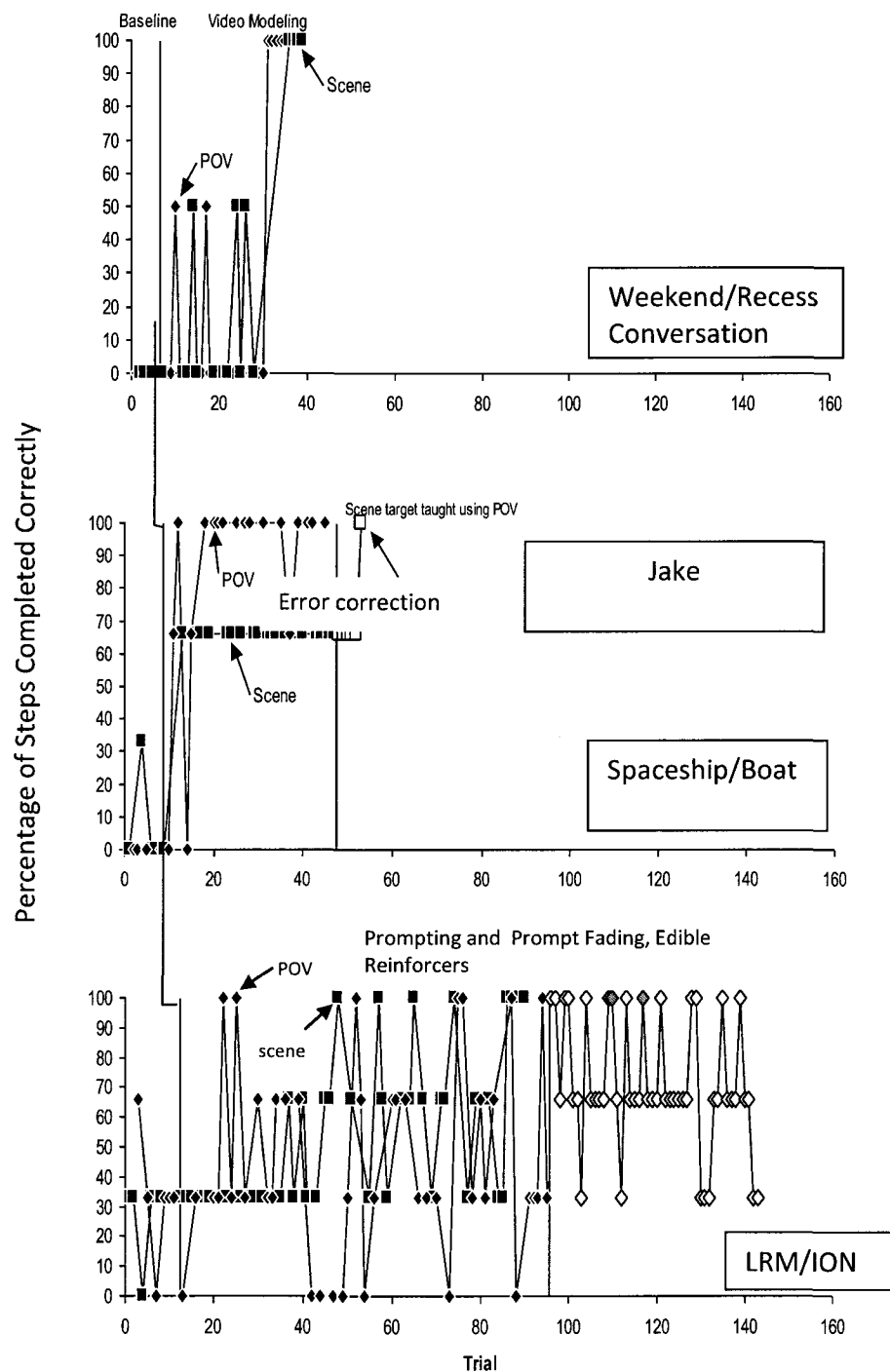


Figure 5. Jake: This figure shows the results of Jake's learning three pairs of target skills. Shaded data points in the prompting and edible reinforcement condition of panel three indicate independent responses.

