

Quotation, Demonstration, and Iconicity*

Kathryn Davidson
kathryn.davidson@gmail.com

manuscript draft to appear in *Linguistics & Philosophy*
14 May 2015

Abstract

Sometimes form-meaning mappings in language are not arbitrary, but iconic: they depict what they represent. Incorporating iconic elements of language into a compositional semantics faces a number of challenges in formal frameworks as evidenced by the lengthy literature in linguistics and philosophy on quotation/direct speech, which iconically portrays the words of another in the form that they were used. This paper compares the well-studied type of iconicity found with verbs of quotation with another form of iconicity common in sign languages: classifier predicates. I argue that these two types of verbal iconicity can, and should, incorporate their iconic elements in the same way using event modification via the notion of a context dependent demonstration. This unified formal account of quotation and classifier predicates predicts that a language might use the same strategy for conveying both, and I argue that this is the case with role shift in American Sign Language. Role shift is used to report others' language and thoughts as well as their actions, and recently has been argued to provide evidence in favor of Kaplanian (1989) "monstrous" indexical expressions (Schlenker 2014a-b). By reimagining role shift as involving either (i) quotation for language demonstrations or (ii) "body classifier" (Supalla 1982) predicates for action demonstrations, the proposed account eliminates one major argument for these monsters coming from sign languages. Throughout this paper, sign languages provide a fruitful perspective for studying quotation and other forms of iconicity in natural language due to their (i) lack of a commonly used writing system which is otherwise often mistaken as primary data instead of speech, (ii) the rich existing literature on iconicity within sign language linguistics, and (iii) the ability of role shift to overtly mark the scope of a language report. In this view, written language is merely a special case of a more general phenomenon of sign and speech demonstration, which accounts more accurately for natural language data by permitting more strict or loose verbatim interpretations of demonstrations through the context dependent pragmatics.

* Many thanks to the Yale Semantics Reading Group, Ivano Caponigro, Diane Lillo-Martin, Beatrice Santorini, Philippe Schlenker, Sun-Joo Shin, Ronnie Wilbur, Sandro Zucchi, and audiences at NELS 2014 at MIT, the 2015 LSA in Portland, U. Chicago Center for Gesture, Sign and Language, Stanford Language and Cognition Group, NYU Semantics Group, and Harvard Linguistics for helpful suggestions, and two anonymous reviewers for L&P whose feedback significantly improved this manuscript. Any responsibility for mistakes is, of course, my own.

1. Introduction

We use language to talk about the way that the world (or another world) is. Since one of the things that exists in our world is language, we sometimes want to *use* language to talk *about* language. Many times the languages being used and being talked about have the same form (e.g. English), so we need to find ways to distinguish the language used and the language mentioned. In spoken language we often do this by changing our tone of voice, or prosody, or making use of “air quotes”, and in written language we use quotation marks. We can do this to talk about the general properties of the language (1) or about a specific token of another’s use of the language (2). In sign languages, we can make use of a strategy called *role shift*, in which the signer breaks eye gaze and may move his/her body to signal that the words used belong to somebody else (3).¹

(1) “Linguistics” has eleven letters.

(2) This morning John said “I’m happy.”

(3) JOHN_a SAY IX₁ HAPPY^a ‘John said “I’m happy.”’ (American Sign Language (ASL))

The above are all strategies for indicating that the language in quotes/role shift is a *performance* of some sort: the speaker (or writer, or signer- let us consider these all to be the same for now) is not merely reporting the content of what someone said, but also something about *how they said it*. Consider that (1)-(3) are not equivalent to (4)-(6): replacing the context of the quote in (1) with a paraphrase changes its truth value (4), while the direct reports in (2)-(3) tell us something that the indirect reports in (5)-(6) do not, namely, something about the shape of John’s utterance, such as the fact that his utterance contained two words, and the form of those words.

(4) “The study of language” has eleven letters.

(5) John said that he is happy.

(6) JOHN SAY IX-3 HAPPY. ‘John said that he’s happy.’

In a sense, then, we can say that direct speech reports (the material in quotation marks and/or role shift) are interpreted *iconically*: they take on the form of what they are meant to represent.

Until very recently, iconicity has played a limited role in the study of spoken and written language. From the earliest semiotic work by Saussure (1916) and Pierce (1931), the arbitrariness of form-meaning pairings in spoken language has been emphasized, in contrast to any iconicity. This focus on arbitrariness is reinforced when it is impossible to guess what someone is saying in a foreign language, a point further driven home in every vocabulary-based test or language trivia game in which the meanings of words must be guessed. Thus, when the spoken language linguist is asked for examples of iconic language, she will typically mention exotic words such as animal sounds (“meow”, “quack”) and other onomatopoeia (“crack”, “whirr”). Perhaps she will point to recent work demonstrating iconicity at the level of phonemes, with, for example, “smallness” or “nearness” being associated with front vowels or palatal consonants (Hinton et al. 2006, Haynie et al. 2014, a.o.). These, however, feel rather like exceptional cases of iconic language. One of the main arguments of this paper is that iconicity

¹ The “a” index refers to a shared locus for John and the source of role shifted perspective. In other words, the subject John and his utterance are associated to a similar area in signing space.

need not be so marginalized in spoken language, since direct language reports like those in (1)-(3) are pervasive: from newspapers to campfire stories, one frequently wants to report someone's language in the way that they said it, preserving at least some aspects of the manner in which they said it. Furthermore, it seems to be no accident that language reports are one of the most pervasive forms of iconicity in spoken and written language: the way in which language is made (with the mouth, or on a typewriter, etc.) is especially conducive to mimicking another's spoken or written language compared to non-linguistic sounds (e.g. wind howling), and so it should not be surprising that much of the iconicity present in spoken and written language is found in direct language reports.

Sign languages occur in a different mode, with different articulators (hands, face, body) and different receptors (eyes, or the hands for deaf-blind signers) than spoken or written language. They happen, in fact, in a modality in which it is quite easy to convey iconic information beyond language, such as how someone threw a ball, a path someone drove down a road, or the manner in which they used a hammer (all of which are frequently conveyed with the co-speech gestures of the hands in spoken English). Given the modality difference between sign and speech, it is perhaps not surprising that sign languages include significant iconic language. In fact, many people imagine sign languages to be more iconic than they actually are, mistakenly assuming that signs are nearly always iconic and/or gestural. The strongest version of this belief is disproven immediately by anyone who doesn't know a sign language and is faced with someone signing to them: they surely realize that signs are not *transparent*, i.e., it is not immediately obvious what signs mean. This is quite unlike simple gestures that accompany speech, or even pantomime, in which anyone can understand the meaning given the appropriate context. There are, however, significant numbers of *translucent* sign for which the motivation is usually available as soon as their meaning is known, such as American Sign Language (ASL) signs MILK² (which looks like someone milking a cow's udder with one hand) and VOTE (which looks somewhat like the action of stuffing a ballot in a box). In their foundational work on sign language processing (1978), Klima and Bellugi reported that non-signers could not guess the English translation of 81 of 90 common ASL signs presented, but after being given the translation, for approximately half of the signs non-signers were able to make the connection between the sign and the meaning (e.g. that the sign VOTE was related to the act of stuffing a ballot). Consider, though, that even *with* translations, about half of these very common signs had opaque meanings and origins. Moreover, sign languages make use of iconicity in different ways for the same sign, as in the sign for *tree*, for which the ASL sign uses the arm and hand to stand for the entire tree and the fingers the branches, Danish sign language outlines the branches, and Chinese sign language outlines the trunk. Thus, while iconicity is present in single signs in some ways more obviously than in spoken languages (especially in non-arbitrary, translucent and transparent signs) it is nevertheless a more limited aspect of the lexicons of the world's sign languages than most non-signers expect.

Above the level of individual words, we saw iconicity in direct language reports above in both English (2) and ASL (3). Iconicity through language reports is found, then, in both language modalities. In addition, another type of iconic construction in sign languages above the level of the individual word can be found in so-called "classifier predicates" (also known by various

² Videos of single lexical signs can be viewed in any online dictionary of American Sign Language

other names depending on the analysis they are given, including “classifier constructions” or “depicting verbs”). It is through these predicates that iconic information is communicated about, for example, how a ball was thrown, the path of a car, or how to use a hammer. These types of predicates are found across the world’s (historically unrelated) sign languages (Zwitserslood 2012), but have received little attention in the formal linguistics literature because of the difficulty of integrating what many have called their “gestural” components. In a recent presentation, Zucchi et al. (2012) proposed an analysis of these iconic constructions as event modifiers in which the classifier predicate includes a demonstration of an event (such as driving, hammer, etc.). Interestingly, the very same notion of demonstration has also taken on a life in the literature on quotation in spoken languages, where quotation is argued to be much more flexible than simple verbatim mentioning of another’s language (Clark and Gerrig 1990). In both cases, speakers (and writers and signers) can use language to describe something while *at the same time* conveying that the form of language should be interpreted iconically, as long as they have a way to mark the iconic portion so that their interlocutor knows the scope of the language that is used performatively.

This paper unifies these two lines of work, suggesting that direct language reports can, and should, be analyzed in the same way as the more obvious forms of iconicity seen in sign language classifiers, as event descriptions. In Section 2 I begin by discussing the unfortunate influence that studying written language has had on ignoring aspects of natural language speech reports that cannot be conveyed on the page, and urge a sharp turn towards separating written from spoken language in the study of natural language quotation. By doing so, the performative/demonstrational aspects of quotation become more apparent, shedding the *verbatim assumption* (Clark and Gerrig 1990) found in literature on quotation and supporting a demonstration/performance view of quotation. In Section 3 I discuss the event description-by-demonstration view of classifier predicates in sign languages. I provide a series of new arguments in favor of this general approach coming from the structure of ASL as well as ASL/English bilingual language, highlighting similarities between quotation and classifiers.

