



LUND UNIVERSITY

From sign to action : studies in chimpanzee pictorial competence.

Studies in chimpanzee pictorial competence

Sonesson, Göran; Hribar, Alenka; Call, Joseph

Published in:
Semiotica

DOI:
[10.1515/sem-2013-0108](https://doi.org/10.1515/sem-2013-0108)

2014

[Link to publication](#)

Citation for published version (APA):

Sonesson, G., Hribar, A., & Call, J. (2014). From sign to action : studies in chimpanzee pictorial competence. *Studies in chimpanzee pictorial competence. Semiotica*, 198, 205-240. <https://doi.org/10.1515/sem-2013-0108>

Total number of authors:
3

General rights

Unless other specific re-use rights are stated the following general rights apply:

Copyright and moral rights for the publications made accessible in the public portal are retained by the authors and/or other copyright owners and it is a condition of accessing publications that users recognise and abide by the legal requirements associated with these rights.

- Users may download and print one copy of any publication from the public portal for the purpose of private study or research.
- You may not further distribute the material or use it for any profit-making activity or commercial gain
- You may freely distribute the URL identifying the publication in the public portal

Read more about Creative commons licenses: <https://creativecommons.org/licenses/>

Take down policy

If you believe that this document breaches copyright please contact us providing details, and we will remove access to the work immediately and investigate your claim.

LUND UNIVERSITY

PO Box 117
221 00 Lund
+46 46-222 00 00

Alenka Hribar, Göran Sonesson* and Josep Call

From sign to action: Studies in chimpanzee pictorial competence

Abstract: Many studies of children and apes realized in psychology address issues that are highly relevant to semiotics, but they often do so indirectly, or they use a terminology that is confusing and/or too vague from a semiotical point of view. The studies reported here, however, follow the paradigm of these psychological studies, but they are couched in an explicit semiotical terminology. They involve three classical semiotical issues: the nature of the sign, as opposed to other meanings; degrees and/or types of iconicity and their relevance for understanding; and the importance of temporal focus in different kinds of semiotical resources. The studies all involve one subject, the chimpanzee Alex, and all issues were studied looking at the actions accomplished by the subject after being exposed to different semiotic resources.

Keywords: sign; picture; iconicity; imitation; video; animal cognition

***Corresponding author: Göran Sonesson:** Lund University. E-mail: goran.sonesson@semiotik.lu.se

Alenka Hribar: Max Planck Institute for Evolutionary Anthropology. E-mail: hribar@eva.mpg.de

Josep Call: Max Planck Institute for Evolutionary Anthropology. E-mail: call@eva.mpg.de

Semiotics, whether Peircean, Saussurean or in-between, has tended to be synchronic, or rather achronic. But if we want to understand the difference between the semiosis occurring in human beings and in other animals, we have to turn to diachronic studies in the widest sense, which includes evolution and child development. Given a more specific concept of sign than that mostly used in semiotics, we can ask whether other primates are able to use signs; and we can investigate the impact of different kinds of iconicity, as well of the temporal differentiations available in different semiotic resources. In the following, we are going to investigate the availability of the concept of sign to higher primates, as well as, more specifically, what difference different kinds of iconicity, and notably those involving temporal division, make to the result.¹

¹ Research for this article was partly financed within the framework of the European Community Research grant “Stages in the Evolution and Development of Sign Use” (2004–2008). Sonesson worked on the final version of this paper while being employed by the Centre for Cognitive

1 Introduction

Psychological experiments have rarely been used within semiotics proper, with the notable exceptions of the early work involving pictorial semiotics by Lindens (1976) and Krampen (1983). Nevertheless, it seems to us that experimental studies are particularly apt to elucidate the fundamental issues of semiotics, in particular in relation to the evolution and development of signs and other meanings. In recent decades, a number of psychologists have addressed semiotic issues with the help of experiments. Two groups have made important contributions to the field: on the one hand, Judy DeLoache and her collaborators, who study, notably, the capacity of children for understanding how to retrieve a hidden object corresponding to a picture or a scale model (DeLoache 2000; DeLoache and Burns 1994), a set-up that was later replicated with apes (Kuhlmeier and Boysen 2002); on the other hand, the work accomplished by Michael Tomasello (1999, 2008) and collaborators, which is dedicated to the emergence of meaning in both children and apes on a much wider scale. From the point of view of semioticians, of whichever conviction, the terminology in these studies seems seriously misleading, and the concepts offered for study appear to be insufficiently analyzed. But these are no doubt the pioneering contributions to experimental semiotics. Nevertheless, an explicit semiotical framework has so far been used only by Persson (2008) and Lenninger (2009).

1.1 The semiotic framework

A primary difficulty consists in the difference of terminology between psychology and semiotics. Many psychologists, like DeLoache (1995: 67), claim that an “entity that someone intends to stand for something other than itself” is a “symbol.” In DeLoache’s own work, “symbols” in this sense are exemplified by pictures, videos, and scale-models. In this article, we will follow the practice in semiotics of using “sign” as the general term, and reserving “symbol” for signs that are highly conventionalized or otherwise rule-bound. In this sense, pictures, videos, and scale-models are primarily iconic, although they may of course contain sym-

Semiotics, financed by the Bank of Sweden Tercentenary Foundation. Hribar and Call were all the time employed by Wolfgang Köhler Primate Research Centre in Leipzig. We want to thank Sonesson’s colleagues at CCS, notably Chris Sinha, Elainie Madsen, Tomas Persson, Joel Pathermore, and Jordan Zlatev, for perspicuous comments on an earlier draft of this paper. We also want to thank the editor of this volume, as well as an anonymous reviewer, for stimulating remarks.

1 bolic (as well as indexical) features. Indeed, we will take it for granted that all, or
 2 most, signs contain iconic, indexical, and symbolic aspects, with one of these
 3 being normally more prominent, or dominant, in the Prague school sense of orga-
 4 nizing the other aspects for their own purpose (such as indexicality in the pre-
 5 dominantly iconic photograph; cf. Sonesson 1994).

6 However, we will need a more explicit definition of the notion of sign, which
 7 takes into account our intuitive understanding of what a sign is, exemplified in
 8 the linguistic sign, but readily generalizable to pictures, videos, and scale-
 9 models, and the like. For such a purpose, the definition must be considerably
 10 more specific than the one ordinarily employed in semiotic theory, notably by
 11 Peirce and his followers, for which all meaningful relations are signs, but which
 12 is at the same time much more general than the Saussurean notion, which tends
 13 to restrict the notion of sign to language and some other systems which are in
 14 some way similar to language. At the same time, it needs to be more specific than
 15 both the Saussurean and the Peircean sign concepts in that it clearly defines the
 16 requirements for two objects being called expression and content, while this is
 17 never done in the work of Peirce, and only by example by Saussure. Such a notion
 18 of sign was first formulated by Sonesson (1989, 1992, 2007, 2009), taking his in-
 19 spiration from the Piagetian idea of differentiation, and Husserl's definition of
 20 appresentation.

21 As we will see, the *sign*, in this sense, is a kind of meaning, but not all mean-
 22 ings are signs. Perception is clearly meaningful to animals and human infants
 23 alike (though not necessarily in the same way), but the capacity for *sign use* is a
 24 much more exclusive property. It is precisely because the capacity for using signs
 25 may be expected to distinguish later stages in evolution and development, that it
 26 is important to separate signs from other meanings.

27 The sign is a meaning that is made up of two parts, traditionally known as
 28 *expression* and *content*. That the sign consists of two parts implies that the parts
 29 are separated. In Piaget's (1967 [1945], 1967, 1970) terms, they are "differentiated
 30 from the point of view of the subject." Contrary to what Piaget suggests, a thing
 31 that is immediately continuous to another or which is a part of another in the
 32 common sense world may very well be differentiated within the sign (cf. Sonesson
 33 1989, 2007). In a parallel fashion, things that are similar to each other can be dif-
 34 ferentiated within the sign. Thus, there can be indexical (contiguity-based) and
 35 iconic (similarity-based), as well as symbolic (convention-based) signs.² If I see a
 36

37 ² For readers with a semiotic background, we would like to point out that we are well aware of
 38 this not being the Peircean definition of symbols, but it seems to us more coherent within a three-
 39 some of signs. Similarly, expression, content, and referent are of course not the members of the
 40 Peircean triad.

branch sticking up over the house and conclude that there is a tree behind the house, this is a mere indexicality; but the marks on the ground left by the animal are indexical signs, clearly separated from the (part of) the animals having produced them.

Indeed, a further differentiation may have to be made for certain purposes. The marks on the ground tell me “an elk was here before,” and this is a fact distinct from the marks, as well as from the elk, which is now somewhere else. Similarly, the color configurations making up the photographs that accompany this article are distinct from the perceptual impressions of Alenka Hribar, but even they are here now wherever the reader is, while Alenka herself is probably still back in Leipzig. This is why we really have to separate three parts of the sign, *expression*, *content*, and *referent*, where content is the standpoint taken on the referent by the sign user, as codified in some semiotic resource.