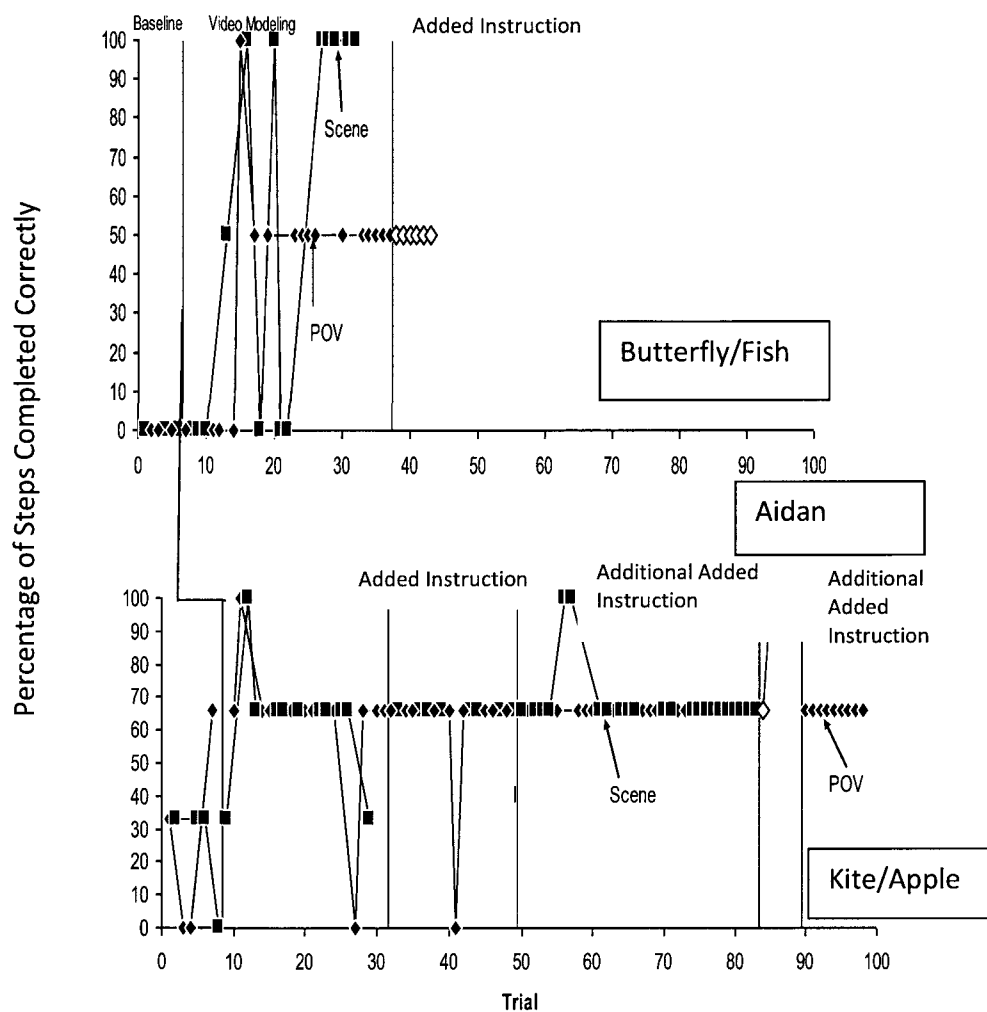
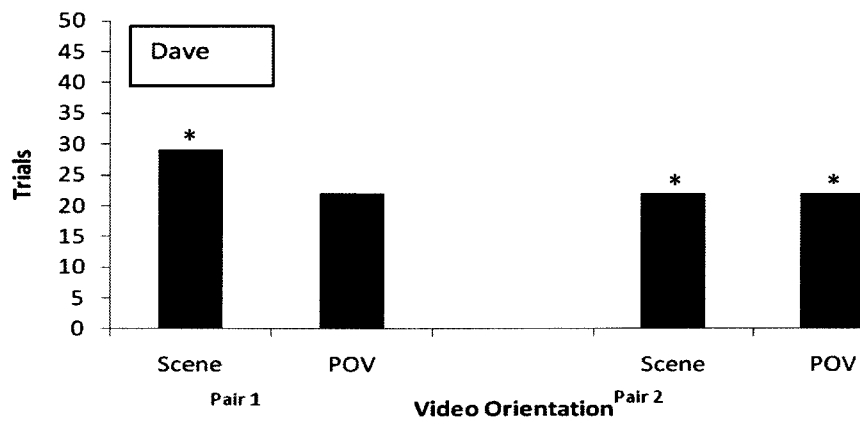
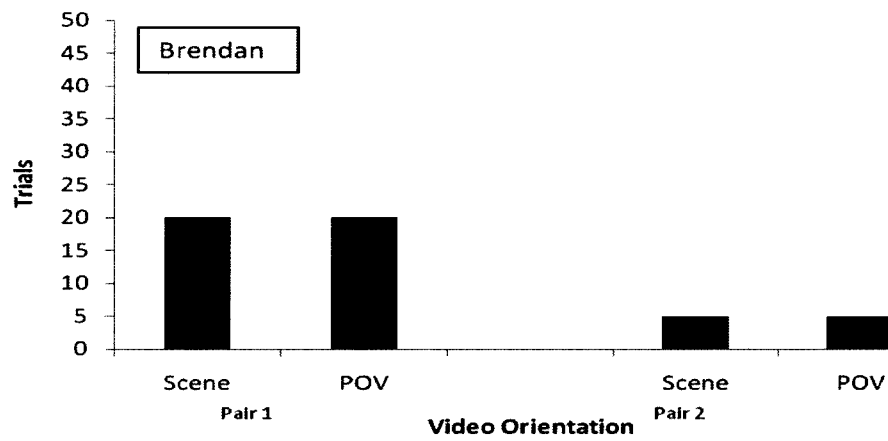
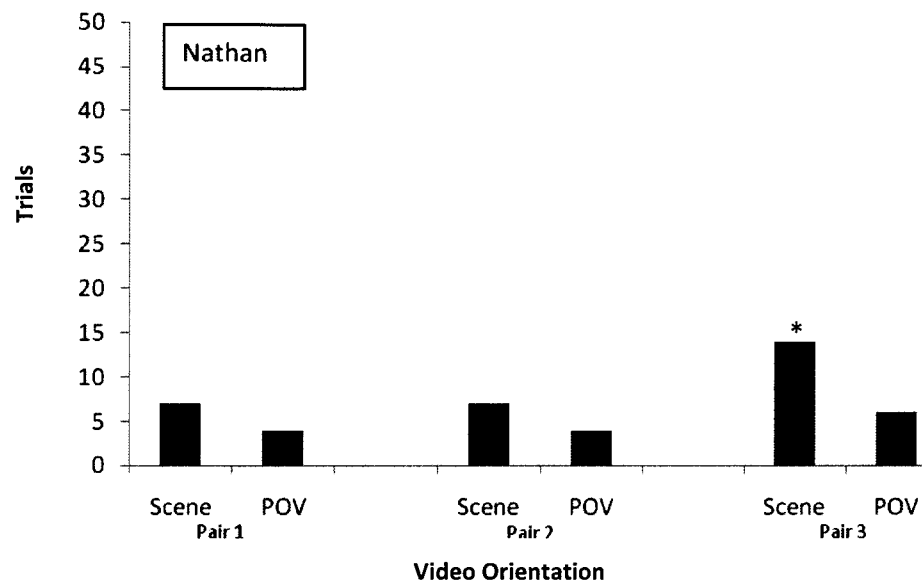


Figure 6. Aidan: This figure shows the results of Aidan learning two pairs of target skills (two craft tasks and two drawing tasks). This figure also shows the results of the pair of target skills that was discontinued due to multiple treatment interference.