In Section 4 I apply the event description analysis of classifiers in Section 3 to a puzzle in sign linguistics having to do with role shift called “constructed action” or “action role shift”. Action role shift involves what might be called an “action quotation”: it uses role shift, as in (3) above, but entails that someone performed an action, not that they talked or thought about doing it. I argue that these are actually classifier constructions, and that not only can the event description view of classifiers account for their iconic behavior, but that these should be *expected* in sign languages given that classifiers and quotations both involve demonstrations. Moreover, in this analysis of action role shift (contrary to Schlenker 2014a-b) it is possible to analyze such constructions without resorting to Kaplanian “monsters” that allow indexical expressions like *I*, *here*, *tomorrow* to shift their context of interpretation, undermining one major argument in favor of such indexical monsters coming from sign languages. Finally, in Section 5 I revisit the issue of quotation in spoken language with increased emphasis on performance and decreased emphasis on verbatim representation. I discuss how neither traditional “classifiers” nor “action reports” involve making use of anything that cannot already occur in the spoken language modality through predicates similar to the English “be like”. In highlighting “be like”, I am following Lillo-Martin 2012 in suggesting that linguistic and philosophical research on quotation has been too focused on the written language modality, which has much more stringent requirements concerning the verbatim representation of the reported language and the verbs that

can introduce it. Instead, I suggest that one can offload the level of detail of verbatim representation to a Gricean pragmatic component, while retaining a simple semantics of event modification.

Throughout the paper, we see that sign languages provide a fruitful lens by which to study quotation and general iconicity for three reasons: (i) the sign language literature has more well-developed theories of iconic communication; (ii) sign languages often directly mark the scope of language reports via role shift; (iii) sign languages do not have widely used writing systems, providing a clear example of natural language unburdened by the confound of writing that limits representation to characters on a page (whether alphabetic, ideographic, etc.).

2. On Quotation in Writing versus Natural Speech

The semantic contribution of quotation has received significant attention throughout the previous century both in linguistics and philosophy, based on the compositional puzzle of distinguishing *used* language from *quoted/mentioned* language. A traditional view attributed to Quine (1940) is that quotation marks are a name-forming operator: they act as a function that takes whatever is inside of them as an argument and return the name of that argument. So, (1) would be interpreted as ‘the name of linguistics has 11 letters’, while (2) would be something like ‘John said a sentence, the name of which is “I am happy”’.

- (7) “Linguistics” has 11 letters.
- (8) This morning John said “I am happy.”

This name-based theory of quotation accounts for the generally robust pattern in English that linguistic material outside the quotation cannot interact in a structural relationship with the material inside the quote. For example, one cannot interpret a *wh*-word in the matrix clause as if it were in the quotation (9), and negation in the matrix clause cannot license a negative polarity item in a quotation (10). Moreover, all of the indexical expressions whose meaning depends on the context (e.g. *I*, *here*, *tomorrow*) fail to take their meaning from the current context of speech: “I” in (11a) is not interpreted as the speaker, but as John, whereas without a quote in (11b), one would need to use a third person pronoun, and “I” must refer to the speaker.

- (9) a. #What did John say “I put ____ there”?
b. What did John say he put ____ where?
- (10) a. #John didn’t say “I ever ate pie.”
b. John didn’t say he ever ate pie.
- (11) a. John_i said “he_{*}/I_i ate the pie.”
b. John_i said that he_i/I_{*}_i ate the pie.

Under this classical name view, the speaker doesn’t *use* quoted material in the typical sense that language is used to describe the world, but rather just *mentions* it, i.e. utters the name of an utterance that someone had previously made, verbatim. Here a quotation is a single atomic unit that is not subject to further linguistic structure and may not even be processed as such.

Despite the appeal of such a simple atomic theory, philosophers beginning with Davidson's influential paper (1979) have noted several disadvantages to this "name" analysis of quotation. For one thing, names of things in the world are generally arbitrary, so there should be no reason to expect a regular relationship between a term within quotes and the term without quotes. For example, the name of the field of linguistics could have just as easily been spelled *T-e-d* instead of the same word with quotes "*L-i-n-g-u-i-s-t-i-c-s*". Clearly, this can't be quite right since we can put quotations around things that we have never heard before and they are able to be interpreted with no more difficulty than previously quoted items. In other words, the relationship between quote and quoted *is not arbitrary*, while names typically are.

In an abrupt departure from proper name based theories, Davidson proposed instead that a quotation involves a demonstrative pronoun. Under this view, a demonstrative (e.g. *this*, *that*) exists covertly in the linguistic form and points to/indicates its referent (the quote). Examples (1)-(2) are paraphrased in (12)-(13) according to their underlying representation in the demonstrative view.

(12) linguistics ←**that** has 11 letters.

(13) I am happy ←**that** John said.

Davidson himself notes (1979), "Here, the quotation isn't a communicative act, it is something like a gesture or like an object in the world." Like the proper name theory, Davidson's demonstrative theory expects the token to be a verbatim representation of the previous utterance, but the link between quoted and quote is no longer arbitrary.

More recent accounts of quotation have focused on cases in which the relationship is even less arbitrary: sometimes the quotation actually participates in the language, i.e. is both used and mentioned. Why would this be beneficial? Consider the naturally occurring example (14): the author of this news piece is both providing the words that Stiviano's lawyer used ("peppercorn of a fact") while incorporating the meaning of the expression into the sentence; replacing the quoted material with a demonstrative is ungrammatical (15). (Note that (14) might arguably even contain a negative polarity item "peppercorn of a fact" licensed by the negation in the main clause, in contrast to (10)) Further oft-cited cases are Davidson's own (16) and Abbott's (17), in which the quotation is not even a syntactic constituent.

(14) Stiviano's lawyer has not denied the part about the gifts, although he says there is not a "peppercorn of a fact" that any fraud was involved.

(New York Times May 1, 2014,
on basketball team L.A. Clippers' owner Donald Sterling)

(15) "peppercorn of a fact" ←**that** Stiviano's lawyer has not denied the part about the gifts, although he says there is not a ____ that any fraud was involved.

(16) Quine said that quotation has "a certain anomalous feature." (Davidson 1979)

(17) She allowed as how her dog ate "strange things, when left to its own devices."

(Abbott 2005)

The participation of language into some of the structure of the main sentence, while still conveying something about the original form of the quotation, has lead to recent proposals that a quotation makes two separate contributions to meaning simultaneously: one descriptive and one

more performative (Potts 2007). This is again a departure from the perhaps more intuitive idea that the semantic contribution of a quotation is only that someone used the words in the quote, and if someone said an utterance, pragmatics tells us that they also meant to convey the content of that utterance (von Stechow 2010), though it may be necessary for both aspects to be part of the semantics for full empirical coverage (e.g. (14)-(17)). Maier (2014) uses a different implementation via dynamic semantics, but his analysis shares this newer intuition that there is both a use and a mention aspect to quotation, and, importantly, that quotations cannot be simply unanalyzable chunks of material opaque to the rest of the sentence.

On natural language

At this point, it is worth noting that most examples in the existing literature are based on written/textual language and not on naturally acquired face-to-face language. Many of the examples given above work also as grammatical speech, but distinguishing the two becomes important for understanding the sign language data presented later in this paper since there is no widely used written form of ASL (and so spoken, not written, language is the appropriate comparison). By this I do not mean to imply that quotation is only a written language phenomenon, but rather that writing may differ in some ways from speech when it comes to quotation, and many natural language linguists and philosophers would claim to be more interested in the latter. In fact, the preference for written language versus speech in the philosophy and linguistics of quotation seems to vary quite a bit in the literature: as one example, the very interesting analysis of “unquotation” in Shan (2010) of examples like “the place that [Obama] visited” could at least on the surface be argued to be a mostly text-based phenomenon.³ One may, of course, choose to define quotation as what occurs in written language, but I would argue for the purposes of understanding what is happening in sign languages and in naturally acquired language more generally that spoken language is the most reasonable comparison.

One consequence of taking natural language to be a primary source of data for quotation/direct language reports turns out to be a greater focus on examples with verbs not of saying, but of overall imitation, as in (18). On the one hand, this kind of embedding predicate (e.g. *be like*) follows the traditional diagnostics for quotation described earlier such as a lack of wh-extraction (19), lack of NPI licensing (20), and shifted indexicals (“I” in all examples refers to the little girl, not the speaker), suggesting that it is also not porous to the outer structure in the same way as traditional quotation.

- (18) The little girl was like “[with a beaming smile] I got an A”
- (19) *What was little girl like “[with a beaming smile] I got ___”?
- (20) *The little girl was never like “[with a beaming smile] I got any A.”

On the other hand, *be like*, *be all*, and other attitude embedding verbs frequently employed in spoken language are much less likely to follow the verbatim assumption: it seems quite possible

³ Importantly, both Shan 2010 and Maier 2014 do extend analyses of unquotation and/or mixed quotation to other examples that may not be explicitly marked as unquotation (such as non-constituent quotation), and in doing so do not seem to be restricted in principle to the written language modality.

that the little girl did not use the exact words reported in (18), but rather said something that conveys the same content, and said it in a similar way as the speaker (via words, but also intonation, expression, etc.) This intuition will be the key to the remainder of this paper.

Quotation as Demonstration

One of the most drastic departures from the verbatim assumption was put forth by Clark and Gerrig (1990), who analyze quotation as a *demonstration*, i.e. a kind of *performance*. They quite explicitly took natural language phenomena as their subject instead of the written form. Under a demonstrational theory of quotation (unfortunately named quite similarly to Davidson's *demonstrative* theory, but different in content!), the speaker performs some aspect of an utterance or other action. In written language, demonstration is nearly impossible to distinguish from a verbatim mention of language, but in spoken language, speakers can use all manner of vocal inflection, affect, facial expressions, gestures, and words to create a demonstration of another's behavior from a separate event (see also Tannen 1989). The most canonical examples involve events of saying, but this need not be so: consider (21)-(24).

(21) My call was like "feed me!"

(22) Bob saw the spider and was like "ahh! [in a scared voice]."

(23) Bob saw the spider and was like "I can't kill it!"

(24) Yesterday I saw Mary studying for her finals, and she was all "I'll never be prepared"

Clearly, the cat in (21) did not use the English words "feed me," but rather, the speaker is demonstrating some aspect of the cat's behavior in the event of the cat's asking for food, or imitating the cat's thoughts that most likely did not take the form of English words. Similarly, Bob in (22)-(23) may or may not have actually used as many words (none, or four), but the sentences are false if Bob didn't act scared at all (22) or didn't indicate that he wasn't willing or able to kill the spider (23). Even facial expressions can be part of the demonstration, such that the speaker of (24) is indicating that Mary had an upset face when she was complaining about not being prepared by using the same facial expression in the demonstration (here the notation of the duration of facial expressions is borrowed from Schlenker's (2014a-b) use of the same in sign languages). In a demonstration, some salient properties (determined by the context of utterance) of the original event are replicated in the demonstration.