Differentiation, however, is not a sufficient criterion. Each time we categorize two different things as belonging together (as opposed to categorical perception in which they are seen as instances of the same type), the items brought together must first be told apart (“differentiated”). Still, categorization as such is not a form of sign use. The sign is a whole made up of several parts, and therefore, there is necessarily some relationship between these parts. There is a *double asymmetric relationship* between expression and content. First, from the point of view of immediacy, expression is more accessible to consciousness than content. In the second place, content is more in focus (more prominent, more important) than expression. The founder of phenomenology Edmund Husserl (1939) formulated the definition of the sign more or less in these terms. Nevertheless, this does not preclude other relations between expression and content being symmetric. It is common to suppose a substitutive relationship, which is a symmetric relation, between expression and content, but this may be misleading, since expressions are rarely used for the same purpose and in the same context as their contents

Bates (1979: 43) has also suggested that the sign (our expression) and its referent (i.e., the content) must be conceived as being similar and yet separate for a sign relationship to obtain. Daddesio (1995: 117) observes that there are three possible cases: the organism fails to grasp any relation at all between two items; it reacts in the same way to both items; or, finally, the organism recognizes the two items as distinct but related. Daddesio’s second case is that of categorization, which is important for perception. However, it would be wrong to conclude from the fact that an individual treats the two items as being distinct, that the particular relationship between the items is necessarily one of sign function (cf. Sonesson 2004, 2009). Given a prototype conception of categories, *a* and *b* may be treated as different just because they are differently central to the category of which they are perceived to form a part. Or they may be attended to differently,

1 merely because one contains more, and more interesting perceptual properties,
2 than the other.

3 The problem of separating the expression and the content becomes particu-
4 larly acute in the case of an iconic sign, in which, by definition, expression and
5 content must share at least some properties. Persson (2008: 10–15) has distin-
6 guished three modes of attending to pictures: *surface mode*, in which only “pat-
7 terns, shapes, and colors, on the surface of the picture” are perceived (a picture of
8 an apple is seen as patches of red and yellow); *reality mode*, in which the picture
9 is seen as part of reality, instead of being about reality (the picture of an apple is
10 seen as an apple); and, finally, *pictorial mode*, which involves both “an *expecta-*
11 *tion of separation*” and “an *expectation of likeness*” (the surface covered with
12 patterns, shapes, and colors is seen as being *about* an apple).³ *Mutatis mutandis*,
13 the case of non-iconic signs is the same, though instead of an “expectation of like-
14 ness,” there would be a more general “expectation of aboutness.” It would be
15 natural to think, however, that the expectation of separation (or rather: differen-
16 tiation) cohabits more uneasily with an expectation of similarity than with the
17 mere expectation of aboutness.

18

19

20 1.2 Review of research relevant to the sign function

21

22 In order to understand signs such as pictures, videos or models, one must under-
23 stand the duality of the sign artefact, i.e., that pictures and videos are 2-D objects
24 in themselves as well as expression of something else, which are usually 3D
25 objects. This involves experiencing both the similarity and the difference between
26 the picture and the object depicted, and grasping the asymmetric relationship
27 between them.

28 According to DeLoache (2004), children gradually learn to understand this
29 duality. Children as young as 5 months old look longer at a doll than at its picture
30 (DeLoache and Burns 1994). However, at 9-months of age children manually ex-
31 plore pictures and images of still and moving objects on a television screen as if
32 they were real objects, i.e., they grasp, pat, and rub them. But if they are pre-
33 sented simultaneously with a real object and with its picture, they preferentially
34 pick a real object over the corresponding depiction (DeLoache et al. 1998; Pier-
35 routsakos and DeLoache 2003; Pierroutsakos and Troseth 2003).

36

37

38 ³ As Persson remarks, Fagot et al. (2000) independently made a similar distinction between
39 picture processing in terms of “independence,” “confusion,” and “equivalence.” These terms,
40 however, appear to be less clear.

Likewise, apes and monkeys have been shown to demonstrate an ability to discriminate between real objects and the corresponding pictures (Parron et al. 2008; Imura and Tomonaga 2003). When picture-naïve baboons and chimpanzees were presented with a real banana piece and the picture thereof, they preferred the real banana. The gorillas did not show this preference. When they were presented with a choice between a picture of a banana and a picture of a pebble, the chimpanzees almost uniformly chose the banana picture. The results for the gorillas were less clear-cut. Some baboons and gorillas even ate the picture, whereas the chimpanzees did not (Parron et al. 2008). These results suggest that the gorillas and at least some baboons did not see the pictures as signs of bananas. Although the chimpanzees did not mistake the picture of a banana for a real banana and ate it, it is still unclear whether they processed the pictures as signs referring to bananas. Small children also try to grasp and even eat pictures. Similarly, young children will imitate actions seen on a television screen, but their imitation level following a live demonstration of the same action is higher (Barr and Hayne 1999; Hayne et al. 2003; Meltzoff 1988). This shows that the picture and its referent are seen as different, not necessarily that they are seen as sign and referent. There may be other explanations; one could speculate that the real doll and the real banana are seen as more prototypical instances of their respective categories; alternatively, that they may simply be more interesting because of having more perceptual predicates (Sonesson 2009). Doves have been known to react differently to a picture and its referent (Cabe 1980).

Even though children start to react differently to a real object and the corresponding picture at around the age of 1 year, they require about another 1.5 years to develop the ability to use pictures and videos as a source of information guiding their search behavior in a real setting. In these studies, children are for instance shown on a video how a toy is being hidden under a chair, and then they have to find this toy in the real room (DeLoache and Burns 1994; Schmitt and Anderson 2002; Troseth 2003a; Troseth and DeLoache 1998). Even 2-year-old children were able to use the video presentation of a hiding event to find the toy after gaining some experience with television (Troseth 2003b). This suggests that children may be able to understand the sign function of photos and videos at an earlier age, if they have had a lot of experience with the relevant medium.

Interpreting pictures and videos appears to be surprisingly difficult: experiments by DeLoache and her collaborators (e.g., DeLoache and Burns 1994) suggest that pictures are understood later than language (around 2.5 years). The problem may be that iconicity gets in the way of the sign function. This interpretation is consistent with another of DeLoache's findings, according to which scale models are even more difficult to understand than pictures. Children begin to understand the sign function of the scale model at around the age of 3 years

(DeLoache 2000; DeLoache et al. 1991). However, 3-year-olds still fail to perceive the dual nature of the model, if its salience as an object is increased (DeLoache 2000). It should be noted that the task set by DeLoache involves more than the recognition of the picture as a picture – it requires an action: fetching the hidden object. Attempts to repeat the task, without the element of hiding, however, does not change the results fundamentally (cf. Lenninger 2009). Without verbal scaffolding, pictures are understood even later, according to Callaghan (2000) and Callaghan and Rankin (2002). Other facilitating elements, not thematized by DeLoache, are various kinds of indexical scaffolding used in the experiments, involving pointing as well as creating neighborhood relations between the picture and the depicted object.

Kuhlmeier et al. (1999) presented chimpanzees with a hiding task involving four possible hiding places, similar to what DeLoache and Burns (1994) used with children. Two chimpanzees were shown a photograph of either the furniture where the reward was hidden (e.g., a chair), of the whole room with the hiding place marked, or of all four hiding places the correct one being pointed out. Under these circumstances, however, the older chimpanzee was reliably able to find the reward in the real room after she had seen the hiding place in the photos, but the younger one failed.⁴ Kuhlmeier et al. (1999) and Kuhlmeier and Boysen (2001, 2002) also showed that chimpanzees were able to use the information they were given by a scale model (i.e., color, shape or position of the hiding place) to find the hidden reward in a real enclosure. Their performance level was higher when object cues were present (color and shape) than when only spatial ones were offered (Kuhlmeier and Boysen 2002; Poss and Rochat 2003).

1.3 Theory and research on iconicity

Callaghan (2000) asked 2.5-year-olds and 3-year-olds to match stimuli to one of two choice objects. The stimuli used were of four different types that differed in iconicity (in what was intended to be an increasing order): “graphic symbols,” “pen symbols,” “color symbols,” and “replica symbols.” While 2.5-year-olds failed the task with all stimuli, 3-year-olds matched all the signs correctly to the referent. But 3-year-olds’ performance was significantly poorer in the “graphic

⁴ A notable difference between DeLoache’s testing situation and that of Kuhlmeier and colleagues, however, is that the former involved an unknown location, whereas the latter took place in the familiar cage. Since familiarity has often turned out to be an important factor in other studies of children and apes, this difference is perhaps not negligible. Cf. Lenninger (2009) for an attempt to eliminate this difference.

condition” than in other conditions, suggesting that the “level of iconicity” (which was the lowest for the “graphic symbols”) had an effect on children’s performance. In another matching task, two objects with the same basic level verbal label were paired, so that the children could not simply match the label with the correct object when making the choice. But when 2.5-year-olds were presented with choice objects that had different verbal labels, and only one of the choice objects matched the correct object to the verbal labels, their performance rose above chance level. Callaghan argues that both verbal and image-based representations are used when processing graphic symbols of objects, but that younger children might rely more on verbal presentations.

A home-raised chimpanzee, named Viki, who had been trained (unsuccessfully) to master spoken language, also was required to match a real object to one of two choice pictures (Hayes and Hayes 1953). The correct picture was a sign of an object of the same class as the real object. Viki’s picture stimuli were of two types: realistic color pictures and black-and-white line drawings (comparable to the “pen symbols” in Callaghan study). She was successful with both types of choice stimuli. Whether Viki knew the labels for all the choice stimuli is impossible to know, and so there still remains the possibility that object labels helped Viki with matching real objects to pictures. More recently, however, testing the famous bonobos Kanzi and Panbanisha, Persson (2008: 245–276) showed that they were able to map lexigrams to pictures, and vice-versa, even in cases of low degrees of “realism.”