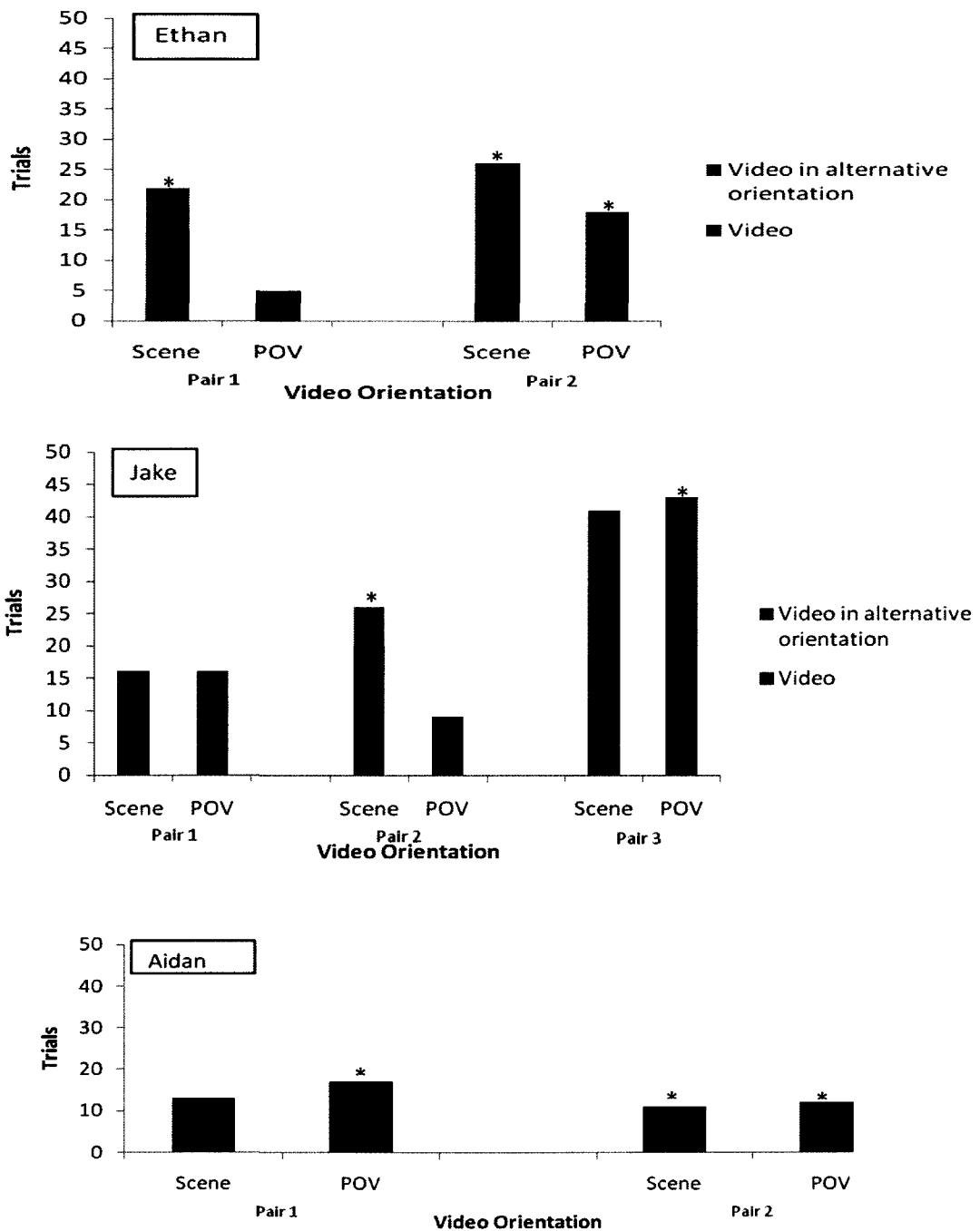


Figure 7. Trials to Criterion: This figure depicts trials to criterion for all targets taught in both camera angle perspectives. Asterisks above bars indicate those skills that participants did not master independently in video modeling (i.e., original comparison or unlearned target taught in alternative orientation). For unmastered targets, the total number of trials conducted with the participant for the unlearned target is presented in the graph.