A clear downside to analyzing speech reports as some kind of performance or demonstration in which at least some --but not all-- aspects of the speech report needed to be accurate/verbatim is that it is difficult to make precise the required mappings between the original context and the speech report. How much of the demonstration/report should be interpreted as part of the original event? Clark and Gerrig provide numerous arguments that, in fact, a surprisingly flexible theory of demonstration is required for natural language quotation/speech reports. As they note, sometimes the form of the words that one uses is the salient part of the demonstration and will be interpreted as having been the same in the event that is being reported; this is usually the case with attitude verbs like *say*, *uttered*, *whispered*, etc. Sometimes, on the other hand, it will be other aspects of the speaker's behavior, as in the *be like* examples in (21)-(24) above, that are meant to be communicated. A theory that hard codes the relationship between the original token and the quotation will be doomed to fail. Importantly, while this lack of hard-coding may be

more obvious for *be like*, it also holds for *say*: (25)-(26) could both be describing the same event, but the relevance of the original language to the new context of the speech report will affect whether or not the speaker intends for it to be conveying the original language or not. For example, the speaker might follow up (25) with a comment about how sad it was that John was using English in France, or about how happy she was to hear someone speak English, but she might just as well have used (25) to convey the same information as (26) if the language that John was using was not relevant.

- (25) We met under the Eiffel tower, and the first thing he said was “my name is John.”
 (26) We met under the Eiffel tower, and the first thing he said was “je m’appelle John.”

Importantly, the aspects of the demonstration that should be considered to be part of the intended interpretation vary by context: the mapping from the demonstration to its interpretation should thus fall within the explanatory domain of pragmatics, not semantics.⁴

Demonstrations in speech and action reports in English

Missing from Clark and Gerrig’s (1990) proposal is a formal compositional semantics for demonstrations. I propose that our semantics include in the ontology a demonstration type *d*. The idea stems from Potts’ (2007) category *S*, a new type for a natural language object (a “Sentence”). I suggest that *d* is more flexible and can account for more natural language data beyond the words used, which is essentially what an *S* is bound to do: one might use *d* in a proper superset of the uses of *S*. The key will be to define *demonstration-of* (shorthand: *demonstration*) as a predicate that takes demonstrations (type *d*) and events (type *e*) as its arguments, which can be lexicalized in English via the expression *like* (27), inspired by Landman and Morzycki (2003)’s analysis of similar expressions in German and Polish. To foreshadow later sections in this paper, this is quite close to a semantics proposed by Zucchi et al. (2012) for classifiers, and, I will argue, should be extended to sign language *role shift* as well, both of which will have semantics similar to the English demonstrational modifier *like* as used below. Some relevant properties of speech reports (i.e. quotations) are given in (28), while (30) provides a compositional account of the English speech report in (29) using a Neo-Davidsonian event-based semantics.

- (27) Definition: a demonstration *d* is a *demonstration of e* (i.e. *demonstration(d, e)* holds) if *d* reproduces properties of *e* and those properties are relevant in the context of speech
 $[[\text{like}]] = \lambda d \lambda e [\text{demonstration}(d, e)]$
- (28) *Properties* of a speech event include, but are not limited to words, intonation, facial expressions, sentiment, and/or gestures
- (29) John was like “I’m happy”
- (30) $[[\text{“I’m happy”}]] = d_I$ (a particular demonstration involving two words, of type *d*)

⁴ For further arguments for a pragmatic approach, see also Maier 2014, who takes quotation to be verbatim, but where “verbatim” is modeled as a context-dependent relation of similarity between form.

$$\begin{aligned}
[[\text{like}]] &= \lambda d \lambda e [\text{demonstration}(d, e)] \\
[[\text{like} [\text{"I'm happy"}]]] &= \lambda e [\text{demonstration}(d_1, e)] \\
[[\text{be like "I'm happy"}]] &= \lambda x \lambda e. [\text{agent}(e, x) \wedge \text{demonstration}(d_1, e)] \\
[[\text{John was like "I'm happy"}]] &= \lambda e. [\text{agent}(e, \text{John}) \wedge \text{demonstration}(d_1, e)] \\
\text{Existential closure (Heim 1982, Kratzer 1998)} &= \exists e. [\text{agent}(e, \text{John}) \wedge \text{demonstration}(d_1, e)]
\end{aligned}$$

Because the demonstration by the speaker/writer of (29) of “I’m happy” is itself an event, we might want to eliminate type d altogether in lieu of simply the demonstration predicate taking two events: $\text{demonstration}(e_1, e_2)$. However, it will likely be counterintuitive for some readers to have a sequences of written marks on a page, for example: “-I-‘-m-_-h-a-p-p-y-”, be an event that can demonstrate an event of another’s writing those words. Nevertheless, we want to include written quotations as a type of demonstration. Furthermore, demonstrations would be proper subsets of events: there are events without communicative purpose that would not be demonstrations of other events. Therefore, I will continue to use d as a new type for (spoken, written, signed, gestural, etc.) demonstrations, but the reader who prefers can conceive of these as a special kind of event.

The “be like” quotative takes demonstrations of all kinds, but some verbs of quotation are more explicit about the type of event (e.g. a “saying” event) and so build that into the lexical semantics, as in (31)-(32). *Say* requires that the quotation be a demonstration specifically of the *words* used, because words are the relevant properties of a saying event. In this way we arrive at the verbatim requirement on quotation through a combination of lexical semantics (saying events involve words) and pragmatics (if d_1 is a demonstration of a saying event, the relevant properties are words).

(31) John said “I’m happy”

$$\begin{aligned}
(32) \quad [[\text{I'm happy}]] &= d_1 \\
[[\text{say}]] &= \lambda d \lambda x \lambda e [\text{agent}(e, x) \wedge \text{demonstration}(d, e) \wedge \text{saying}(e)] \\
[[\text{said "I'm happy"}]] &= \lambda x \lambda e. [\text{agent}(e, x) \wedge \text{demonstration}(d_1, e) \wedge \text{saying}(e)] \\
[[\text{John said "I'm happy"}]] &= \lambda e. [\text{agent}(e, \text{John}) \wedge \text{demonstration}(d_1, e)] \\
\text{Existential closure} &= \exists e. [\text{agent}(e, \text{John}) \wedge \text{demonstration}(d_1, e) \wedge \text{saying}(e)]
\end{aligned}$$

Although *be like* and *say* are common introductions for language reports, demonstrations need not be used only in language reports. Examples of demonstrations that aren’t based on English words can be seen in (33)-(34), with a compositional account of (33) given in (35).

(33). Bob was eating like [gobbling gesture with hands and face].

(34) The flowers bloomed like [gesture of hands].

- (35) Bob was eating like [gobbling gesture].
 $[[\text{gobbling gesture}]] = d_1$
 $[[\text{like}]] = \lambda d \lambda e [P(e) \ \& \ demonstration(d,e)]$
 $[[\text{like} [\text{gobbling gesture}]]]$
 $= \lambda e [P(e) \ \& \ demonstration(d_1,e)]$
 $[[\text{eat}]] = \lambda e. [eating(e)]$
 $[[\text{eats like} [\text{gobbling gesture}]]]$ (Intersective predicate modification)
 $= \lambda e. [eating(e) \wedge demonstration(d_1,e)]$
 $[[[\text{agent}] \text{ eats like} [\text{gobbling gesture}]]]$ (Introduction of external argument)
 $= \lambda x \lambda e. [agent(e,x) \wedge eating(e) \wedge demonstration(d_1,e)]$
 $[[\text{Bob eats like} [\text{gobbling gesture}]]]$
 $= \lambda e. [agent(e,bob) \wedge eating(e) \wedge demonstration(d_1,e)]$
 Existential closure
 $= \exists e. [agent(e,bob) \wedge eating(e) \wedge demonstration(d_1,e)]$

The conjoining of the *like* phrase (“like [gobbling gesture]”) and the main predicate *eat* returns the intersection of events that are both eating events done by x and events that d demonstrates through a standard rule of intersective predicate modification (Davidson 1967, for recent detail Morzycki 2014).⁵ This is important because if the verb *eating* is omitted (36), then there is no entailment that the event was an eating event, but there is an entailment that there was some event that had Bob as its agent and the gobbling gesture is a demonstration of that event. In other words, the demonstration in (35) is licensed by *like*, not *eating*. Note that this contrasts with (31) and even (36) in which quotative verbs introduce demonstrations as arguments; the demonstrational phrase in (35) is an optional modifier.

- (36) Bob was like [gobbling gesture].
 $[[\text{gobbling gesture}]] = d_1$
 $[[\text{Bob was like} [\text{gobbling gesture}]]] = \exists e. [agent(e,bob) \wedge demonstration(d_1,e)]$

At this point, one might fear that a semantics for speech reports has gone too far out of the domain of compositional semantics via the use of a demonstration, but in fact the forms above can serve as a bridge between the semantics and the pragmatics by allowing interlocutors to interpret different aspects of a demonstration depending on what is relevant in the context via the flexible definition of a demonstration. With regard to truth conditional standing, it seems that while the accuracy of a demonstration and what is demonstrated in a given demonstration can be taken care of by the pragmatics, while the use of a demonstration itself is asserted, not implicated, material. To see this, consider (37).

- (37) He was like “This isn’t fair” [said in a whiney voice].

⁵ This account is an oversimplification of a number of issues specific to manner adverbials, such as whether manners should be added to the ontology (Dik 1975; Piñón 2007, a.o.). If so, the final semantics could include a manner type m and possibly a further definable relation between events and manners R (i). For recent overview of the merits and problems with approaches along these lines, see Gehrke and Castroviejo (2015). Here we stick to basic intersection for simplicity.
 (i) $\exists e. [eating(e) \wedge agent(e,bob) \wedge \exists m. [R(e,m) \ \& \ demonstration(d_1,m)]]$

One might felicitously counter (37) with the statement “No, I don’t think he wasn’t whining when he said it wasn’t fair, I think he was just trying to make the field fair for others.” In other words, the assertion that “ d_i is a demonstration of the event” is available for denial by the interlocutor. Of course, demonstrations themselves are rather flexible, so in many cases one might take issue with the wording (we already know that demonstrations need not be verbatim) or the manner (whining, above).