Morris (1971 [1946]) seems to have been the first to conceive iconicity to be a question of degrees: a film, for instance, is more iconic of a person than is a painted portrait because it includes movement. Moles (1981) constructed a scale comprising thirteen degrees of iconicity from the object itself (100%) to its verbal description (0%). Such a conception of iconicity has been argued to be problematic, not only because distinctions of different nature appear to be amalgamated, but also because it takes for granted that identity is the highest degree of iconicity and that the illusion of perceptual resemblance typically produced, in different ways, by the scale model and the picture sign are as close as we can come to iconicity besides identity itself (cf. Sonesson 1998; Kendon 2004: 2). A more neutral way of describing the case may well be to say that the original perceptual appearances have been submitted to different kinds of transformations (cf. Groupe µ 1992; Sonesson 2004).

A specific aspect of iconicity involves the possibility of static signs corresponding to temporal reality (and other temporal signs). One of our studies takes its inspiration from a widespread conception in classical aesthetics, notably in the work of Gotthold Ephraim Lessing (1964 [1766]), according to which a dynamical element, such as an action, is more easily identified in a static picture con-

figuration when it is shown in the penultimate phase of the action, just before the action is complete (cf. Sonesson 2004). From our point of view, there are two stages to this hypothesis: the first stage would be to see whether a single temporal phase of any kind is iconically as efficacious as the whole sequence; the second stage would then have to determine whether, for instance, the penultimate phase is more efficacious than the final one.

1.4 Aims and means of the present study

In our study, we tested whether a nursery-reared chimpanzee, named Alex, would be able to understand that pictures and videos of the experimenter demonstrating an action represented a real demonstration and so would imitate the action presented in the picture or video. The pictures and videos differed in levels or kinds of iconicity. Alex had previously been trained with the “Do as I do” procedure on a few actions demonstrated by an experimenter, but he had never before been presented with pictures or videos of demonstrations. The home-reared chimpanzee Viki was reportedly able to imitate an action presented to her in the form of a video, a black-and-white photo or a line drawing (Hayes and Hayes 1953). However, this was never systematically tested; and the report does not provide any methodological details. The present sequence of experiments can be seen as a remake of the Viki study with tighter controls. At the same time, our study systematically uses the ability to imitate the behavior rendered in the pictures and videos as an indication of the presence of picture understanding. In the end, this led to the introduction of a further kind of variation: since actions are depicted, and since they can be complete or not, we wanted to see whether the alternative of rendering the final or penultimate phase of the action sequences made any difference.⁵

2 “Do as I do” training

Contrary to a widespread misconception, contemporary consensus has it that apes are not good at “aping”: that is, in other words, imitating (cf. Tomasello

⁵ When introducing this variation to the study, we were not aware of a similar suggestion having been made by Persson (2008: 61), precisely when discussing the Viki case: “One must infer what happened before a static view, and what will happen just after it, in order to read *clapping* and *patting* into the relations between body parts.” However, Persson himself claimed to have been inspired to this observation by Sonesson (1989).

1999, 2008; but cf. De Waal 2009). That said, there is an established technique, “Do as I do,” which consists in training subjects to attempt to copy, on command, what a model does (cf. Custance et al. 1995). Before any of our studies could be conducted, this training needed to be completed. Numerous repetitions are required for the training to work (in our case, 92 sessions over seven months).

2.1 Method

2.1.1 Subject

A six-year-old chimpanzee (*Pan troglodytes*) named Alex participated in this and all subsequent studies reported here. He is housed at the Wolfgang Köhler Primate Research Centre at Leipzig Zoo in Germany. Until the age of three, he was raised in a nursery together with two other chimpanzees; but at the time of the study, he was living in a social group with two juvenile and three adult chimpanzees, with access to spacious indoor and outdoor enclosures. The chimpanzees are fed several times per day and are never food deprived; water is available to them *ad libitum*. Prior to this study Alex was video-, photo- and drawing-naïve. Alex was tested in an indoor testing room that was familiar to him.⁶

2.1.2 Actions

Alex was trained to copy 24 actions: nine were familiar to him and fifteen were novel (Table 1). The novel actions were taken from Custance et al.’s (1995) study. For a transfer test, we used 45 novel actions also taken from the Custance et al.’s study. The actions were categorized into eight groups (Table 2). One of the actions – “touch armpit” – was omitted, since time was lacking for more training sessions, and we estimated that the number of actions was sufficient for the purpose of our study.

⁶ For those accustomed to studies of human subjects, and/or other animal studies, it may seem strange that only one subject is tested. However, that just one or a few subjects are involved is quite currently the case in studies of primates, since the availability of subjects, mostly from zoos, is rather limited. Increased caution is required, then, in drawing general conclusions.

Table 1: List of training actions and their descriptions

Old actions

1. Clap	the palms were hit together several times
2. Hit head	the top of the head was hit with right hand several times
3. Index to cheeks	index fingers were pressed to both sides of the cheeks
4. Shake lip	the lower lip was taken with thumb and index finger and then moved vigorously
5. Hit stomach	the stomach was hit several times with right hand
6. Shake ear	the right ear was taken with thumb and index finger and then moved vigorously
7. Smack	a sound was made with a tongue
8. Protrude tongue	the tongue was put out of the mouth
9. Fish face	mouth was closed and the cheeks were suck in

Novel actions

10. Shake hand	the right hand was shaken loosely from the wrist
11. Raise one arm	the right arm was put into the air
12. Stamp foot	the right foot stamped the floor several times
13. Pat stomach	the stomach was patted alternately with the palm of each hand several times
14. Raise two arms	both arms were put into the air simultaneously
15. Touch chin	the right index finger was placed on the chin
16. Praying hands	both palms were touching each other
17. Wipe face	the palm of one hand was wiped down over the face several times
18. Slap floor	the floor was slapped several times with the right palm
19. Raise foot	the right foot was raised from the floor
20. Wipe floor	the right palm was wiped from side to side across the floor several times
21. Touch armpit	the left arm was raised and the right index finger was placed on the left armpit
22. Grab wrist	the left wrist was grasped by the right hand
23. Swing arm	the right arm was swing back and forth several times
24. Wipe hands	the palms were wiped together several times

2.1.3 Procedure

While Alex was still in the nursery, he already received some training in the “Do as I do” procedure. His caretakers imitated him occasionally, and encouraged him to imitate them back. Subsequently, this “copying game” was made more formal by one caretaker introducing a clicker and small food rewards. Most of the actions at that time (actions 1–9 in Table 1) were “invented” by Alex on his own, and the caretaker repeated them, but others were demonstrated to him and repeated

Table 2: List of novel transfer actions and their descriptions. Actions that Alex imitated correctly are marked as such in the third column and good approximations are given as “Approx”

Action	Description of an action (after Custance et al., 1995)	Result
<i>Facial actions</i>		
1. Protrude lips	The front teeth were clicked together several times	Imitated
2. Lip smacking		Imitated
3. Teeth chattering		
4. Puff out cheeks		
5. Close eyes		
<i>Single hand actions</i>		
6. Open hand	The hand was held up with all the digits splayed apart	Approx.
7. Finger wiggling	The fingers were sequentially curled and straightened	Imitated
8. Stiff wave	The hand was waved stiffly from the wrist	
9. Raised index	The hand was held in a fist except for an extended index finger	
10. Hitchhiker's thumb	The hand was held in a fist except for an extended thumb	
11. Circle	All the fingers were arched over so their tips touched the tip of the thumb	Imitated
<i>Symmetrical hand actions</i>		
12. Index fingers touch	The tips of both index fingers were held together	
13. All digit tips touching	The fingers of both hands were splayed apart and bent with the tips of all the equivalent digits touching	
14. Interlink fingers	The palms of both hands were held together and the fingers were interweaved and curled over	
15. Roll hands	The hands were held in fists with one in front of the other and alternately circled around one another several times	
16. Peek a boo	Both hands were held side by side in front of the face. They were then moved apart to reveal the face and brought back together again.	
<i>Asymmetrical hand actions</i>		
17. Clap back of hand	The right palm slapped the back of the left hand several times	
18. Two fingered clapping	The first two fingers of the right hand slapped the left palm several times	
19. Palm point	The tip of the right index finger touched the left palm	Imitated
20. Palm punch	The right hand formed a fist which punched the left palm	
21. Grab thumb	The left thumb was grasped by the right hand	
22. Rabbit hole	The left hand formed a circle and the right index finger was placed inside it	

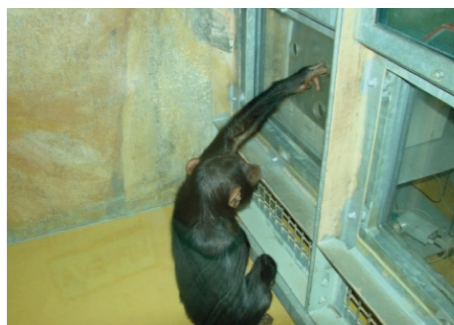
1 Table 2 (cont.)