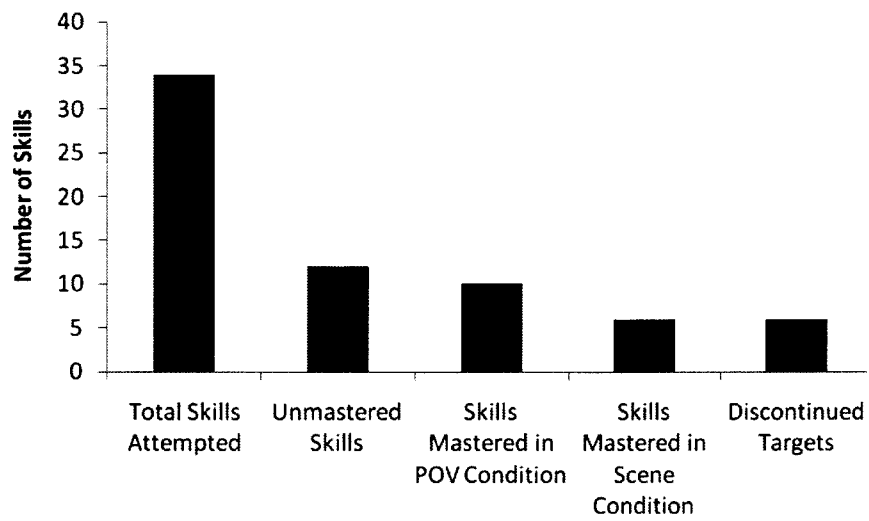


Figure 8. Targeted Skills: This figure depicts the total number of skills and subcategories of those skills with respect to performance outcomes and interventions.

Table 1

Descriptive Information About Video Models and Attending Behavior

Participant	Target	POV			Scene		
		Percent Trials Attending Data Collected	Average		Percent Trials Attending Data Collected	Average	
			Percentage Attending	Duration (s)*		Percentage Attending	Duration (s)*
Nathan	1	100.0%	82.9%	27.8	100.0%	84.9%	30.5
	2	100.0%	87.2%	29.4	100.0%	69.5%	24.3
	3	87.5%	70.2%	47.3	50.0%	74.1%	49.5
Brendan	1	35.0%	93.0%	29.1	25.0%	92.7%	33.5
	2	100.0%	95.5%	29.1	100.0%	97.9%	33.5
Dave	1	95.5%	52.4%	20.0	93.1%	58.9%	19.7
	2	95.5%	59.0%	25.0	95.5%	58.9%	23.0
Ethan	1	100.0%	97.6%	29.4	73.3%	70.2%	33.0
	2	43.6%	65.1%	28.6	45.7%	57.3%	33.8
Jake	1	81.3%	64.6%	15.1	58.8%	56.8%	12.9
	2	42.1%	88.4%	25.8	30.0%	92.7%	22.3
Aidan	3	30.2%	59.2%	23.1	31.7%	59.2%	30.2
	1	100.0%	81.8%	30.2	100.0%	81.0%	31.9
	2	67.5%	89.6%	22.5	59.1%	87.0%	23.5

*Duration was calculated per trial from the time the child was asked to orient to the

television to the end of the video. The number presented in the table is the average duration

across all trials.

Table 2

Number of Trials of Instruction

Participant	Task	Perspective	Number of Trials in		Number of Trials In Video Modeling with Instructions	Number of Trials in	
			Video Modeling	Modeling		Supplemental Instruction	
Nathan	Make a Butterfly	Scene	7		N/A		N/A
	Make a Fish	POV	4		N/A		N/A
	Make an Ice Cream	Scene	7		N/A		N/A
	Make a Clown	POV	4		N/A		N/A
	Which one is different (food)? Which one is different (animal)?	Scene	14		N/A		21 (e)
Brendan		POV	8		N/A		N/A
	Make a Lud	Scene	20		N/A		N/A
	Make a Ping	POV	20		N/A		N/A
	Make a Bok	Scene	5		N/A		N/A
	Make a Fye	POV	5		N/A		N/A
Dave	Draw a Fish	Scene	29		N/A		35 (b,d)
	Draw a Balloon	POV	22		N/A		N/A
	Make a Truck	Scene	22		N/A		21 (d)
	Make a Boat	POV	22		N/A		6 (d)
	Make a Bok	Scene	15		N/A		7 (a)
Ethan	Make a Fye	POV	5		N/A		N/A
	Make a Lud	Scene	26		20		5 (d)
	Make a Ping	POV	18		21		6 (d)
	Recess Script	Scene	17		N/A		N/A
	Weekends Script	POV	16		N/A		N/A
Jake	Draw a Boat	Scene	20		N/A		6 (a)
	Draw a Spaceship	POV	19		N/A		N/A

Table 2 - continued

Aidan	Inside/Next to/On Top	Scene	41	N/A	N/A
	Left/Right/Middle	POV	43	N/A	48 (b,e)
	Make a Butterfly	Scene	14	N/A	N/A
	Make a Fish	POV	17	7	N/A
	Draw a Kite	Scene	11	10	23 (c), 6 (d)
	Draw an Apple	POV	12	8	20 (c)

a: Video in other modality

d: In vivo modeling with error correction

b: Edible reinforcers

e: Prompting and prompt fading

c: Added Instruction

REFERENCES

- American Psychiatric Association. (2000). *Diagnostic and statistical manual of mental disorders* (4th ed. text revision). Washington, DC.
- Alberto, P.A., Cihak, D.F., & Gama, R.I. Use of static picture prompts versus video modeling during simulation instruction. *Research in Developmental Disabilities*, 26, 327-339.
- Apple, A. L., Billingsley, F., & Schwartz, I. S. (2005). Effects of video modeling alone and with self-management on compliment-giving behaviors of children with high-functioning ASD. *Journal of Positive Behavior Interventions*, 7, 33-46.
- Bandura, A. (8th Ed.). (1977). *Social learning theory*. Oxford, England: Prentice-Hall.
- Bandura, A. & Barab, P.G. (1971). Conditions governing nonreinforced imitation. *Developmental Psychology*, 5, 244-255.
- Bellini, S., & Akullian, J. (2007). A meta-analysis of video modeling and video self-modeling interventions for children and adolescents with autism spectrum disorders. *Exceptional Children*, 73, 264-287.
- Centers for Disease Control and Prevention (2007). CDC releases new data on autism spectrum disorders (ASDs) from multiple communities in the United States. Retrieved from <http://www.cdc.gov/od/oc/media/pressrel/2007/r070208.htm>.
- Charlop, M. H., & Milstein, J. P. (1989). Teaching autistic children conversational speech using video modeling. *Journal of Applied Behavior Analysis*, 22, 275-285.