One final point about demonstration involves its relationship to quotes of words that are also used, as in (38), without attitude embedding verbs like “be like”. On the one hand, one might have the intuition that these are frequent in written text but not natural language. Kasimir (2008) reports mixed findings on whether gestural “airquotes” in a spoken English version of (38) are evidence for an influence of writing on similar spoken language examples. On the other hand, Maier (2014) and Anand (2007) have both argued strongly against the view that such mixed quotations are confined to written language.

(38) Joffrey said he was “honorable”.

In any case, in examples like (38) where words are demonstrated as they are used, it seems unlikely that the demonstration is introduced via the lexical semantics of any of the lexical items. Rather, we can tentatively suggest that the speaker of (38) is just exercising a choice, always available, to demonstrate a word while also using it, which can be implemented in a variety of ways. The interlocutor may realize the word as separately marked, and pragmatically interprets it as not the speaker’s own words, but someone else’s (and thus receives its formal interpretation from the other context). This may be able to be resolved in the pragmatics of marking off the word as special, not in any particular semantics of spoken or written quotation marks.

Further remarks on spoken vs. written language

I close this preliminary discussion of English with a few final words about written language versus spoken language before moving on to iconic demonstrations in sign languages. First, the demonstration theory of quotation is flexible enough to handle speech reports that involve more than just words, but of course it *can* include just words. Since written language is essentially only letters (or at least, the communicative content is typically assumed to be the alphabetic characters, syllabaries, ideographs, etc. and not, e.g. the serifs in a particular font, individual ink drops, etc.), then written language quotations are typically interpreted as demonstrations of another’s words. This, however, is just a special case of the more general phenomenon of natural language demonstration in speech reports.

I expect that the written form of quotation has become the de facto type of quotation for two reasons: first, much of academic discourse takes place in written papers, and so examples like Quine quoted by Davidson appear in our own academic discourse. Secondly, and probably more importantly, the typographic convention of using quotation marks in written languages makes the scope of the quotation immediately visible and easier to have crisp judgments about, seemingly more clear than any suprasegmental prosodic or intonational material in spoken languages. The helpfulness of the quotation mark convention has even spread back into spoken conversation

through the use of air quotes (or “finger dance”) in spoken English⁶ to mark the extent of quoted material in conversation, which seems to provide additional information about the boundaries of the demonstration (in other words, it’s clearer what is intended as used and what as mentioned).

Given these reasons for the frequency of studying written quotation over speech, there are compelling reasons to look at language reports in sign languages like ASL. For one thing, there is not a written form of the language that is widely accepted, and so the conflation between natural language and its written form is less available. Second, many -if not all- sign languages have an especially visible form of marking the boundaries of utterances, discussed earlier in (3) called *role shift*. The semantic/pragmatic contribution of role shift and its use in both speech and action reports will be the focus of Section 4. Before moving to role shift, though, I will focus on the third reason, the fact that iconicity has long been a topic of interest in sign languages. In particular, classifier constructions are a common example of iconic demonstration that can be quite naturally analyzed with an event semantics involving demonstration, and as such will be the focus of the next section, Section 3.

3. Classifiers Predicates in Sign Languages

I have argued for a semantics using event description to incorporate the iconic aspects of demonstrations used in spoken language reports, including both gestural performances such as “eat like [gesture]” as well as more traditional quotations. Sign languages are probably better known for their ability to convey iconic information, and in fact, a similar semantics to that given in Section 2 has been proposed in a talk by Zucchi, Cechetto, and Geraci (2012) to incorporate verbal iconicity in sign languages in so-called “classifier predicates”. This section will focus on classifier predicates, especially an implementation of their semantics with a demonstration analysis via event description. I argue throughout that there are important and previously overlooked similarities between classifiers and quotation that point toward a closely analogous account. In Section 4, we will see that a more unified view correctly predicts that the same kind of linguistic element could be used for both quotation and classifiers, which I argue is the case in sign language *role shift*.

First, to begin to understand classifier predicates we must attend for a moment to sign language phonology. Each ASL sign can be broken into three parameters: a handshape, a location, and a movement. The sign MOTHER, for example, is produced with a 5 handshape, at the chin location, and with a contact movement. Vary any one of these three parameters and the sign is no longer the same sign, i.e., these parameters can be changed to form minimal pairs (Stokoe 1960)(39). All three parameters are produced simultaneously.

- (39) a. MOTHER: **5**(handshape)-**Chin**(location)-**Contact**(movement)
b. FATHER: **5**(handshape)-**Temple**(location)-**Contact**(movement)
c. THINK: **1**(handshape)-**Temple**(location)-**Contact**(movement)

⁶ Brief investigation into this phenomenon indicates that it is not widespread to other spoken languages: French speakers report that they do not gesture « this » nor Japanese speakers gesture 「this」 to indicate quotation. Why a gesture for “this” marking may have caught on in English is, as far as I know, a puzzle.

In classifier predicates, by contrast, the location and movement appear to form a single gestural unit that conveys dynamic information about the spatial arrangement and/or movement of referents (40)-(41).

- (40) Classifier predicate for location of a woman:
1(handshape)-[**short downward motion toward a location a**](location+movement)
 Notated as *CL-1(located in point a)*
- (41) Classifier predicate for movement of a bicycle:
3(handshape)-[**wavy path from a to b**](location+movement)
 Notated as *CL-3(wavy path from a to b)*

Syntactically, classifier predicates can be used as verbs in sentences with overt subjects, as in (42a-b), although ASL also permits null subjects, so classifier predicates can additionally form a sentence on their own when the context provides a salient subject (42c). These predicates can convey, for example, the placement of the woman, man, and bike in (42) and (43), their movement toward each other in (42), and the shape of a table in (44).

- (42) a. WOMAN CL-1(Located in point a). ‘There is a woman here.’
 b. MAN CL-1(Located in point b). ‘There is a man here.’
 c. [Right hand] CL-1(straight movement toward center)
 [Left hand] CL-1(straight movement toward center)
 ‘The two people walk toward each other.’
 Overall interpretation: ‘The man and woman walked toward each other.’
- (43) MAN CL-1(Located in point a). ‘There is a woman here.’
 BICYCLE CL-3(Located in point b). ‘There is a bicycle here.’
- (44) TABLE CL-B(Outline of the shape of the table) ‘There is a table like this.’

Classifier predicates have received little formal semantic analysis despite their typological frequency throughout sign languages of the world: Zwitserlood (2012) reports that all but possibly one sign language productively uses such constructions. The term “classifier” has been commonly used because signer’s handshape in the classifier construction (the part that I notate with “CL”) is dependent on a variety of semantic and physical criteria, not unlike classifiers in spoken languages. For example, the 1 handshapes used in (42) (“CL-1”) are used because the subjects are a human man and later a woman, but when the subject is a bicycle (43) the handshape changes (“CL-3”). Small animals, flat objects like paper, long thin objects, and airplanes are among other categories that receive their own classifiers. Although categories of precisely this kind are common across different types of classifier languages (e.g. Mandarin Chinese), the morphosyntax of sign language classifier predicates pattern most closely with predicate classifiers (e.g. Navajo) in that the classificatory handshapes are part of the verbal, not the nominal, structure (Allan 1977, Aikhenvald 2000, Zwitserlood 2012). There is some iconicity in these classificatory handshapes in American Sign Language: the B handshape for a table is a flat hand and the F handshape for pens and other thin cylinders is shaped to hold thin objects. However, if this were all to classifier constructions we would chalk up these classifiers to merely

being among the translucent, and perhaps in some cases transparent, morphemes in sign languages.

It is, rather, the combination of location and movement parameters of classifiers that is striking and presents a challenge for compositional semantics. This is because the information expressed with the combination of location and movement appears to be infinite and interpreted iconically, more similar to gestures than signs. In fact, there is considerable empirical evidence, outlined in detail by Zucchi et al. (2012), that the discrete linguistic status of the *handshapes* of classifiers is quite different from gestural iconicity found in the *movement+locations* of classifiers. Like other linguistic morphemes, classifier handshapes can be arbitrary, and behavioral studies show that they are interpreted categorically. By contrast, location and movement are interpreted gradiently and iconically in behavioral studies of classifiers and there are seemingly infinite options (Emmorey and Herzig 2003). Classifier predicates are also sometimes called “depicting verbs” (Liddell 2003, Dudis 2004), which has the advantage of taking the focus away from the handshape that does the classifying and focusing instead on the ability of the movement and location parameter combination to “depict” in an analog and iconic way (i.e., to demonstrate).

The separation of handshape as linguistic and movement+location as gestural is not uncontroversial: Supalla (1982) in particular has argued for an extensive morphological decomposition of classifier predicates (including location and handshape) into a large but finite number of affixes. Of course, a finite set of morphemes can give rise to an infinite number of potential descriptions if there are unrestricted ways to combine them and no restrictions on their shape, so one analysis can be difficult to distinguish from the other⁷, but the experimental evidence cited above suggests that there really is something extra-linguistic about the depiction. At the other end of the language-to-gesture spectrum, some have eschewed a formal linguistic analysis of classifier predicates altogether, taking the view that they are essentially non-linguistic (DeMatteo 1977, Macken Perry and Haas 1993, and Cogill-Koez 2000). In general, most linguistic investigations of classifier/depicting that capitalize on the pictorial nature of these verbs do not analyze them as fully part of the linguistic system.