3 Action	Description of an action (after Custance et al., 1995)	Result
4 <i>Touch parts of the body in sight</i>		
5 23. Shoulder	The right hand was placed on the left shoulder	Imitated
6 24. Elbow	The right index finger was placed on the left elbow	
7 25. Stomach	The right hand was placed on the stomach	Imitated
8 26. Thigh	The right hand was placed on the right thigh	
9 27. Knee	The right hand was placed on the left knee	
10 28. Foot	The right hand was placed on the left foot	
11 <i>Touch parts of the body out of sight</i>		
12 29. Back of head	The right hand was placed on the back of the head	
13 30. Top of head	The right hand was placed on the top of the head	Imitated
14 31. Nose	The tip of the right index finger was placed on the end of the nose	Imitated
15 32. Ear	The tip of the right index finger was placed on the right ear	Imitated
16 33. Clap behind	Both hands were clapped together behind the back	
17 34. Elbow behind	The right hand was brought around the back to touch the left elbow	
19 <i>Face/head related actions</i>		
20 35. Whistle	One long whistle was blown	Approx.
21 36. Mouth pop	The right hand formed raise index finger with back of the hand facing towards the actor; the end of the index finger was placed against the inside of the left cheek, and the hand was jerked from the wrist out of the mouth making a popping sound	
22		
23		
24		
25 37. Lip wobbling	The lips were protruded and the side of the right index finger was rubbed up and down over them	Imitated
26		
27 38. Mouth pull	The corners of the mouth were stretched using the index fingers	Approx.
28		
29 39. Look up	The head was tipped up and down once	
30 40. Look right	The head was turned to the right once	Imitated
31 <i>Whole body actions</i>		
32 41. Jump	Standing upright, the demonstrator jumped once	
33 42. Turn around	The demonstrator turned trough 360 degrees using several small steps	
34 43. Flap arms	Both arms were waved up and down as if imitating a bird	
35		
36 44. Hug self	The arms hugged the upper body as it was twisted back and forth from the waist	
37		
38 45. Foot to foot	Each foot was alternately raise and lowered several times	
39		
40		

often as he showed an interest in them. The whole previous training was done in a play situation with the caretaker being in the room with Alex. In the present study (during the training and all subsequent experiments, except in the Experiment 5a), however, the experimenter (*E*) and Alex were seated opposite each other, separated by a Plexiglas panel and mesh. While an action was demonstrated, *E* would ask Alex to repeat the action by saying, “do this.” When Alex reproduced the demonstrated action, *E* pressed a clicker, which served as a secondary reinforcer. After every three to five clicks, Alex was given a food reward – raisin, grapes or bananas – which served as a primary reinforcer. Importantly, Alex learned novel actions (actions 10–24 in Table 1) through shaping technique, which means that in the early stage of learning a novel action, he was at first rewarded for rough approximations to that action, and later on he would be rewarded only for ever closer approximations, and eventually, when he was proficient in an action, he was rewarded only for correct imitations. The training sessions lasted from 5 to 25 minutes, but the average was 15 minutes.



a



b

Fig. 1: (left) Alex imitating action Raise two arms. (right) Alex in front of the Plexiglas panel (see section 7)

The “Do as I do” training consisted of four phases: 1) Training of the 24 actions. 2) *Evaluation* of the training actions, which consisted of three sessions. In each session all training actions except “touch armpit” were demonstrated once. 3) A *transfer test* with novel actions: Forty-six novel actions were mixed with 22 old (training) actions. One session consisted of 24 (or 20) actions: 12 (or 10) transfer and 12 (or 10) training actions. After four sessions during which all 46 novel actions had been tested once, the order of actions was randomized and another four sessions completed. Therefore, eight sessions were conducted in total, and each action was tested twice. The procedure was exactly the same as in the training phase, except that if Alex did not respond immediately with a correct action, *E*

1 repeated the demonstration two more times. 4) Training with video, photo, and
2 drawing demonstrations. Alex was familiar with the “Do as I do” procedure from
3 live demonstrations, but not from demonstrations of an action being presented
4 on a computer. The computer was a 15-inch laptop, which rested on *E*’s lap; the
5 demonstration was in the form of a video, photo, or drawing.⁷ Alex received eight
6 training sessions with three actions that he previously copied in the transfer test
7 when presented to him live: lip wobbling, nose touching, and ear touching. (De-
8 scriptions of the actions are in Table 1.) These actions were demonstrated to him
9 live, in a video, in a photo or in a drawing.

12 2.2 Results and discussion

14 It took Alex 92 sessions presented over seven months to copy 23 out of the 24
15 training actions reliably. In the *evaluation phase*, Alex correctly reproduced 80%
16 or more of the demonstrated actions. In the *transfer test*, Alex correctly copied 12
17 out of the 45 novel actions (27%), which were from six different groups (see Table
18 2). We concluded that he had copied the action if he produced the demonstrated
19 action in one of his first three responses. Three more actions could be considered
20 good approximations of demonstrated actions: open hand (he lifted his hand a
21 little bit toward the panel; whistle (he blew air and produced a sound); mouth
22 pull (he put both index fingers on the inside of the lower lip). Our findings are
23 consistent with the results from the Custance et al.’s. (1995) study. In the last two
24 sessions of the training with video, photo, and drawing demonstrations, Alex cor-
25 rectly reproduced the demonstrated action in 100% of the cases where it was
26 demonstrated live or in a video, 76% of the cases where it was presented as a
27 photo, and 21% of the cases where it was presented as a drawing.

30 3 Experiment 1 – video, photo, drawing

32 3.1 Method

34 3.1.1 Materials and actions

36 We used twenty actions, which Alex reliably copied in the training phase. The
37 actions were classified into two groups: iterative (repetitive) and terminative

40 ⁷ See the “materials” section of Experiment 1 for a description of videos, photos, and drawings.

(non-repetitive) actions. The *iterative* actions are actions that keep repeating in a sequence, or, put in other terms, actions that consist of a number of action parts that are more or less identical and that follow each other without any fixed limit (e.g., *shake hand* or *stamp foot*). *Terminative* actions, on the other hand, are actions having a clear end state that are performed only once, or, in other words, actions that consist in bringing the parts of the body involved unto an end state (e.g., *touch nose* or *praying hands*). The actions were demonstrated to Alex by the experimenter either live, or in the form of videos, photos, and drawings, on a 15-inch computer monitor, held in the lap by *E* who sat motionless. The videos were between two and five seconds long, projected in a loop, so that a single video strip might repeat itself a few times during a demonstration. Photos were in color and showed *E* in the process of demonstrating an action. If an action was *terminative* (e.g., *touch nose*), the photo showed the action at its end state. If it was *iterative* (e.g., *shake hand*), it was presented as a random sequence. The drawings were black-and-white sketches of the photos (vectoralized from the photos in Photoshop). Like the photos, the drawings were showed by *E* during the demonstration of an action. Depending on the extent to which the experimenter's body was visible on the videos, photos, and drawings showed during a demonstration, the experimental material was divided into two groups: 1) "whole body" demonstrations showing the demonstrator's entire body performing an action and 2) "zoomed in" demonstrations showing only part of the demonstrator (see Table 4). During each demonstration, the videos, photos, and drawings were presented for 5–15 seconds in full-screen mode on the laptop.

3.1.2 Procedure

There were four conditions, depending on how a demonstration was presented: 1) live condition, 2) video condition, 3) photo condition, and 4) drawing condition. The procedure was the same as in the training phase. During a demonstration, Alex was asked to repeat the demonstrated action. If he performed the correct action, *E* pressed the clicker. If he did not respond or performed the wrong action in the "live" condition, *E* demonstrated the action again and asked Alex to repeat the action; *E* performed a maximum of three demonstrations of one action. In the other three conditions, *E* asked Alex two more times "do this" while a video, photo, or drawing kept showing the demonstration. All twenty actions were presented four times in one session that lasted approximately fifteen minutes, once in each condition. Conditions were blocked in groups of five actions. In total, two sessions were conducted over two consecutive days.



Fig. 2: Examples of the photos and drawings used in Experiments 1 – a photo and a drawing of (left) Protrude tongue action, and (right) Raise foot action

3.1.3 Coding and analysis

All sessions were filmed. Alex’s responses from all the videos were coded. For the analysis, we considered only the first response he produced after a demonstration. The responses were classified into four types: 1) Correct action (the action demonstrated to him), 2) Known action (one of the actions that had previously been demonstrated to him), 3) New action (an action previously not demonstrated), and 4) No response (no reaction was observed). It is impossible to estimate the likelihood of Alex performing the correct action by chance, since, in theory, there are an infinite number of potential actions he could produce. However, Alex was previously rewarded for 24 training actions plus 12 transfer actions; therefore we will take these 36 actions as a pool of all the possible actions he could choose

from. Thus, if we consider only his first produced actions, there is 2.78% possibility that Alex produced the correct action by chance. Therefore, when we run a Binomial test to investigate whether Alex successfully copied the demonstrated actions above chance level in each condition, we tested his performance against $p = 0.0278$. A second coder, blind to the condition as well as to the action demonstrated, scored 32 actions (20% of Alex's first responses) performed by Alex to assess inter-observer reliability. The second coder scored whether Alex performed one of the 20 test actions, a new action or nothing. Observers' agreement was excellent (Cohen's kappa = 0.97).

3.2 Results and discussion

Table 3 presents Alex's responses (correct actions, known actions, new actions, and absence of action) to demonstrations in the four conditions. Alex reliably copied demonstrated actions above chance in the live (77.5%), video (40%), and photo (25%) conditions (Binomial test: $p < 0.05$ for all conditions), but not in the drawing condition (7.5%; Binomial test: $p > 0.05$). Alex's success in correctly copying a demonstrated action differed across the four conditions (Friedman test: $\chi^2 = 27.841$, $p < 0.001$, $df = 3$, $N = 20$). Alex's performance in the live condition was significantly better than in the other three conditions (Wilcoxon test: $z > 2.8$, $p \leq 0.004$, $N = 20$, all cases). There was also a significant difference between the video and drawing conditions ($z = 2.469$, $p = 0.016$, $N = 20$), but no differences

Table 3: Experiment 1: Types of Alex' responses to the demonstrated actions

Condition	Session	Responded action			
		Correct (out of 20)	Known	New	Nothing
Live	1	17*	2	1	0
	2	16*	1	3	0
Video	1	7*	8	5	0
	2	9*	7	4	0
Photo	1	5*	8	4	3
	2	5*	8	4	3
Drawing	1	1	9	5	5
	2	2	4	8	6

Note. * Binomial test, $p < 0.001$

1 were found between the video and photo conditions, nor between the photo and
2 drawing conditions. If we compare Alex's performance with *iterative* and *termina-*
3 *tive* actions, we find no significant differences between them in any condition.