- Charlop-Christy, M. H., Le, L., & Freeman, K. A. (2000). A comparison of video modeling with in vivo modeling for teaching children with autism. *Journal of Autism and Developmental Disorders, 30*, 537-552.
- Charlop-Christy, M. H., & Daneshvar, S. (2003). Using video modeling to teach perspective taking to children with autism. *Journal of Positive Behavior Interventions, 5*, 12-21.
- Cohen, H., Amerine-Dickens, M., & Smith, T. (2006). Early intensive behavioral treatment: Replication of the UCLA model in a community setting. *Journal of Developmental & Behavioral Pediatrics, 27*, S145-S155.
- Cox, M.V. (1978). The development of perspective-taking ability in children. *International Journal of Behavioral Development, 1*, 247-254.
- Delano, M. E. (2007). Video modeling interventions for individuals with autism. *Remedial and Special Education, 28*, 33-42.
- Dillon, C. & LeBlanc, L. (in press). Teaching children with autism to mand for information using video modeling. *Education and Treatment of Children*.
- Dillon, C., LeBlanc, L., & Geiger, K. (2009). *A review of procedural variations in conducting video modeling: What we know, what we think we know, and what we need to find out*. Paper presented at the Association for Behavior Analysis conference. Retrieved January 26, 2010, from www.abainternational.org.
- Eikeseth, S., Smith, T., Jahr, E., Eldevik, S. (2007). Outcome for children with autism who began intensive behavioral treatment between ages 4 and 7: A comparison controlled study. *Behavioral Modification, 31*, 264-278.

- Eldevik, S., Eikeseth, S., Jahr, E., & Smith, T. (2006). Effects of low-intensity behavioral treatment for children with autism and mental retardation. *Journal of Autism and Developmental Disorders*, 36, 211-224.
- Gilliam, J. (2006). GARS-2: Gilliam Autism Rating Scale - Second Edition. Austin, TX: PRO-ED.
- Goldsmith, L.A., & LeBlanc, L.A. (2004). Use of technology in interventions for children with autism. *Journal of Early and Intensive Behavior Intervention*, 1, 166-178.
- Green, G. (1996). *Early behavioral intervention for autism: What does research tell us?* Southborough, MA: New England Center for Autism.
- Green, G. (2001). Behavior analytic instruction for learners with autism: Advances in stimulus control technology. *Focus on Autism and Other Developmental Disabilities*, 16, 72-85.
- Hill, A., Bölte, S., Petrova, G., Beltcheva, D., Tacheva, S., & Poustka, F. (2001). Stability and interpersonal agreement of the interview-based diagnosis of autism. *Psychopathology*, 34, 187-191.
- Hine, J.F. & Wolery, M. (2006). Using point of view video modeling to teach play to preschoolers with autism. *Topics in Early Childhood Special Education*, 26, 83-93.
- Howard, J. S., Sparkman, C. R., Cohen, H. G., Green, G., & Stanislaw, H. (2005). A comparison of intensive behavior analytic and eclectic treatments for young children with autism. *Research in Developmental Disabilities*, 26, 359-383.
- Kanner, L. (1943). Autistic disturbances of affective contact. *Nervous Child*, 2, pp. 217-250.

- LeBlanc, L. A., Coates, A. M., Daneshvar, S., Charlop-Christy, M. H., Morris, C., & Lancaster, B. M. (2003). Using video modeling and reinforcement to teach perspective-taking skills to children with autism. *Journal of Applied Behavior Analysis, 36*, 253-257.
- Lord, C., Risi, S., Lambrecht, L., Cook, E. H., Leventhal, B. L., DiLavore, P. C., Pickles, A. & Rutter, M. (2000). The autism diagnostic observation schedule-generic: A standard measure of social and communication deficits associated with the spectrum of autism. *Journal of Autism and Developmental Disorders, 30*, 3, 205–223.
- Lovaas, O. I. (1987). Behavioral treatment and normal educational and intellectual functioning in young autistic children. *Journal of Consulting and Clinical Psychology, 55*, 3-9.
- Mazefsky, C.A., & Oswald, D.P. (2006). The discriminative ability and diagnostic utility of the *ADOS-G*, *ADI-R*, and *GARS* for children in a clinical setting. *Autism, 10*, 533-549.
- McCoy, K. & Hermansen, E. (2007). Video modeling for individuals with autism: A review of model types and effects. *Education and Treatment of Children, 30*, 183-213.
- Montgomery, J.M., Newton, B., Smith, C. (2008). Review of *GARS-2*: Gilliam autism rating scale-second edition. *Journal of Psychoeducational Assessment, 26*, 395-401.
- Nikopoulos, C. K., & Keenan, M. (2004). Effects of video modeling on social initiations by children with autism. *Journal of Applied Behavior Analysis, 37*, 93.