In presentation in 2012, Zucchi and his colleagues proposed an analysis that does incorporate the descriptive/iconic properties of classifiers within a formal system through event modification very similar to that proposed in Section 2. In this view, classifiers are comprised of a linguistic component (a classifying handshape like $cl_{i,vehicle}$ and a light verb such as MOVE) and a nonlinguistic component illustrating the movement of the verb (the location and movement that modify the verb)(45). (Note: I use their notation in (45); in my notation, $cl_{i,vehicle} = Cl-3$).

$$(45) \quad [[cl_{i,vehicle}-MOVE_e]] =$$

$$\text{Syntax: } [IP \exists e [IP pro_i [I' \lambda_i [I' I [VP [V cl_{i,vehicle} MOVE_e]]]]]]$$

$$\text{Semantics: } \exists e [move(x_{i,vehicle}, e) \wedge similar_{S \rightarrow L}(dthat[the\ movement\ of\ cl_{i,vehicle}], e)]$$

⁷ Even here though there is a difference between countably infinite (combinations of morphemes) and uncountably infinite (points in 2 or 3 dimensional space), with experimental evidence by Emmorey and Herzig (2003) seemingly favoring a real number/uncountably infinite interpretation.

I defer from going into too much detail concerning Zucchi et al's, notation, since the details are based on slides from an oral presentation, as a manuscript is still forthcoming. However, the idea here is to have a sense for important similarities between this and the account in Sec. 2, specifically (i) an event-based semantics and (ii) analog interpretation of the gestural movement. Their semantics for a classifier predicate like $cl_{i,vehicle}-MOVE_e$ includes a function *similar* which incorporates the demonstration as context dependent event modification, noting "In the context of utterance of classifier predicates, signers assume that there is a correspondence between regions of the signing space and other regions of real space in which the object in the class denoted by the classifier moves." As they note, "according to this interpretation the classifier predicate denotation is only defined in contexts in which the classifier moves, thereby providing a referent for the indexical term formed by the [demonstration] operator."

I follow Zucchi and his colleagues in assuming that rather than classifier predicates being entirely gestural or being entirely made of discrete morphemes, they are instead semantically light verbs that require a demonstrational argument-- precisely like the quotative predicates that we saw in Section 2. What are the possible verbs that underlie classifier predicates? An exhaustive study is probably beyond the scope of this paper, but a good starting point would be some existing categorizations of types of classifiers based on their handshapes. For example, Engberg-Pedersen (1993) has categorized classifier handshapes into four categories based on what the handshape represents: (i) a whole entity (the subject of the sentence), (ii) the handling of an instrument, (iii) the extension and surface of something, or (iv) a body part/limb. The motions+locations portions of the classifier predicate can then be categorized into four corresponding types: motion, manner, extension, and position (Engberg-Pederson 1993, Zwitserlood 2012). Under a demonstrational analysis of classifier predicates, each verbal root can have a semantics nearly identical to the English *be+like*, except that in addition to being a demonstration of an agent in an event they also allow a demonstration of an agent's *movement* (MOVE), *location* (LOCATE⁸), *manner* (MANIPULATE), and *extent* (EXTEND-TO). In other words, a re-glossing of (42)-(44) could be given as in (46) (intended as a full discourse) and (47)-(48), with a detailed account of (46a) in (49). The match between classifier and subject (i.e. that CL-1 be used for humans, CL-3 for vehicles, etc.) can be noted in the lexical semantics of that classifier either in the semantic structure or possibly (if future work suggests this to be a better solution) as a presupposition, similar to gender in English pronouns and potentially height in ASL pronouns (Schlenker et al. 2013).

- (46) a. WOMAN CL-1-(LOCATE). 'There is a woman here.'
 b. MAN CL-1-(LOCATE). 'There is a man here.'
 c. Right hand: CL-1(MOVE)
 Left hand: CL-1(MOVE)
 'The two people walk toward each other.'

⁸ One might be inclined to consider locations as *states* rather than *events*, but at least most examples of purported examples of these in ASL seem to be easily interpretable as events. For example, descriptive classifiers (showing where things are in an area, or where parts are on an object) all seem to work as eventive (X was placed at Y). Predicates that are uncontroversially stative seem to make rather strange classifiers and are instead fully lexicalized, although this is merely an observation that deserves further empirical investigation beyond the scope of this paper.

Overall interpretation: ‘The man and woman walked toward each other.’

- (47) MAN CL-1-(LOCATE). ‘There is a woman here.’
BICYCLE CL-3(LOCATE). ‘There is a bicycle here.’
- (48) TABLE CL-B(EXTEND-TO) ‘There is a table like this.’
- (49) WOMAN CL1-LOCATE(placing the hand at location a).
[[placing the hand at location a]] = d_l
[[CL-1-LOCATE]] = $\lambda d \lambda x \lambda e [theme(e,x) \wedge human(x) \wedge locating(e) \wedge demonstration(d,e)]$
[[CL-1-LOCATE(hand at location a)]]
= $\lambda x \lambda e. [theme(e,x) \wedge human(x) \wedge locating(e) \wedge demonstration(d_l,e)]$
[[WOMAN CL-1-LOCATE(hand at location a)]]
= $\lambda e. [theme(e,woman) \wedge human(x) \wedge locating(e) \wedge demonstration(d_l,e)]$
Existential closure
= $\exists e. [theme(e,woman) \wedge human(x) \wedge locating(e) \wedge demonstration(d_l,e)]$

Argument alternations

One particularly nice piece of support for iconic event modification in classifier predicates comes from morphosyntactic findings regarding argument alternations. It has been noticed that different classifier types produce different argument structures in classifier constructions, even with the same verb of movement (Schick 1987, Kegl 1990, Janis 1992, Zwitserlood 1996, Benedicto and Brentari 2004). Benedicto and Brentari (2004) provide the examples in (50a-b), which both express the movement of a book via classifier predicates. The only difference is that in (50a) a “whole entity” classifier is used (for the book), while in (50b), a “handling” classifier is used (for the holder of the book). They show that the former contains a single (internal) argument, while the second involves both an (external) agent and an (internal) theme, based on the distribution of various tests for agentivity and argumenthood, such as the quantificational domain of NONE (51), which quantifies over the books in both cases but not the agent.

- (50) a. BOOK CL-B(move down)
‘The book fell down’
b. BOOK CL-C(move down)
‘Someone put the book down’
- (51) a. BOOK CL-B(move down) NONE
‘None of the books fell down (on its side)’
b. BOOK CL-C(move down) NONE
‘S/he didn’t put any book down (on its side)’
#‘Nobody put the book down (on its side)’

The event-based semantics account provides the flexibility to account for these argument alternations in the classifier predicate structure. Furthermore, these alternations suggest that the argument structure of the predicate should arise from the choice of classifier handshape, not the light verb. As such, I suggest that size and shape classifiers take three arguments: a

demonstration, and event, and an experiencer (52), while handling classifiers take four: a demonstration, an event, an agent, and a theme (53).

(52). BOOK CL-B-MOVE.

[[the gestural movement of the hand]] = d_i
[[CL-B-MOVE]] = $\lambda d \lambda x \lambda e [theme(e,x) \wedge flatobject(x) \wedge moving(e) \wedge demonstration(d,e)]$
[[CL-B-MOVE(the movement of the hand)]]
= $\lambda x \lambda e [theme(e,x) \wedge flatobject(x) \wedge moving(e) \wedge demonstration(d_1,e)]$
[[BOOK CL-B-MOVE(the movement of the hand)]]
= $\lambda e [theme(e,book) \wedge flatobject(x) \wedge moving(e) \wedge demonstration(d_1,e)]$
Existential closure
= $\exists e. [theme(e,book) \wedge flatobject(x) \wedge moving(e) \wedge demonstration(d_1,e)]$

(53) BOOK CL-C-MOVE.

[[the gestural movement of the hand]] = d_i
[[CL-C-MOVE]]
= $\lambda d \lambda y \lambda x \lambda e [agent(e,x) \wedge theme(e,y) \wedge chunky(y) \wedge moving(e) \wedge demonstration(d,e)]$
[[CL-C-MOVE(movement of the hand)]]
= $\lambda y \lambda x \lambda e [agent(e,x) \wedge theme(e,y) \wedge chunky(y) \wedge moving(e) \wedge demonstration(d_1,e)]$
[[BOOK CL-C-MOVE(movement of the hand)]]
= $\lambda x \lambda e [agent(e,x) \wedge theme(e,book) \wedge chunky(book) \wedge moving(e) \wedge demonstration(d_1,e)]$
Existential closure
= $\exists e \exists x [agent(e,x) \wedge theme(e,book) \wedge chunky(book) \wedge moving(e) \wedge demonstration(d_1,e)]$

Other alternations stemming from the classifiers (for example, the body part/handling alternation also discussed by Benedicto and Brentari) can build the structure from the lexical entry of the classifier itself, which determines the light verb and is the source for both the argument structure and the presuppositional information about the referent (e.g. a human, vehicle, flat object, etc.).

Support for the demonstration argument from bimodal bilingualism

Another type of data in support of the demonstrational theory of classifiers proposed here is the interesting fact that bimodal (i.e. sign and speech) bilingual children and adults sometimes use sound effects when signing classifiers in spontaneous dialogue. One example collected from a corpus of natural dialogue can be seen in (54), in which a 3-year-old child (bilingual in English and in ASL) uses English words simultaneously with many ASL signs, but uses a sound effect during a classifier construction.

(54) *ASL Sign:* GOLF CL-1(path of ball going up) BALL CL-1(path of ball going up)
English Whisper: golf (sound-effect) ball (sound-effect)
‘In golf the ball goes high up, the ball goes like this...’
(3 year-old bimodal bilingual in spontaneous play with Deaf father,
reported in Petroj et al. 2014)

This is anecdotally reported to be frequent behavior for both bilingual children and adults, but is not a requirement of bimodal bilingual language: another example of this sort of “code blending” involves English words simultaneously with the classifier construction (55).⁹

(55) *English speech*: And my mom's you know walking down.

ASL sign: CL-V(walking down stairs)

(Adult in spontaneous conversation with another bimodal bilingual adult, reported in Emmorey et al. 2008)

While data of this sort cannot be conclusive about the structure of each at the quantitative level (since mixing in either direction is *possible*), it's notable that in (55), English seems to be the dominant language in the interaction with the classifier providing a supporting role. Perhaps the ASL demonstration stood in for what may have been a co-speech gesture in a hearing non-signer. In contrast, the child in (54) is whispering the English lexical items, which has been argued elsewhere to be a form of *ASL-dominant* code-blending (Petroj et al. 2014). It may well be the case that in (54) the signer is taking (most of) the structure of the clause from American Sign Language, which involves a classifier construction, and so the speech stream can merely demonstrate, while in (55) the underlying structure is the English words and the classifier functions more as a co-speech gesture.