4 The fact that the subject consistently performed better on the live condition,
5 and that there is a decrease both in success and correctness from the live condi-
6 tion to the video condition, as well as from the video condition to the photo condi-
7 tion, seems to confirm, and extend to apes, the idea, voiced by Callaghan (2000)
8 among others, that there is a kind of "scale of iconicity" involved, at least if we
9 exclude replicas, on one extreme, and verbal description on the other.⁸ Perhaps
10 we should rather talk about familiarity here, relative to the direct experience in
11 the world of perception (perhaps this is Callaghan's "realism"). It seems obvious,
12 in any case, that this is not a question of mere quantity of properties correspond-
13 ing between the sign and its target (Moles' 0 to 100%), but of certain properties
14 being essential. A more thorough variation of properties would be needed to es-
15 tablish this, but this must be left for another study.

16 The lack of difference regarding iterative and terminative actions deserves to
17 be noted. We will return to this distinction in Experiment 4.

18 Overall, the results suggest that – in some sense or other – Alex understood
19 that the videos, photos, and drawings represented the actions that the experi-
20 menter wanted him to imitate, even though he was unable to copy the actions in
21 the drawing condition.⁹ Indeed, the very fact that results were different across the
22 different conditions strongly suggests that some process of interpretation was go-
23 ing on. Moreover, as the experimenter was sitting still and resting the computer
24 on her lap, it seems implausible that Alex could confuse the image of the experi-
25 menter seen on the laptop with the real experimenter – especially since all ac-
26 tions that Alex correctly reproduced in the video, photo, and drawing conditions
27 were presented in the "zoomed in" version, thus only showing a part of the ex-
28 perimenter. It is possible to conceive the "zoomed in" variant as being a kind of
29 attention focusing device. However, when we tested this assumption in an addi-
30 tional experiment (not reported here), the results showed that the type of a demon-
31 stration did not have any influence on Alex's performance. If he reproduced an

32

33

34 ⁸ Replicas may be 100% iconic, but it is more difficult to see them as signs than, for instance,
35 pictures, as DeLoache's experiments have shown. Sonesson (1994) distinguished two kinds of
36 iconicity here, primary and secondary iconicity, depending on the iconic relation or the sign rela-
37 tion being most directly accessible.

38 ⁹ The reason for this conclusion is that three parts of the time Alex responded to the drawing
39 with an action, albeit often with a wrong one. Since the drawings used here were vectorialized
40 from photographs in Photoshop, we do not really know how this result compares to manually
produced drawings.

Table 4: Experiment 1: List of the demonstrated actions in the four conditions and how often Alex copied them correctly

Demonstrated Action	Demonstration Type	Action Type	Condition			
			Live	Video	Photo	Drawing
Clap	Zoomed in	Iterative	1	0	0	0
Grab wrist	Zoomed in	Terminative	0	1	2	2
Hit head	Zoomed in	Iterative	2	2	0	0
Index to cheeks	Zoomed in	Terminative	2	1	1	0
Pat stomach	Zoomed in	Iterative	2	1	0	0
Praying hands	Zoomed in	Terminative	2	0	2	1
Protrude tongue	Zoomed in	Terminative	2	2	2	0
Raise foot	Whole body	Terminative	2	0	0	0
Raise one arm	Zoomed in	Terminative	2	0	0	0
Raise two arms	Zoomed in	Terminative	1	0	0	0
Shake ear	Zoomed in	Iterative	2	2	0	0
Shake hand	Zoomed in	Iterative	2	2	0	0
Shake lip	Zoomed in	Iterative	2	2	2	0
Slap floor	Whole body	Iterative	2	0	0	0
Stamp foot	Whole body	Iterative	2	0	0	0
Swing arm	Whole body	Iterative	0	0	0	0
Touch chin	Zoomed in	Terminative	1	0	0	0
Wipe face	Zoomed in	Iterative	2	2	1	0
Wipe floor	Whole body	Iterative	2	0	0	0
Wipe hands	Zoomed in	Iterative	2	1	0	0

Note. 0 – Alex did not correctly copy the action, 1 – Alex copied the action in one of the sessions, 2 – Alex copied the action in both sessions

action when it was presented in a “zoomed in” photo or video, he also correctly reproduced this action when it was presented in a “whole body” photo or video. Therefore, it appears that the difference of results observed earlier did not depend on the potential attention-focusing device of “zooming-in” on the action.

4 Experiment 2 – black-and-white photos

The difference in Alex’s performance with colored photos and black-and-white line drawings in Experiment 1 could arise from the difference in colors. Maybe colored photos gave Alex more information than black-and-white drawings – or they were “more iconic,” in the sense of Callaghan (2000). To test this, we pre-

sented Alex with colored as well as black-and-white photos of the experimenter demonstrating an action.

4.1 Method

4.1.1 Materials and actions

We used 10 colored and 10 black-and-white photos presenting the same 10 actions (see Table 5). The colored photos were the same photos used in previous experiments and were of “zoomed in” type, and the black-and-white photos were identical to the colored ones except that they were in grey-scale.¹⁰

4.1.2 Procedure

There were three conditions: 1) Live condition 2) Colored photos condition 3) Black-and-white photos condition.

Two sessions were conducted. In each session each action was demonstrated three times, once in each condition. At the beginning of each session first all ten actions were demonstrated live. After that both conditions alternated blocked in 5 actions.

4.2 Results and discussion

Results are presented in Table 5. Alex performed significantly different on the three conditions (Friedman test: $\chi^2 = 9.652$, $p = 0.004$, $df = 2$, $N = 10$). Post hoc tests showed that Alex copied significantly more actions in the live condition compared to the black-and-white photos condition (Wilcoxon test: $z = 2.460$, $p = 0.016$, $N = 10$), but not compared to the color photos condition (Wilcoxon test: $z = 1.633$, $p = 0.250$, $N = 10$). However, the crucial comparison is between Alex’s

¹⁰ Ideally, we should of course have used novel actions for both sets of photos, but this was impossible to do, since no more actions were available. This means that Alex had seen half of the material from Experiment 2 twice before in Experiment 1. However, if we compare Alex’s performance on both experiments with statistical analysis, we find that this fact had no importance for the results (for photos Wilcoxon test: $z = 1.134$, $p = 0.5$, $N = 12$; for videos Wilcoxon test: $z = 1.732$, $p = 0.25$, $N = 12$).

Table 5: Experiment 2: List of actions used in the three conditions and how often Alex copied them correctly

Actions	Condition		
	Live	Colored photos	Black-and-White photos
Grab wrist	2	2	1
Hit head	2	0	0
Index to cheeks	2	2	2
Praying hands	1	1	1
Protrude tongue	2	2	2
Shake ear	2	2	1
Shake Lip	2	1	0
Touch chin	2	2	1
Wipe face	2	2	1
Wipe hands	2	0	1

Note. 0 – Alex did not correctly copy the action, 1 – Alex copied the action in one of the sessions, 2 – Alex copied the action in both sessions

performance on the color and black-and-white photos conditions where we did not find a significant difference (Wilcoxon test: $z = 1.633$, $p = 0.219$, $N = 10$). Thus, when compared to the live condition, the color did seem to have an influence on Alex’s performance – Alex copied significantly less when demonstrations were presented in black-and-white photos but not when they were presented in colored photos; however, if we look only at the two photo conditions this difference was not significant. Therefore, the results suggest that the colors of the photos do not have a huge impact on Alex’s performance. However, our data could be taken to suggest that there is something like a “scale of iconicity,” certainly in the relation between reality and pictures, and perhaps also between colors and black-and-white photos.

5 Experiment 3a – end state versus incomplete action: Live demonstrations

The purpose of Experiment 3, of which the first condition is described below, was to investigate the idea, common in classical aesthetics, according to which a static

1 version of an action is more readily interpretable from the penultimate, rather
2 than the final, phase of the action.

5.1 Method

5.1.1 Actions

9 We used 10 actions selected from the list of the training and transfer actions (from
10 the “Do as I do” training phase): Touch chin, Touch ear, Touch nose, Index to
11 cheeks, Touch head, Touch stomach, Grab wrist, Touch armpit, Praying hands,
12 Protrude tongue (for the descriptions see Tables 1 and 2). All actions were termi-
13 native and consisted of a movement of a body part (i.e., for action “touch head”
14 – raise a hand on top of your head) and an end state (the hand is on top of the
15 head).

5.1.2 Procedure

20 All actions were demonstrated live but there were three possible variations of the
21 demonstration: 1) Full condition: Alex saw a full demonstration of an action with
22 a movement and an end state. 2) Incomplete condition: Alex saw a movement of
23 a body part but no end state of an action. *E* stopped the action just before the end
24 state. 3) End state condition: Alex only saw a demonstration of an end state of an
25 action with no movement of a body part.

26 In one session all actions were demonstrated three times – once in each con-
27 dition. First all 10 actions were demonstrated in Full live condition, then the two
28 other conditions alternated, blocked in 5 actions. There were 2 sessions conducted
29 in two consecutive days.

5.2 Results and discussion

35 Results are presented in Table 6. Alex’s performance did not differ on the three
36 conditions (Friedman test: $\chi^2 = 3.920$, $p = 0.150$, $df = 2$, $N = 10$). In the Full and End
37 state conditions he copied 9 out of 10 actions and in the Incomplete condition he
38 copied 8 actions.