- Norman, J.M., Collins, B.C., & Schuster, J.W. (2001). Using an instructional package including video technology to teach self-help skills to elementary students with mental disabilities. *Journal of Special Education Technology, 16*, 5-16.
- Reichow, B. & Wolery, M. (2009). Comprehensive synthesis of early intensive behavioral interventions for young children with autism based on the UCLA Young Autism Project model. *Journal of Autism and Developmental Disorders, 39*, 23-41.
- Sallows, G. O., & Graupner, T. D. (2005). Intensive behavioral treatment for children with autism: Four-year outcome and predictors. *American Journal on Mental Retardation, 110*, 417-438.
- Shabani, D.B., Katz, R.C., Wilder, D.A., Beauchamp, K., Taylor, C.R., & Fischer, K.J. (2002). Increasing social initiation in children with autism: Effects of a tactile prompt. *Journal of Applied Behavior Analysis, 35*, 79-83.
- Sherer, M., Pierce, K.L., Paredes, S., Kisacky, K. L., Ingersoll, B., & Schreibman, L. (2001). Enhancing conversation skills in children with autism via video technology: Which is better, "Self" or "Other" as a model? *Behavior Modification, 25*, 140-158.
- Shipley-Benamou, R., Lutzker, J.R., & Taubman, M. (2002). Teaching daily living skills to children with autism through instructional video modeling. *Journal of Positive Behavior Interventions, 4*, 165-175.
- Sindelar, P.T., Rosenberg, M.S., Wilson, R.J. (1985). An adapted alternating treatments design for instructional research. *Education and Treatment of Children, 8*, 67-76.

- Smith, T., Eikeseth, S., Klevstrand, M., & Lovaas, O. I. (1997). Intensive behavioral treatment for preschoolers with severe mental retardations and pervasive developmental disorder. *American Journal on Mental Retardation*, 102, 238-249.
- Spradlin, J.E., & Brady, N.C. (1999) Early childhood autism and stimulus control. In P.M. Ghezzi, W.L. Williams, & J.E. Carr (Eds.), *Autism: Behavior analytic perspectives* (pp.49-65). Reno, NV: Context Press.
- Sturmey, P. (2003). Video technology and persons with autism and other developmental disabilities: An emerging technology for PBS. *Journal of Positive Behavior Interventions*, 5, 3-4.
- Volkmar, F., Chawarska, K., & Klin, A. (2005). Autism in infancy and early childhood. *Annual Review of Psychology*, 56, 315-336.
- Yeargin-Allsopp, M., Rice, C., Karapurkar, T., Doernberg, N., Boyle, C., & Murphy, C. (2003). Prevalence of autism in a US metropolitan area. *Journal of the American Medical Association* 289, 49-55.

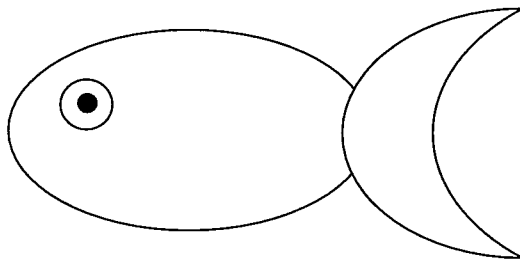
Appendix A

Craft Tasks

The following are illustrations of each of the craft tasks identified as targets for various participants. The crafts are identified by the type of video modeling used to teach the craft task and the participant or participant(s) who were taught to create each item.

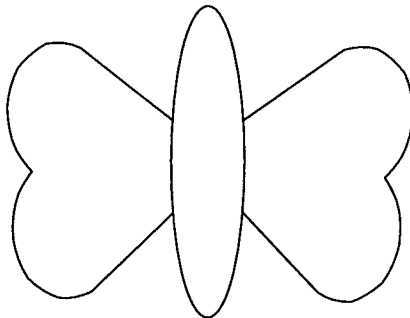
Fish (POV):

Taught to: Nathan, Aidan



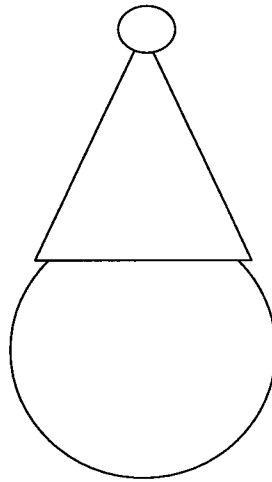
Butterfly (scene):

Taught to: Nathan, Aidan



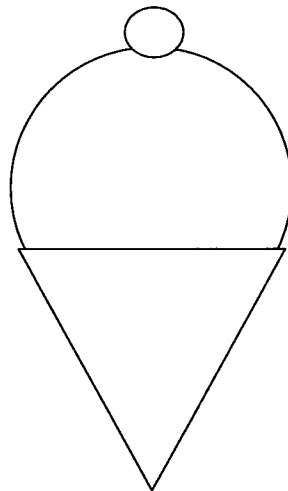
Clown (POV):

Taught to Nathan



Ice Cream (scene):

Taught to Nathan

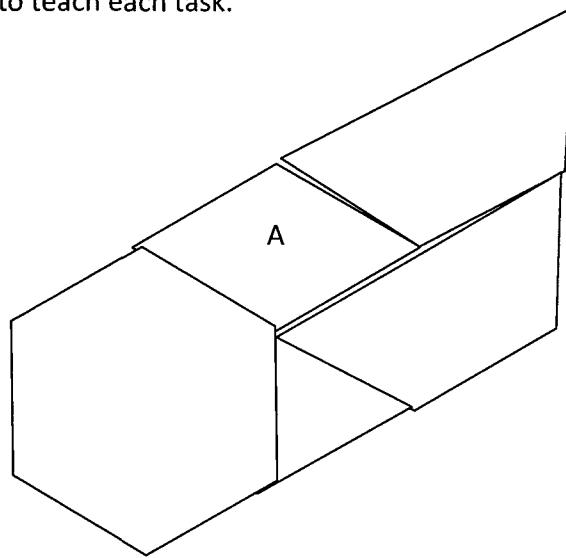


Appendix B

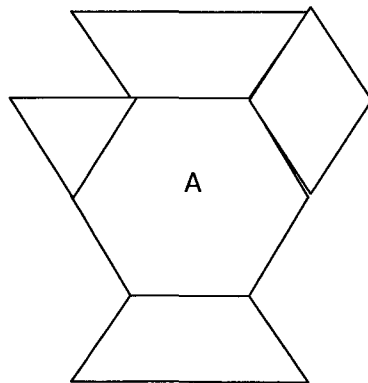
Tangram Tasks

The following are illustrations of the tangram tasks taught to Brendan and Ethan. The anchor piece (the piece placed in the middle of the magnetic board to start the task) is labeled with an “A.” The tangram shapes are also labeled with the type of video modeling used to teach each task.