Support for the flexibility depending on language context comes from brief inspection of four hours of naturally occurring (free play sessions) of bimodal bilingual child speech (a single hearing child who has Deaf parents, 2 hours of play at age 2;03 years and 2 hours again at age 3;00 years)¹⁰. Of 48 total classifier predicates produced by the child in these four hours, 20 were accompanied by no speech at all (these were all with a Deaf interlocutor, usually the parent), and 14 were accompanied by an English verb (12 of these 14 were with hearing interlocutors). Another 9 occurred with sound effects, and these were with an equal mix of Deaf and Hearing interlocutors. This suggests that it is not a rarity for bimodal bilingual children to blend classifiers with what are essentially vocal gestures in English, but that it is one of many options and depends on the interlocutors and thus perhaps the primary language of the conversation. Only 5 other classifier predicates were produced with other material or were unintelligible. Clearly, there is further interesting empirical work to be done in this domain, and which may prove fruitful for both language and gesture research.

Further issues

One final observation on classifier constructions involves the analysis of examples in which the signer's two hands are used to produce two demonstrations at one time, or when one hand is producing a demonstration and the other hand is signing lexical content that is not necessarily part of a demonstration. An example from Liddell (1980) is given in (56) (using his notation).

⁹ The author is grateful to Karen Emmorey for sharing these observations.

¹⁰ The author is grateful to Diane Lillo-Martin and Deborah Chen Pichler, and Ronice Müller de Quadros for permitting publication of this data from the Bimodal Bilingual Bi-national child corpus. For more information, see Lillo-Martin and Chen Pichler (2008).

- (56) Hand 1: TABLE B-CL _____ (palm down)____
 Hand 2: MAN PAINT[underside of B-CL]
 ‘The man painted the underside of the table’ (Liddell 1980)

Interestingly, Liddell’s own discussion of this example uses vocabulary nearly identical to that used in the philosophical and linguistics literature to discuss quotations: he wrote how this “mixed” construction behaves like the “pure” locative construction with respect to the syntax of interrogative statements (i.e. the placement of *wh*-words in a *wh*-interrogative and question markers in polar interrogative) and general word order. This seems to stem from the fact that in both these classifier constructions and in quotations there is an important aspect of imageability/iconicity that can combine in a variety of ways with the less iconic, descriptive content, and sometimes seems to have properties of both. Further discussion of two-handed examples like (56) is beyond the scope of this paper, but point toward further fertile ground for future work on classifier demonstrations.

The argument put forward thus far has been that classifiers and quotation share important properties and both use demonstrations as event modification. An interesting empirical prediction then arises that there might be languages that unify the two more overtly by using the same linguistic construction for both quotations and classifiers. In the next section I argue that this is precisely the case for *role shift* in ASL and other sign languages. As we will see, viewing role shift as both quotation and as classifier has an additional advantage in providing an account for the seemingly “shifted” interpretation of indexical expressions in some cases of role shift, and by doing so removes one major argument for the existence of so-called “monstrous” (Kaplan 1979, etc.) interpretations of indexicals under sign language role shift.

4. Role Shift as Quotation and Classifier

One of the most well-studied phenomena in the sign linguistics semantics literature, and one that I will claim instantiates both quotation and classifier predicates, is known descriptively as *role shift*. Role shift involves a change in facial expressions and eyegaze and also possibly head position and body lean¹¹ to essentially “become” another character in order to report their language (57), thoughts (58), or actions (59) directly. For notation in this paper, the referent toward which the role is shifted (the other’s perspective) is written with corresponding indices (e.g. *a* in (57)) and the line above glosses for signs shows the duration of the role shift. Like written quotation, role shift overtly marks the boundaries of the report (the “mentioned” language, e.g. “I’m looking at it” in (57)) from the rest of the sentence that is not a report (the “used language”, e.g. “Mary said”).

¹¹ See Cormier, Smith, and Sevcikova (in press) for a recent and very thorough overview of the articulators that participate in the broad category of what they term *constructed action*, and how different options and degrees may indicate subcategories. I’ve chose to continue to use the term *role shift* in this paper for reasons of consistency with existing formal semantic literature on the topic, which has been primarily interested in the *shifted* perspective and semantic consequences.

(57) *Language reports*

(*indirect*)

a. MARY_a SAY IX_a A-WATCH-B. ‘Mary said that she was watching it.’

(*direct*)

b. MARY_a SAY IX₁ 1-WATCH-B_a. ‘Mary said “I was watching it.”’

(58) *Attitude reports*

(*indirect*)

a. MARY_a THINK IX_a HAPPY. ‘Mary thought that she was happy’

(*direct*)

b. MARY_a THINK IX₁ HAPPY_a. ‘Mary thought “I was happy”’

(59) *Action reports*

(*indirect*)

a. MARY_a A-WATCH-B. ‘Mary was watching it’

(*direct*)

b. MARY_a 1-WATCH-B_a. ‘Mary was watching it’

Lexical material under role shift shares a variety of syntactic/semantic similarities with quotation, including the fact that indexical expressions (those words whose meaning is dependent on the context, such as the first person pronoun IX₁ ‘I’/‘me’) is interpreted with respect to a different context than the current speech context. In other words, IX₁ in (57) refers to Mary, not the speaker, just like “I” in its quotation translation. However, there also differences between quotation as traditionally conceived and role shift in sign languages that has led many researchers to differentiate the two.

In this paper I will actually be arguing against most recent literature (Zucchi 2004, Quer 2005, Herrmann and Steinbach 2012) to say that RS actually *is* quotation in language and attitude reports, given the new view suggested here that spoken language quotation is a more flexible notation than what is found merely in written language. I argue that when we consider “natural” (signed and spoken) language, we find more flexibility in the types of demonstrations permitted with “quotation”, more akin to clearly demonstrational quotatives like the English “be like”, and that this is like role shifted language and attitude reports in ASL (Lillo-Martin 1995). Of course, at least one example of role shift is very clearly not quotation of any sort of demonstrated language: action reports (59). What I suggest instead is that these action reports are instead a case of a classifier predicate with an underlying verb of “body imitation”, a classification mostly ignored in recent literature but originally suggested decades ago by Supalla (1982). Since we saw in Sections 2-3 that both quotation and classifiers use demonstrations as event modifiers, it is no longer puzzling that both might make use of the same role shift strategy in ASL, even if they are not both quotation as traditionally conceived. In this section I will focus first on language reports, which I argue are essentially quotation, and second on action reports, which I argue are classifiers, the difference between the two being how the demonstration is introduced (as an argument of a quotative predicate or as a predicate modifier). In the final subsection, I discuss a different proposal by Schlenker (2014) who argues that action reports are a case of “supermonsters” that shift the context for indexical evaluation, and discuss the merits of the current analysis with respect to the behavior of indexical expressions.

Language and Attitude Reports as Quotation

There has been lengthy debate in the sign language semantics literature about whether language reports under role shift are essentially quotation (Supalla 1982, Lillo-Martin 1995, Zucchi 2004, Quer 2005, Herrmann and Steinbach 2012, Huebl 2012, a.o.). Part of the confusion, I suggest, stems from the traditional focus in philosophy and linguistics on *written* language as a source of quotation data, and hence on verbatim reports, which does not always seem to hold for sign language role shift. However, there are at least two additional reasons, as far as I can tell, that one might want to argue that language reports using role shift differ from quotation. The first of these concerns syntactic dependencies, where linguistic material under role shift seems to interact with the non-role shifted part of the sentence, and in doing so does not always behave as single unanalyzable chunks (“pure quotation”). Schlenker (2014a-b) provides a particularly clear linguistic grammaticality judgment paradigm testing three syntactic dependencies in language reports with role shift: wh-movement, negative polarity items, and ellipsis, and so I will focus on his data here. Wh-extraction appears to pose the most striking problem: Schlenker (2014a-b) reports grammatical examples of wh-extraction in which the position associated to the wh-word would otherwise be located in the role shifted clause (60). Crucially, both indexical expressions under role shift in this example, IX-1 (his notation, my IX₁) and HERE, are evaluated with respect to some context that is not the context of speech (IX-1 is interpreted as John, HERE as LA).

(60) *Context: The speaker is in NYC; the listener was recently in LA with John.*

BEFORE IX-a JOHN IN LA [= while in LA]

RS _____

WHO IX-a SAY IX-1 WILL LIVE WITH ____ HERE **WHO**?

‘While John was in LA, who did he say he would live with there?’

(ASL, from Schlenker 2014a)

Since both written and spoken English do not usually permit wh-extraction from inside of a quotation, Schlenker argues that (60) must be a case of indirect, instead of direct, reported speech. However, since both indexical expressions (IX-1 and HERE) differ in their interpretation from typical indirect speech, he suggests that RS is the overt expression of an operator whose purpose is to shift the context for interpretation of indexicals- a so-called “monster”, based on evocative terminology used by Kaplan (1989). I will focus on Schlenker’s analysis of examples like (60) later in this section, but for now, it is enough to notice that there is a divergence with respect to the interpretation of wh- words in quotation and language reports under role shift.

A second source of difference between quotation and role shift has been the variable behavior of indexical expressions, and in particular the fact that in Catalan sign language (Quer 2005) and German sign language (Hermann and Steinbach 2012), some indexical expressions can be evaluated with respect to the context of the speaker while others are evaluated with respect to a different context, even in the same clause under the same extent of role shift. This behavior notably seems to violate the “shift together” constraint proposed for indexical shift by Anand and Nevins (2004) to account for spoken language data. A third difference between quotation and role shift has been argued to be action reports (59), which I will put aside for now and discuss in

Putting aside action reports, we are left with data from *wh*-expressions and mixed indexicals as arguments against a quotational analysis of role shift, but as it happens, these are the areas with the most inter-signer and inter-language variability regarding role shift across sign languages. For example, Schlenker reports that the *wh*-data he found in ASL does not seem to hold for French Sign Language (LSF), where extraction is not allowed with any reported attitudes (both language or other attitudes) that involve role shift. Other research points toward some ASL signers not accepting *wh*-extraction from RS even within ASL, or at least differences depending on the verbs that precede the role shift: while doxastic attitude embedding predicates (e.g. *say*, *argue*, *claim*) don't allow extraction from role shift, proffering attitude embedding predicates (e.g. *imagine*, *think*) do seem to permit it (61)(Koulidobrova and Davidson 2014), although the extent of role shift in these examples seems to vary as well.