39 In themselves, these results have no bearing on the original hypothesis,
40 which, as formulated, concern static representations of actions. The negative

Table 6: Experiments 3a and 3b: List of actions used in the three conditions of each experiment and how often Alex copied them correctly.

Demonstrated Action	Experiment 4a – live			Experiment 4b – photos		
	End state	Incomplete	Full live	End state	Incomplete	Live
<i>Terminative actions</i>						
Ear	1	1	2	2	2	1
Grab wrist	2	2	1	2	2	2
Index to cheek	2	0	2	0	1	0
Nose	2	1	2	2	1	2
Praying hands	2	2	2	1	0	2
Protrude tongue	2	2	2	0	1	2
Stomach	2	1	2	0	0	2
Top of the head	2	1	0	0	1	2
Touch armpit	1	1	2	2	2	2
Touch chin	0	0	1	1	1	0
<i>Iterative actions</i>						
Pat stomach				0	0	1
Slap floor				0	1	0
Stamp foot				0	0	0
Swing arm				0	0	1
Wipe face				1	2	2

Note. 0 – Alex did not correctly copy the action, 1 – Alex copied the action in one of the sessions, 2 – Alex copied the action in both sessions

results may however be interesting for another reason. They suggest that Alex is able to anticipate elements lacking in the sequence of action with reference to the model of action he has previously learned. Tomasello (1999) has argued that apes learn by means of *emulation* (copying of goals) rather than by true *imitation* (copying of means).¹¹ If the final phase of a terminative action is identified with its goal (which is at least one possible interpretation), it would seem that apes are capable of attending to the means part of the action. Since this is not a learning context, this result, so far, can only be taken as suggestive. In the present context, however, these results are only interesting taken together with the results of the following experiment.

¹¹ Tomasello (2008), however, puts the emphasis on apes' learning by means of ritualization of intention-movements.

6 Experiment 3b – end state versus incomplete action: Photo demonstrations

To establish the second part of Lessing's hypothesis, as described in section 1.3, we will have to demonstrate that static representations of the penultimate phase of an action are most iconically efficacious in relation to the whole sequence of action than the final phase.

6.1 Method

6.1.1 Materials and actions

We used 30 photos presenting 15 actions (for a list of actions see Table 6; examples see Figure 3). Half of the photos presented the 15 actions at their end states and were used in the "end state photo" condition, and the other half presented the same 15 actions a moment before they reached the end state, which were used in the "incomplete photo" condition. Ten actions were terminative actions, and thus were executed only once and had a clear end state (e.g., touch nose; these were the same actions as in Experiment 3a). The other 5 actions were iterative actions, but since iterative actions are made up of a sequence of terminative sub-actions, they were shown at the moment of reaching the end-point of one of the sub-actions (e.g., stamp foot – at the moment the foot hit the floor).

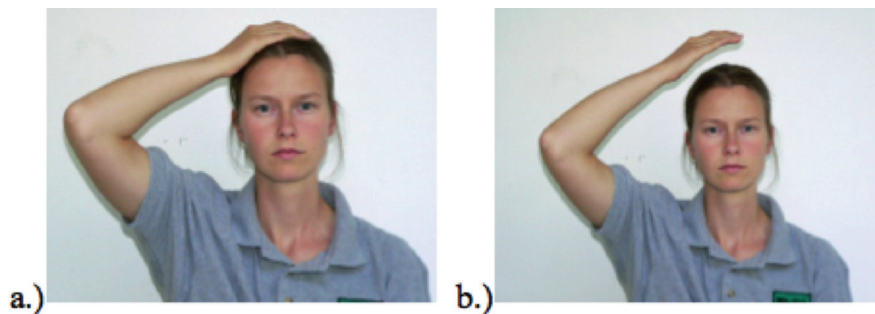


Fig. 3: An example of photos used in Experiment 3b – a) photo of the end state demonstration of the Top of the head action, and b) photo of the incomplete demonstration of the Top of the head action

6.1.2 Procedure

The procedure was the same as in all previous Experiments. There were three conditions: 1) Full live condition (Alex saw a full demonstration of an action performed by the experimenter). 2) Incomplete photo condition 3). End state photo condition.

Two sessions were conducted. In one session all actions were demonstrated three times – once in each condition.

6.2 Results and discussion

Results are presented in Table 6. Alex's performance neither differed between the three conditions (Friedman test: $\chi^2 = 3.282$, $p = 0.196$, $df = 2$, $N = 15$) nor was there any difference in the performance of the iterative and terminative actions in any of the conditions (Mann-Whitney test: $z < 1.62$, $p > 0.13$, $N = 15$, all cases). Moreover, in Experiment 3a, we saw that the different versions of the live condition, of which only the full version was repeated here, did not differ significantly.

As far as apes are concerned – or, more exactly, in the case of Alex – Lessing's hypothesis is not supported by our results.¹² The latter are, however, interesting for other reasons: it suggests that not only when presented in live action, but also in a static view, the action can be identified without having reached its conclusion. What is meant by identifying the action is of course an important issue, to which we will return in the concluding discussion. However, we seem to have established the first part of the hypothesis, as we analyzed it above (Section 1.3): a single static view may be as iconically efficacious as a whole sequence.

7 Experiment 4 – discrimination task

In Experiment 1 Alex failed to copy actions if they were demonstrated to him through line drawings. Here we wanted to test whether the drawings conveyed enough information for him to discriminate them. If so, it seems that it is the copying task as such that causes the problem. In addition, we were interested to see, whether Alex would be faster to discriminate between stimuli showing E in the process of demonstrating an action than Alexandra did, a chimpanzee who

¹² As far as we know, Lessing's hypothesis has never been tried out experimentally with human beings, which deprives us of important comparative material.

1 had neither experience with the “Do as I do” paradigm nor had ever seen E perform these actions.

5 7.1 Method

7 7.1.1 Subjects

9 Subjects were two chimpanzees: Alex, who participated in all the previously described experiments, and Alexandra, who was housed together with Alex, but neither had any previous experience with the “Do as I do” paradigm, nor had previously seen any of the photos and drawings used in this experiment.

15 7.1.2 Materials

17 We used 27 samples of three different types (see Figure 4). They were all 300×300 pixels large. Nine samples, which were used in the “object” condition, were photos of random objects that neither of the subjects had ever seen before. Nine samples, which were used in the “photo” condition, were photos of the experimenter demonstrating an action. These photos were taken from Experiment 1; however, they were made smaller, and some of them were also cropped so as to fit the size 300×300 pixels. Alexandra had never seen these photos. Nine further samples, which were used in the “drawing” condition, were line drawings of the experimenter demonstrating an action. These drawings were also taken from Experiment 1 and were reduced in size, and some of them were also cropped so as to fit the size 300×300 pixels. Each photo and drawing presented different actions; thus there were 18 different actions shown in the samples. All samples were presented on a 21-inch touch screen.

32 7.1.3 Procedure

34 The touch screen monitor was placed against the front of the cage, where the five holes in a Plexiglas panel allowed the chimpanzees to touch the monitor (Figure 1 [right]). Both chimpanzees had previous experience with the touch screen and two-choice discrimination tasks. In a discrimination task two samples are presented simultaneously on the monitor – one correct sample, which is always rewarded, and one distracter, which is never rewarded. In the present task each correct sample was paired with two alternatives. However, the correct sample was

always presented with only one of the two alternatives at the time, which one being randomly assigned. Each such triad (correct sample and two alternatives) was unique. If the subject touched the correct sample, the pellet machine automatically delivered a pellet, and the next trial could be started. If the subject touched the incorrect sample, the monitor turned green for 5 seconds, after which a new trial was initiated. Before each trial the subjects had to touch the “starting” sample to initiate the trial.

There were three conditions: 1) object condition, 2) photo condition (All actions were known to Alex; however the second triad consisted of actions that Alex hadn’t successfully copied in any of the previous experiments when demonstrated to him in photos), 3) drawing condition (when Alex was presented with these drawings in Experiment 1 he did not copy any of the actions).

One session consisted of 20 trials. One to four sessions could be conducted per day. Both subjects started with the Object condition. When the subjects reached the criterion – defined as choosing the correct sample in 75% cases in one session – they proceeded to the photo condition. Again, when they chose the correct sample in 75% cases or above, they proceeded to the drawing condition. When they also reached the criterion in the drawing condition, the whole cycle of three conditions was repeated for two more times with different triads of samples. If for one triad they did not reach the criterion in ten sessions, they proceeded to the next condition. They both received the same order of triads (and conditions). The method was designed to investigate whether discrimination was possible as well as the speed of learning.

When Alex finished with the first part of the study, he was additionally tested on another three sets of drawings (see drawing condition 2 in Figure 4). These drawings were black-and-white sketch copies of the photos used in the photo condition. In two triads the correct sample was the same action as in photo condition, but in one it was a different one.

7.1.4 Coding and analysis

The dependent measure was a number of sessions needed to reach the criterion. Since we only used Alex’s and Alexandra’s data for the purpose of comparison, we assigned them a score of 11 in cases where they failed to reach the criterion in 10 sessions. To be able to compare Alex’s and Alexandra’s performance (with Wilcoxon test), we pooled all three conditions. Moreover, we compared, separately for Alex and Alexandra, the performance on the three conditions by means of a Kruskal Wallis test. Because we only had three data points in each condition, we were unable to compare them pairwise.