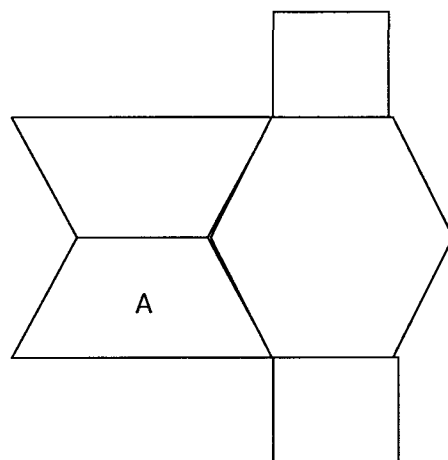
Ping (POV):



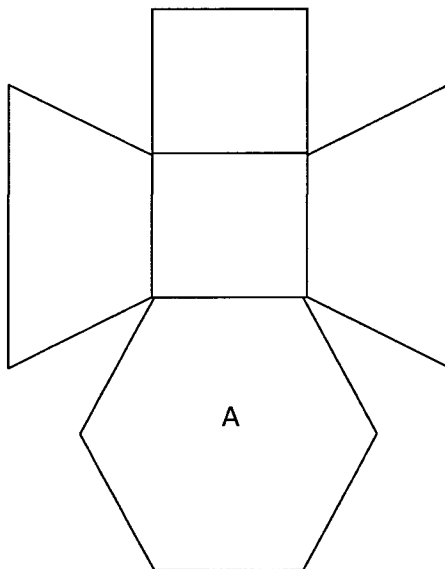
Lud (scene):



Fye (POV):



Bok (scene):



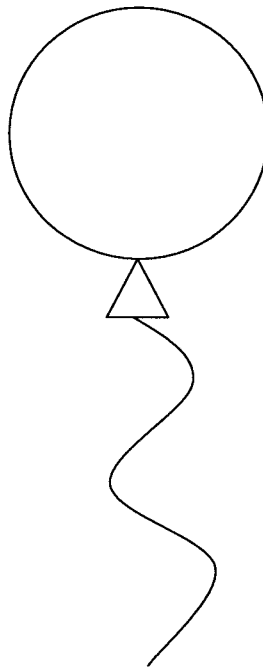
Appendix C

Drawing Tasks

The tasks in this appendix were drawing tasks learned by a variety of participants. Participants had to draw each of the shapes, though not in a specific order, to be considered correct. The illustrations are labeled below with the type of video modeling that was used to teach each task and also the participant who was taught the target skill.

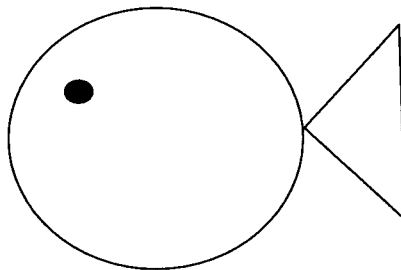
Balloon (POV):

Taught to Dave



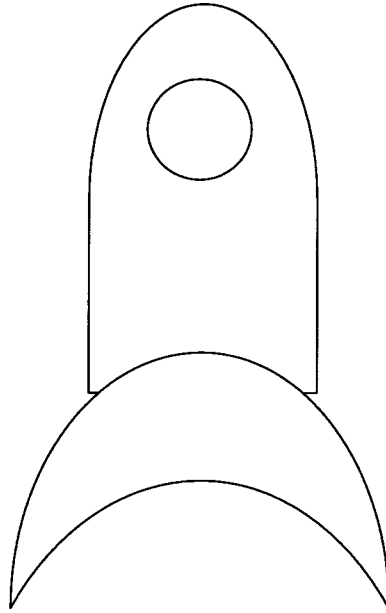
Fish (scene):

Taught to Dave



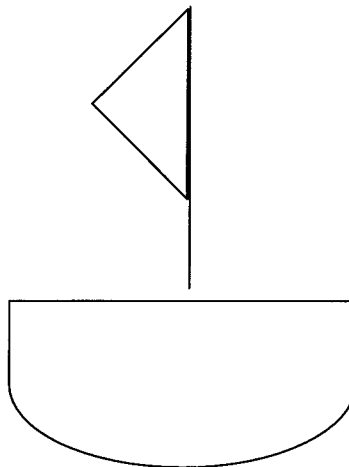
Spaceship (POV):

Taught to Jake



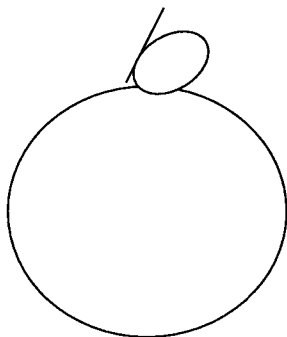
Boat (scene):

Taught to Jake



Apple (POV):

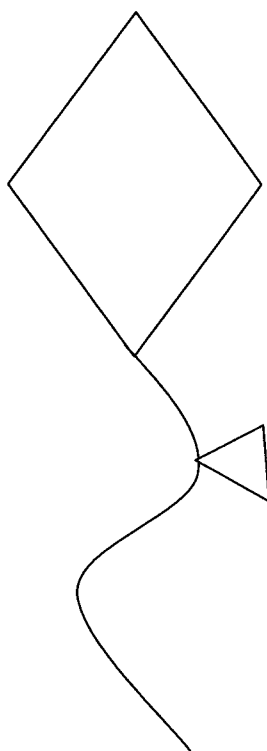
Taught to Aidan



Kite:

Taught to Aidan

,



Appendix D

Conversation Scripts

The following is the conversation script that was taught to Jake. Each script is labeled with the type of video modeling used to teach each script.