- All of the discussed data regarding wh-extraction comes from a very small number of signers, so clearly more work should be done to determine the nature of this variability. On the other hand, as Schlenker himself notes (2014a) and I agree, the status of wh-words certainly leaves open the possibility that role shifted language reports *may* in fact be quotation, since there are many more cases of disallowed wh- traces than we would expect for indirect (non-quotation) reports.

(62) _____RS_a (negative polarity item: ANY)
 *BOB_a NEVER SHOW ANY HEARSOFT
 ‘Bob never showed any sympathy’

(63) _____RS_a (VP ellipsis)
 #IX-2 LOVE OBAMA. JOHN TELL-1 IX-1 NOT
 ‘You love Obama, but John tells me “I don’t [love Obama]”’
 (both from Schlenker 2014b)

23

current variation in wh- behavior among languages and between embedding predicates points toward, I imagine, something more complicated than the verbatim repetition of another's language commonly found in written quotation, although this remains to be seen. However, as we noted in Section 2, spoken quotation shows contextual flexibility with regards to its faithfulness and perhaps signers differ on how willing they are to stretch their demonstrations.

As for cross-language variation found in shifting together indexicals, it seems that the flexible notion of demonstration may, here too, account for the variation found among language reports under role shift in sign languages. For example, the lack of shifting together among variables in Catalan and German sign language be implemented through the demonstration relation as long as the indexicals that shift are iconic for purposes demonstration. This makes the prediction that more iconic indexicals will shift before less iconic indexicals under role shift. Anecdotal data suggest that this may be born out, as sign language pronouns (which are highly iconic) are frequently found to shift in sign languages even when other indexicals like NOW and YESTERDAY do not, with some intermediate examples (e.g. the ASL sign HERE, which is somewhat iconic) leading to intermediate judgments. However, a satisfying answer to this typological prediction will have to be left for future research.

Further restrictions found among signers, languages, and predicates may be due to a variety of factors, whether pragmatic, semantic, or syntactic. However, given the muddled data on wh-movement and the flexibility of a demonstrational account of indexicals, it becomes at least reasonable to consider an analysis of role shift in language and attitude reports as quotation. There is another use of role shift, however, that stands in stark contrast to quotation: action reports, the topic of the next section.

Action Reports as Classifiers

We have so far been ignoring direct *action reports*, as in (59) and below in (64), also called “constructed action” or “action role shift”. While the quotation-like examples we’ve seen so far report someone’s language, role shift actions report someone’s *doing* something, just like non-role shifted clauses. That is, both (64a) and (64b) below entail that the friend performed an action of watching/looking, i.e. did not just think about doing so. Moreover, action role shifts are interpreted iconically: they are used in cases where the signer is trying to demonstrate some aspect of the subjects behavior (their facial expressions, say, or perhaps their attitude), and are only licensed by “iconic” lexical items (Schlenker 2014b), which includes inflecting verbs like WATCH (also glossed sometimes as LOOK-AT) and SHOW in (65).

(64) (Inflecting verb, no action shift)

- _____topic
- a. FRIEND_a OLYMPICS_b a WATCH_b
 ‘My friend watched the Olympics’

(Action report: Inflecting verb with role shift)

- _____topic _____a
- b. FRIEND_a OLYMPICS_b 1 WATCH_b
 ‘My friend watched the Olympics (like this)’

(Lillo-Martin and Quadros 2011)

(65) (*Action report*: Inflecting verb with role shift)

_____rs
IX-b(WOLF) IPHONE SHOW-CL(iconic).

‘The wolf showed the iPhone (like this).’

(Schlenker 2014b)

Action reports appear to constitute a clear case against an analysis of role shift as the sign language version of written quotation, since they involve actions, not words. In other words, there is no way in English to precisely convey (64b) or (65) strictly through the auditory stream without recourse to body movements or gestures. However, unlike when a speaker of English says the verb “watch” while simultaneously looking up with a particular expression or “show” with a gesture, action role shift appears to involve a further linguistic puzzle: the verb in action role shift uses *first person agreement* marking (notated with the “1” subscript on WATCH), even though it is not the signer herself who looked up, but rather the subject of the sentence. In contrast, in English if one says “I watched” with an unusual expression, this must be interpreted as the speaker herself watching, not the subject of the sentence. Although sign language researchers have long noticed the occurrence of first person marking, it has risen to the attention of formal semantics research through recent work by Schlenker (2014a-b), who calls this first person agreement in action role shift a “super monster”—a particularly punchy example of an indexical expression that receives a shifted interpretation in a non-quoted context and can clearly not be analyzed as quotation.

In this paper, I suggest a new stance with regard to action reports and quotation. On the one hand, I suggest that action reports are examples of classifier predicates, not quotation. That is, they are used language, not mentioned language in the way that quotations are. On the other hand, I argued in Sections 2 and 3 that quotations in spoken language and classifier predicates in sign language share an important property of introducing demonstrations as event modifications. It is perhaps not surprising, then, that role shift is the manifestation of two types of event modifying demonstrations in ASL, one for language and one for actions. The connection is indirect, though, via the introduction of demonstrations.

Recall our semantics for classifier predicates, in which they introduce a demonstration that is provided by the combination of movement and location parameters (66). The entire structure is monoclausal: the demonstration is a modification of the main verb LOCATE, EXTEND, etc.

(66). WOMAN CL1-LOCATE.

[[hand at location a]] = d_l

[[WOMAN CL-1-LOCATE(hand at location a)]]

$\exists e. [locating(e) \wedge exp(woman, e) \wedge human(woman) \wedge demonstration(d_l, e)]$

I suggest that the same analysis can be extended to action reports, except that instead of the movement and location parameters providing the demonstration, the signer’s entire body provides the demonstration. This leaves those parameters open to provide a full lexical verb, such as WATCH, SHOW-CL etc. While in other classifier predicates the demonstration is introduced via the classifier, in the case of action reports the demonstration is introduced via the use of role shift, which acts as a modifier like English “like” (67). In a sense, this follows

suggestions by Lillo-Martin (1995, 2012) that the English expression “be like” is the closest equivalent to role shift in sign languages, but instead of a silent main verb “be like” that embeds the action report as a proposition, I tie the two together through the semantics of predicate modification, so that the entire action report is a monoclausal structure. The idea here is that the RS morpheme, in contrast to English “like”, is produced simultaneously with other lexical material, consistent with a tendency toward simultaneous verbal morphology in sign languages versus sequential morphology in spoken languages (Klima and Bellugi 1979). In this way it’s semantically identical to (35) above, *Bob was eating like [gobbling gesture]*, except the manner modifier occurs simultaneously with the verb.¹²

- (67) a. $[[\text{like}]] = \lambda d \lambda e [\text{demonstration}(d, e)]$
 b. $[[\text{rs}]] = \lambda d \lambda e [\text{demonstration}(d, e)]$

- (68). IX-b(WOLF) IPHONE SHOW-CL(icons) rs
 $[[\text{gesturally appropriate handshape and movement}]] = d_1$
 $[[\text{rs}]] = \lambda d \lambda e [\text{demonstration}(d, e)]$
 $[[\text{rs} [\text{gesturally appropriate handshape and movement}]]]$
 $= \lambda e [\text{demonstration}(d_1, e)]$
 $[[\text{SHOW}]] = \lambda e. [\text{showing}(e)]$
 $[[\text{SHOW rs}(\text{gesturally appropriate handshape and movement})]]$
 $= \lambda e. [\text{showing}(e) \wedge \text{demonstration}(d_1, e)]$
 $[[\text{SHOW rs}(\text{gesturally appropriate handshape and movement})]]$ (External arguments)
 $= \lambda y \lambda x \lambda e. [\text{agent}(e, x) \wedge \text{theme}(e, y) \wedge \text{showing}(e) \wedge \text{demonstration}(d_1, e)]$
 $[[\text{IPHONE SHOW rs} [\text{gesturally appropriate handshape and movement}]]]$
 $= \lambda x \lambda e. [\text{agent}(e, x) \wedge \text{theme}(e, \text{iphone}) \wedge \text{showing}(e) \wedge \text{demonstration}(d_1, e)]$
 $[[\text{WOLF IPHONE SHOW rs} [\text{gesturally appropriate handshape and movement}]]]$
 $= \lambda e. [\text{agent}(e, \text{wolf}) \wedge \text{theme}(e, \text{iphone}) \wedge \text{showing}(e) \wedge \text{demonstration}(d_1, e)]$
 Existential closure
 $= \exists e. [\text{agent}(e, \text{wolf}) \wedge \text{theme}(e, \text{iphone}) \wedge \text{showing}(e) \wedge \text{demonstration}(d_1, e)]$

There are a number of supporting data points in favor of a monoclausal analysis involving demonstration. First, I take as a starting point some puzzling data reported by Engberg-Pederson (1993) and confirmed in separate data collection by Schlenker (2014a) in which first person agreement is licensed in the reported action (69b) (and is interpreted shifted to the new context, compared with (69a)) but an overt first person pronoun is not allowed in reported action (69c). Instead, the appearance of a first person pronoun in action reports causes the report to be interpreted as an *attitude report* (69c). In other words, with first person *agreement marking* the interpretation is that the friend did indeed watch the Olympics, while with *an overt first person pronoun* the interpretation is an attitude report, i.e. that the friend thought about watching the Olympics (and may or may not actually have done so).

¹² An anonymous reviewer notes that the proposed semantics predicts the availability of sequential morphology here, that is, the non-iconic verb SHOW followed by a role shift into a gestural depiction of the manner event. These are judged as marginally grammatical in ASL, but far more natural (fitting in with the simultaneous morphology of the language) is to have it simultaneously with the iconic verb.