7.2 Results and discussion

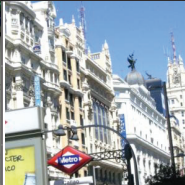
Triad	Correct sample	Distracter 1	Distracter 2	Session		
Object condition					Alex	Alexandra
1				1	1	4
				2	2	2
				3	1	3
Photo condition						
1				1	4	4
				2	2	7
				3	9	/
Drawing condition 1						
1				1	/	/
				2	2	/
				3	1	2
Drawing condition 2						
1				1	2	
				2	5	
				3	4	

Fig. 4: Some examples of pictures used in Experiment 4 and the number of sessions Alex and Alexandra needed to learn the correct sample

Pooling all three conditions, Alex learned the correct samples faster than Alexandra (Wilcoxon test: $z = 2.207$, $p = 0.031$, $N = 9$). Alexandra did not reach the criterion in two drawing triads and in one photo triad, and Alex did not reach the criterion in one drawing triad. Both subjects seemed to be able to discriminate

between the drawings just as good as they discriminated between the photos and objects (Kruskal Wallis test: Alex, $\chi^2 = 2.952$, $p = 0.293$, $df = 2$, $N = 9$; Alexandra, $\chi^2 = 2.550$, $p = 0.325$, $df = 2$, $N = 9$). Considering only the number of sessions that Alex needed to reach the criterion in the triads in the photo condition, Alex learned the correct sample for the triad of actions, which he had never before managed to correctly copy, as fast (or even faster) than for the other two triads. This would suggest, that “knowing” which action the photos demonstrate, did not help him to learn the correct sample faster.

The main conclusion, however, is that it is not lack of discrimination that accounts for Alex’s failure to execute the actions when looking at drawings of the actions. Alexandra’s results will not be discussed here, since they were only used to show that Alex’s experience did make a difference.

8 General discussion

All through these studies, Alex has shown himself to be capable of executing different sequences of action, being prompted by staged action sequences, drawings, and black-and-white and color photographs, no matter if the latter shows the penultimate or the final stage of the sequence. We cannot conclude of course that all apes, or even all chimpanzees, can repeat these feats; however, given certain circumstances (of which we only know what they are not), one case is sufficient to show that they lie, with a Vygotskian term, within their zone of proximal development.

Although there were significant differences between the results in all conditions, Alex performed above chance in all of them, except on the drawings (which he was however able to discriminate). The experiment was repeated with differently sized pictures, as well as with black-and-white as opposed to color photographs, without finding any differences. Finally, the task was conducted with pictures representing still actions with an incomplete goal as well as pictures of the same action in which the goal had been achieved (end state), once again without any significant difference between the two pictorial stimuli, while these had fewer correct responses compared to a live model.

The fact that the success rate in the case of live action, videos, and static pictures were so different would seem to indicate that some kind of interpretative work was going on. In the cases when the action was shown on video, it is not possible to say whether the live illustration of the action and the video were qualitatively different to Alex. Nevertheless, the quantitative difference resulting from using a video instead of a live action as a prompt may be taken to indicate such a qualitative difference. In any case, a still photo serving as a prompt for a real ac-

tion must certainly be considered different from the action, at the same time that it appears to have been taken by Alex to “stand for” it, as suggested by the fact that in Experiment 1 he performed the represented action more often than expected by chance. If so, there is a clear differentiation between expression and content. To suggest that Alex is simply confusing the still photo, and even more the photo of the incomplete action, where the picture prompting the action is two times removed from the action requested (as a sign and as a pre-final phase), seems indeed far-fetched. However, it is less clear whether the double asymmetry characteristic of signs could be attributed to Alex.

It is possible to conclude that picture understanding is within the purview of chimpanzee capacities, and since Alex was neither language-trained, nor engaged in any other form of sign use, we can also suppose that it is possible to understand pictures as iconic signs, quite independently of language. In this sense, Alex stands apart from Viki and Kanzi, whose feats in this domain cannot readily be dissociated from their language training, whether successful, as in the second case, or not, as in the first case. However, we should note that Alex had extensive experience with imitation both prior to the study, and during these studies. This is consistent with there being a close conceptual link between imitation and sign use (Sonesson 2007; Zlatev 2009).

In turn, it is not completely clear why Alex would reproduce actions depicted in complete or incomplete photos but not actions depicted in drawings. If this was simply a question of the amount of information conveyed, another result might be expected. At the beginning of Experiment 1 Alex received an equal amount of training for videos, photos, and drawings. One possibility is that Alex did not see the drawings as representations of the actions that he was required to reproduce, but merely as a series of lines on a white background, i.e., due to the degraded nature of the representation he operated in the “a-mode” of Daddesio (1995) or the “surface mode” of Persson (2008). The fact that he could discriminate between the drawings does not necessarily tell against such an interpretation.

A semiotically “rich” interpretation of this result could be that Alex not only used the picture as a sign for the real-world action, but that he could simultaneously recognize a complete action including its goal state from an earlier phase of its development, i.e., that he was capable to grasp a form of indexicality (in this case temporal contiguity) indirectly through the sign. Indeed, the fact that the static representation of the penultimate phase and of the final one served equally well to initiate the copying behavior on the part of Alex could be given a positive reading. Certain presuppositions, however, would have to be taken for granted. Perception leads to identification because each perceptual moment is saturated with possible earlier phases, which are more or less determined, as well as with possible later phases, which may receive more or less determination. In

phenomenology, the former ones are called *retentions* and the later ones *proten-* 1
tions (cf. Sonesson 1989). Alex had been trained on the complete actions. If the 2
 only thing you are offered is a single phase of these actions, then you have to 3
 pretend and/or retain the others phases, in order to see the actions as being the 4
 same. Some actions are no doubt only a way of getting the members of the body 5
 into a given static position, which is the real bearer of the meaning.¹³ In these 6
 cases, at least, it is natural for the final position to be as successful at suggesting 7
 the action to imitate, as the action as a whole. On the other hand, the fact that the 8
 penultimate phase serves as well to obtain this effect might be taken to suggest 9
 that Alex goes through a more complex kind of interpretative work, perceiving the 10
 single, static phase as being the expression for which the full action is the content. 11

Nevertheless, the major factor that argues against such an interpretation is 12
 the lack of any evidence concerning novel actions. Since all actions involved were 13
 taken from the set of actions on which Alex had been trained earlier on, and Alex 14
 has been known to have difficulties with the imitation of novel actions, we cannot 15
 exclude that a much more simple explanation in terms of conditional learning 16
 could be given. This would suppose that Alex could generalize what he had 17
 learned from the training involving the complete actions, not only to the render- 18
 ing of these actions involving different kinds of iconic transformations, but also 19
 to the different single static phases of such actions. If he is supposed to make this 20
 generalization on the basis of surface mode perception of the pictures, then it is 21
 not clear whether what is perceived is sufficiently similar to allow such a general- 22
 ization. Given our results, it appears more difficult to tell generalizations starting 23
 out from object mode and pictorial mode apart. 24

Yet, even though there might still be room for explaining Alex's performance 25
 in some more parsimonious way not involving signs, the differentiated responses 26
 to different varieties of iconic signs, as well as the successful recognition of the 27
 incomplete action representations, lend support to the "richer" interpretation. If 28
 Alex is yet not privy to true sign use, he certainly seems to be on his way to it. 29
 Further investigations must tell how far such an interpretation can be supported. 30

References 31

Barr, Rachel & Harlene Hayne. 1999. Developmental changes in imitation from television during 35
 infancy. *Child Development* 70(5). 1067–1081. 36

¹³ These would all be what we have called terminative actions above, but all terminative actions 39
 are not necessarily of this type. 40

- 1 Bates, Elisabeth. 1979. *The emergence of symbols: Cognition and communication in infancy*.
2 New York: Academic Press.
- 3 Cabe, Paul A. 1980. Picture perception in nonhuman subjects. In Margaret A. Hagen (ed.), *The*
4 *perception of pictures. Vol. II. Dürer's devices*, 305–343. New York: Academic Press.
- 5 Callaghan, Tara C. 2000. Factors affecting children's graphic symbol use in the third year:
6 language, similarity, and iconicity. *Child Development* 15(2). 185–214.
- 7 Callaghan, Tara C. & Mary P. Rankin. 2002. Emergence of graphic symbol functioning and the
8 question of domain specificity: A longitudinal training study. *Child Development* 73(2).
9 359–376.
- 10 Custance, Deborah M., Andrew Whiten & Kim A. Bard. 1995. Can young chimpanzees (*Pan*
11 *troglydytes*) imitate arbitrary actions? Hayes and Hayes (1952) revisited. *Behavior*
12 132(11–12). 837–859.
- 13 Daddesio, Thomas C. 1995. *On minds and symbols: The relevance of cognitive science for*
14 *semiotics*. Berlin & New York: Mouton de Gruyter.
- 15 DeLoache, Judy S. 1995. Early symbol understanding and use. In *Psychology of Learning and*
16 *Motivation*, vol. 33, Douglas L. Medin, (ed.), 65–114. New York: Academic Press.
- 17 DeLoache, Judy S. 2000. Dual representation and young children's use of scale models. *Child*
18 *Development* 71(2). 329–338.
- 19 DeLoache, Judy S. 2004. Becoming symbol-minded. *Trends in Cognitive Sciences* 8(2). 66–70.
- 20 DeLoache, Judy S. & Nancy M. Burns. 1994. Early understanding of the representational
21 function of pictures. *Cognition* 52(2). 83–110.
- 22 DeLoache, Judy S., Valerie Kolstad & Kathy N. Anderson. 1991. Physical similarity and young
23 children's understanding of scale models. *Child Development* 62(1). 111–126.
- 24 DeLoache, Judy S., Sophia L. Pierroutsakos, David H. Uttal, Karl S. Rosengren & Alma Gottlieb.
25 1998. Grasping the nature of pictures. *Psychological Science* 9(3). 205–210.
- 26 De Waal, Frans. 2009. *The age of empathy: Nature's lessons for a kinder society*. New York:
27 Harmony.
- 28 Fagot, Joël, Julie Martin-Malivel & Delphine Dépy. 2000. What is the evidence for an equivalence
29 between objects and pictures in birds and nonhuman primates? In Joël Fagot (ed.), *Picture*
30 *perception in animals*, 295–320. Hove: Psychology Press.
- 31 Groupe µ. 1992. *Traité du signe visuel: pour une rhétorique de l'image*. Paris: Éditions du
32 Seuil.
- 33 Hayes, Keith J. & Catherine Hayes. 1953. Picture perception in a home-raised chimpanzee.
34 *Journal of Comparative and Physiological Psychology* 46(6). 470–474.
- 35 Hayne, Harlene, Jane Herbert & Gabrielle Simcock. 2003. Imitation from television by 24- and
36 30-month-olds. *Developmental Science* 6(3). 254–261.
- 37 Husserl, Edmund. 1939. *Erfahrung und Urteil: Untersuchungen zur Genealogie der Logik*.
38 Prague: Academia.
- 39 Imura, Tomoko & Masaki Tomonaga. 2003. Perception of depth from shading in infant
40 chimpanzees (*Pan troglodytes*). *Animal Cognition* 6(4). 253–258.
- Kendon, Adam. 2004. *Gesture: Visible action as utterance*. Cambridge: Cambridge University
Press.
- Krampen, Martin (ed.). 1983. *Visuelle Kommunikation und/oder verbale Kommunikation?*
Hildesheim: Olms.
- Kuhlmeier, Valerie A. & Sarah Till Boysen. 2001. The effect of response contingencies on scale
model task performance by chimpanzees (*Pan troglodytes*). *Journal of Comparative*
Psychology 115(3). 300–306.