Weekend (POV):

Researcher: What do you like to do on weekends?

Jake: Go to Bounceland.

Researcher: I like to play with my spaceships.

Jake: That's cool.

Recess (scene):

Researcher: What do you like to do at recess?

Jake: Play on swings.

Researcher: I like to jump rope.

Jake: That's cool.

Appendix E

Sample Data Sheet for Child Data and Procedural Integrity – Baseline

Trial	Child placed one red block on top of the yellow	Child placed one red block below the yellow	Child placed the green block to the left of the yellow	Child placed blue block to the right of the yellow
1	Y / N	Y / N	Y / N	Y / N
2	Y / N	Y / N	Y / N	Y / N
3	Y / N	Y / N	Y / N	Y / N
4	Y / N	Y / N	Y / N	Y / N
5	Y / N	Y / N	Y / N	Y / N
6	Y / N	Y / N	Y / N	Y / N
7	Y / N	Y / N	Y / N	Y / N
8	Y / N	Y / N	Y / N	Y / N
9	Y / N	Y / N	Y / N	Y / N
10	Y / N	Y / N	Y / N	Y / N

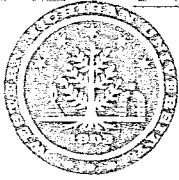
Trial	Researcher gave child appropriate materials	Researcher gave appropriate instruction	Researcher paused 3s to allow for responding	Researcher praised correct responding
1	Y / N	Y / N	Y / N	Y / N
2	Y / N	Y / N	Y / N	Y / N
3	Y / N	Y / N	Y / N	Y / N
4	Y / N	Y / N	Y / N	Y / N
5	Y / N	Y / N	Y / N	Y / N
6	Y / N	Y / N	Y / N	Y / N
7	Y / N	Y / N	Y / N	Y / N
8	Y / N	Y / N	Y / N	Y / N
9	Y / N	Y / N	Y / N	Y / N
10	Y / N	Y / N	Y / N	Y / N

Appendix F

Sample Data Sheet for Child Data and Procedural Integrity - Treatment

Trial	Child placed one red block on top of the yellow	Child placed one red block below the yellow	Child placed the green block to the left of the yellow	Child placed blue block to the right of the yellow
1	Y / N	Y / N	Y / N	Y / N
2	Y / N	Y / N	Y / N	Y / N
3	Y / N	Y / N	Y / N	Y / N
4	Y / N	Y / N	Y / N	Y / N
5	Y / N	Y / N	Y / N	Y / N
6	Y / N	Y / N	Y / N	Y / N
7	Y / N	Y / N	Y / N	Y / N
8	Y / N	Y / N	Y / N	Y / N
9	Y / N	Y / N	Y / N	Y / N
10	Y / N	Y / N	Y / N	Y / N

Trial	Researcher showed child video	Researcher said "Watch This"	Researcher gave child appropriate materials	Researcher gave appropriate instruction	Researcher paused 3s to allow for responding	Researcher praised correct responding
1	Y / N	Y / N	Y / N	Y / N	Y / N	Y / N
2	Y / N	Y / N	Y / N	Y / N	Y / N	Y / N
3	Y / N	Y / N	Y / N	Y / N	Y / N	Y / N
4	Y / N	Y / N	Y / N	Y / N	Y / N	Y / N
5	Y / N	Y / N	Y / N	Y / N	Y / N	Y / N
6	Y / N	Y / N	Y / N	Y / N	Y / N	Y / N
7	Y / N	Y / N	Y / N	Y / N	Y / N	Y / N
8	Y / N	Y / N	Y / N	Y / N	Y / N	Y / N
9	Y / N	Y / N	Y / N	Y / N	Y / N	Y / N
10	Y / N	Y / N	Y / N	Y / N	Y / N	Y / N



Appendix G

HSIRB Approval Letter

Date: December 18, 2007

To: Linda LeBlanc, Principal Investigator
Courtney Dillon, Student Investigator for dissertation
Kaneen Geiger, Student Investigator for thesis

From: Amy Naugle, Ph.D., Chair

A handwritten signature in black ink, reading "Amy Naugle", written over the word "Chair".

Re: HSIRB Project Number: 07-11-01

This letter will serve as confirmation that your research project titled "Evaluating the Effectiveness of Video Modeling for Children with Autism: Preference and Point of View vs. Scene Modeling" has been **approved** under the **full** category of review by the Human Subjects Institutional Review Board. The conditions and duration of this approval are specified in the Policies of Western Michigan University. You may now begin to implement the research as described in the application.

Please note that you may **only** conduct this research exactly in the form it was approved. You must seek specific board approval for any changes in this project. You must also seek reapproval if the project extends beyond the termination date noted below. In addition if there are any unanticipated adverse reactions or unanticipated events associated with the conduct of this research, you should immediately suspend the project and contact the Chair of the HSIRB for consultation.

The Board wishes you success in the pursuit of your research goals.

Approval Termination: November 21, 2008