- Kuhlmeier, Valerie A. & Sarah Till Boysen. 2002. Chimpanzees (*Pan troglodytes*) recognize spatial and object correspondences between a scale model and its referent. *Psychological Science* 13(1). 60–63.
- Kuhlmeier, Valerie A., Sarah Till Boysen & Kimberly L. Mukobi. 1999. Scale-model comprehension by chimpanzees (*Pan troglodytes*). *Journal of Comparative Psychology* 113(4). 396–402.
- Lenninger, Sara. 2009. In search of differentiation in the development of a picture sign. In Eero Tarasti (ed.), *Communication: Understanding/misunderstanding. Proceedings of the ninth congress of the IASS/AIS, Helsinki-Imatra, June 11–17*, vol. 2, 893–903. Imatra: International Semiotics Institute; Helsinki: Semiotic Society of Finland.
- Lessing, Gotthold E. 1964 [1766]. *Laokoon – oder über die Grenzen der Malerei und Poesie*. Stuttgart: Philipp Reclam Jun.
- Lindekens, René. 1976. *Essai de sémiotique visuelle: le photographique, le filmique, le graphique*. Paris: Klincksieck.
- Meltzoff, Andrew N. 1988. Imitation of televised models by infants. *Child Development* 59(5). 1221–1229.
- Moles, Abraham. 1981. *L'image: Communication fonctionnelle*. Brussels: Casterman.
- Morris, Charles. 1971 [1946]. *Signs, language, and behavior*. In *Writings on the general theory of signs* (Approaches to Semiotics 17), 73–398. The Hague: Mouton.
- Parron, Carole, Josep Call & Joël Fagot. 2008. Behavioral responses to photographs by pictorially naïve baboons (*Papio anubis*), gorillas (*Gorilla gorilla*), and chimpanzees (*Pan troglodytes*). *Behavioral Processes* 78(3). 351–357.
- Persson, Tomas. 2008. *Pictorial primates: A search for iconic abilities in great apes* (Lund University Cognitive Studies 136). Lund: Lund University.
- Piaget, Jean. 1967[1945]. *La formation du symbole chez l'enfant: Imitation, jeu et rêve, image et représentation*. Neuchatel: Delachaux & Niestlé.
- Piaget, Jean. 1967. *La psychologie de l'intelligence*. Paris: Armand Colin.
- Piaget, Jean. 1970. *Epistémologie des sciences de l'homme*. Paris: Gallimard.
- Pierroutsakos, Sophia L. & Judy S. DeLoache. 2003. Infants' manual exploration of pictorial objects varying in realism. *Infancy* 4(1). 141–156.
- Pierroutsakos, Sophia L. & Georgene L. Troseth. 2003. Video verité: Infants' manual investigation of objects on video. *Infant Behavior and Development* 26(2): 183–199.
- Poss, Sarah R. & Philippe Rochat. 2003. Referential understanding of videos in chimpanzees (*Pan troglodytes*), orangutans (*Pongo pygmaeus*), and children (*Homo sapiens*). *Journal of Comparative Psychology* 117(4). 420–428.
- Schmitt, Kelly L. & Daniel R. Anderson. 2002. Television and reality: Toddlers' use of visual information from video to guide behavior. *Media Psychology* 4(1): 51–76.
- Sonesson, Göran. 1989. *Pictorial concepts: Inquiries into the semiotic heritage and its relevance for the analysis of the visual world*. Lund: Aris; Lund University Press.
- Sonesson, Göran. 1992. The semiotic function and the genesis of pictorial meaning. In Eero Tarasti (ed.), *Center and periphery in representations and institutions: Proceedings from the conference of the International Semiotics Institute, Imatra, Finland, July 16–21* (Acta Semiotica Fennica 1), 211–156. Imatra: International Semiotics Institute.
- Sonesson, Göran. 1994. Prolegomena to the semiotic analysis of prehistoric visual displays. *Semiotica* 100(2/4). 267–332.
- Sonesson, Göran. 1998. Icon, iconicity. In Paul Bouissac (ed.), *Encyclopedia of semiotics*, 293–297. New York: Oxford University Press.

- 1 Sonesson, Göran. 2004. Perspective from a semiotical perspective. In Rossholm, Göran (ed.),
2 *Essays on fiction and perspective*, 255–292. Bern: Peter Lang.
- 3 Sonesson, Göran. 2007. From the meaning of embodiment to the embodiment of meaning: A
4 study in phenomenological semiotics. In Tom Ziemke, Jordan Zlatev & Roslyn M. Frank
5 (eds.), *Body, language, and mind: Embodiment*, vol. 1 (Cognitive Linguistics Research
6 35.1), 85–127. Berlin & New York: Mouton de Gruyter.
- 7 Sonesson, Göran. 2009. Here comes the semiotic species: Reflections on the semiotic turn in
8 the cognitive sciences. In Brady Wagoner (ed.), *Symbolic transformations: The mind in
9 movement through culture and society*, 38–58. Routledge: London.
- 10 Tomasello, Michael. 1999. *The cultural origins of human cognition*. Cambridge, MA: Harvard
11 University Press.
- 12 Tomasello, Michael. 2008. *Origins of human communication*. Cambridge, MA: MIT Press.
- 13 Troseth, Georgene L. 2003a. Getting a clear picture: Young children's understanding of a
14 televised image. *Developmental Science* 6(3). 247–253.
- 15 Troseth, Georgene L. 2003b. TV guide: Two-year-old children learn to use video as a source of
16 information. *Developmental Psychology* 30(1). 140–150.
- 17 Troseth, Georgene L. & Judy S. DeLoache. 1998. The medium can obscure the message: Young
18 children's understanding of video. *Child Development* 69(4). 950–965.
- 19 Zlatev, Jordan. 2009. Levels of meaning, embodiment, and communication. *Cybernetics and
20 Human Knowing* 16(3–4). 149–174.

Bionotes

21
22 *Alenka Hribar* (b. 1981) is a postdoctoral researcher at the Max Planck Institute for
23 Evolutionary Anthropology (hribar@eva.mpg.de). Her research interest is ana-
24 logical reasoning in great apes and children. Her publications include “Great
25 apes use landmark cues over spatial relations to find hidden food” (with J. Call,
26 2011); “Children's reasoning about spatial relational similarity: The effect of
27 alignment and relational complexity” (with D. Haun & J. Call, 2011); “Understand-
28 ing the functional properties of tools: chimpanzees (*Pan troglodytes*) and capu-
29 chin monkeys (*Cebus apella*) attend to tool features differently” (with G. Sabbatini
30 et al., 2012).

31
32 *Göran Sonesson* (b. 1951) is a professor at Lund University
33 (goran.sonesson@semiotik.lu.se). His research interests include visual semiot-
34 ics, cultural semiotics, evolution and development of semiosis, and epistemology
35 of semiotics. His publications include “The view from Husserl's lectern: Consider-
36 ations on the role of phenomenology in cognitive semiotics” (2009; “New consid-
37 erations on the proper study of man – and, marginally, some other animals”
38 (2009); “Semiosis and the elusive final interpretant of understanding” (2010);
39 and “Semiotics inside-out and/or outside-in: How to understand everything and
40 (with luck) influence people” (2012).

Josep Call (b. 1966) is a senior scientist at the Max Planck Institute for Evolutionary Anthropology <call@eva.mpg.de>. His research interests include primate cognition, animal cognition, and cognitive evolution. His publications include “Apes save tools for future use” (with N. Mulcahy, 2006); “Chimpanzees are rational maximizers in an ultimatum game” (with K. Jensen & M. Tomasello, 2007); “Monkeys like mimics esis” (with M. Carpenter, 2009); and “Methodological challenges in the study of primate cognition” (with M. Tomasello, 2011).

1
2
3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40