- (69) a. FRIEND_a ^{topic}OLYMPICS_b _aWATCH_b
 'My friend watched the Olympics'
- b. FRIEND_a ^{topic}OLYMPICS_b ₁ ^aWATCH_b
 'My friend watched the Olympics'
- c. FRIEND_a (SAY) ^{topic}OLYMPICS_b IX₁ ₁WATCH_b ^a
 'My friend was like "the Olympics, I watch"'
 *'My friend watched the Olympics'

This data is quite unexpected in all previous analyses of role shift in formal frameworks, which take the role shifted material to a full clause embedded under some type of attitude predicate or other type of context shifting operator (Zucchi 2004, Lillo-Martin 1995, Quer 2005, Schlenker 2014a-b). The pattern in (69) is all the more puzzling because of how robust signers' judgments on it are: the same shift toward an attitude report interpretation with an overt first person pronoun holds in both ASL and LSF, reported by different authors and who have consulted on the judgments with different native signers—a more robust pattern than, for example, wh-movement, and one in need of an explanation.

The analysis proposed here in which the action verb (i.e. WATCH, SHOW) is the main predicate immediately accounts for the robust ungrammaticality of an overt first person pronoun in (69c) because this is a single clause, and a first person pronoun would be an extraneous syntactic argument: the subject position is already taken by FRIEND. Of course, an immediate problem facing such an analysis is the first person agreement marking on the verb in which a first person character (the friend) is the subject. Here, I suggest that we are overlooking an important aspect of the first person agreement marking in sign languages: its iconicity. When a signer is performing/demonstration an action, the verb *looks like it does when it is inflected for the first person*. (See Meir et al. (2007) for detailed discussion on the unusual properties of the body for conveying another's actions as a subject in sign languages.) No one denies that iconicity is central to reported action: an especially clear example is provided by Schlenker from LSF, which has two verbs that mean *show*: one that is not iconic, and another that is extremely iconic, down to the shape of the hands holding the object that one is showing (70). The latter easily licenses an action report (70b) while the former leads to questionable grammaticality and interpretation (70a) in the very same sentence structure.

- (70) RECENTLY WOLF_b IPHONE FIND HAPPY. SHEEP IX-b b-CALL-a.
 'Recently the wolf was happy to find an iPhone. He called the sheep.'
- a. RS_b ^{topic}IX-b IPHONE SHOW.
 => unclear inference He [= the wolf] showed [or: said/thought he was showing/would show] the iPhone to him.'

- b. RS_b _____
IX-b IPHONE SHOW-CL.

=> the wolf in fact showed the iPhone to the sheep He [= the wolf] showed the iPhone to him.'

(LSF, reported in Schlenker 2014b)

What does it mean for the main verb to be iconic, and why might this lead to separate interpretations? I suggest that the more iconic a verb is, the more able it is to be interpreted simultaneously with a demonstration. Of course, not every use of an iconic verb has to be a demonstration: this is where *role shift* comes into play. Role shift does to verbs in ASL what “like” does to English verbs: it turns them into predicates of demonstration. Interestingly, English seems to actually be *more permissive* in the verbs that can take “like” and demonstrations, but the difference may be due to the fact that in ASL, the verb is *part* of the demonstration (simultaneously produced with the role shift operator) while in English the verb precedes the demonstration, and so need not (in fact, cannot) be part of the demonstration itself.

In addition to accounting for the behavior of indexical expressions and predicting the iconic interpretation of action reports, there is also an historic motivation for the analysis of action reports as classifiers. In Supalla’s (1982) classification, reported actions are actually a subcategory of classifier constructions, called “body classifiers”, in which the body represents an animate entity, i.e., instead of the handshape taking a value of an entity, instrument, etc., the signer’s entire body can take on the value of another person. The idea, then, that I pursue in this paper is not a new one: *action reports are part of the classifier system, and as such should receive a unified syntactic and semantic analysis with classifiers*. Nevertheless, in recent work they have generally been separated, with reported action falling within the proposed explanatory realm of quotation and reported utterances based on their unified use of role shift, while more prototypical classifier constructions are mostly essentially untouched by those working at the interface of formal syntax and semantics. Another motivation behind this separation seems to have been the observation that body classifiers cannot be combined with all of the same motion or location verbs as other classifiers (Supalla 1982, Schick 1990, Zwitserlood 2012), but this may simply be a case of different selectional restrictions on verbs with otherwise similar semantics, a common difference among verbs in any lexicon.

Role shift for both quote and classifier?

At this point, the reader may be convinced that classifiers and action reports both involve demonstrations, while remaining convinced that *role shift* is a single phenomenon and should instead deserve a unified analysis including language reports. If reported action is just a classifier construction, why does it manifest as role shift, the same thing that occurs in reported language? There may be two related answers to this question, involving pragmatic signaling and perspective taking.

Consider, first, quotation in written language. There is certainly nothing iconic about “”, confirmed by other ways that we can signal quotation such as *italics*, **bold symbols**, <<double brackets>>, or 「these」. In spoken languages, one might use air quotes or prosody. These are all arbitrary systems that serve only to set off the quoted text from the main text. In fact, it seems *sufficient* to have any kind of different marking to set off quotation: one might introduce a

Of course, it hardly seems coincidental that body classifiers (aka action reports) make use of the same sort of role shifting nonmanual marking as speech reports, and for that reason the two have been considered to be closely related in previous literature. This may be more coincidental than it seems, though: in reported speech, one becomes the other speaker because someone else spoke; in reported action, one becomes the other actor because one wants to show *aspects of the behavior of the actor in the speaking event*. A nice contrast can be seen in (71), where the signer can use the simple description (71a), an action report to demonstrate aspects of the agent's looking up at the movie (i.e. that he had to raise his head)(71b), or a body part classifier handshape (for the man's head) turning toward a movie screen (71c).

- Both (71b) and (71c) are used for purposes of demonstration, albeit different demonstrations of the same event; to signal the demonstrated component they are set off, with either role shift or classifier handshapes. Thus, it is not the role shift that licenses demonstrations in either language reports or action reports, but rather the predicate modification (in the linguistic structure) that permits a demonstration, which manifests as role shift. Supporting data in favor of some structural differences between the two come from Emmorey and Reilly (1998), who show that Deaf children with Deaf parents master language reports before action reports, although both seem to reach adult-like status quite late (after age 7), as do classifier predicates. Finally, it is notable that Liddell (2003) has suggested a unification of language and action reports through the “mental spaces” cognitive linguistics framework. As far as modeling the demonstration itself goes, such a framework might very well be the most fruitful way to think about what he calls “constructed dialogue” and “constructed action”, and once the demonstration is provided then it is quite consistent with the analysis suggested here for the larger structure of the entire clauses.

Schlenker (2014a-b) has recently proposed an analysis of role shift that also uses the notion of iconicity to unify language and action reports in role shift, but which differs in a number of ways

from the suggestions in this paper. Foremost among these differences is his claim that role shift is semantically a context shifting operator that can apply to an embedded clause, which then changes the interpretation of indexicals within its domain. The role-shift-as-context-shifting-operator view has been proposed by a number of sign language researchers in recent years (Zucchi 2004, Quer 2005, Herrmann and Steinbach 2012), but Schlenker’s proposal is the most explicit in its unification of attitude and action reports, and in its semantic proposal for the licensing of the context shift operator (for example, Quer 2005 focuses on syntax), so it is the account of context shift that I will focus on for purposes of argumentation.

The idea behind the context shifting analysis of role shift comes from reports in spoken languages where indexical expressions like *I*, *here*, or *tomorrow*, which are typically evaluated with respect to the context of utterance, are instead evaluated with respect to some other context. So, for example, Amharic is one language reported to have a first person pronoun that can sometimes also be used to refer to someone other than the subject (see the schematic in (72a)). In English, this is generally only permitted in instances of quotation but is disallowed in indirect speech (72b), but in some languages it seems that they are permitted even in examples that clearly do not involve quotation, as tested by wh-movement, etc.

- (72) Situation to be reported: John says: ‘I am a hero’
 a. Amharic: John_i says that I_i am a hero
 b. English: John_i says that {he_i is/*I_i am} a hero (Schlenker 2003)

The indexical expressions found in these languages have been called “monsters” since Kaplan (1979) claimed that they would not be found in natural language outside of examples of quotation. Faced with these monsters in various languages, a number of proposals have been made that propose various operators that act on linguistic material and have the effect of changing the context in which they are evaluated- they “shift” the context (Schlenker 2003, Anand and Nevins 2004, a.o.).

We saw above that language reports in sign languages show some indications of what we may now call “monstrous” behavior, i.e. permitting shifting indexicals along with non-quotational behavior like wh-movement, but that none of the language report data provide a knock-out argument against a more simple quotational analysis, a view shared also by Schlenker (2014a-b). Action reports, on the other hand, initially appear to be an ideal case of a “super-monster”: impossible to analyze as pure quotation, but clearly a shifted indexical expression in the first person agreement marking. Moreover, they use the same nonmanual marking (role shift) that is used in language reports, and so one might plausibly want to give a unified analysis of the two. With this in mind, Schlenker (2014a-b) proposes that all role shift is the overt realization of the scope of a context shifting operator, and that although the attitude report data in sign languages provides somewhat unclear conclusions, the action report data provides a clear piece of evidence for such an operator and parsimony requires they be analyzed in a unified way.

Recall that there are two key facts that an analysis of action reports must cover: first, at least one seemingly indexical expression (first person marking) has a shifted interpretation with action role shift, and second, that action role shift is limited to *iconic* situations, when the signer is using a verb that is in some way iconic. In Schlenker’s account, the first person indexical occurs

because of the operator that shifts that context of evaluation for indexicals within its domain. This operator is licensed by the use of iconic lexical material (this can include inflectional verbs, classifiers, etc.), which unites the shifted indexical phenomenon with the iconic interpretations. ASL generally permits null subjects, so the material under role shift is proposed to be a clause with a null subject.

Despite the attractiveness of the context shift account in extending an existing idea from spoken languages to new sign language data, I think a few problems lend cause for concern in the application to reported actions and undermine the evidence for their being a “supermonster”. The first is that while first person agreement is happily licensed under reported actions and receives a shifted interpretation, an overt first person pronoun is not grammatical with a reported action interpretation, as we saw earlier in (69). Instead, the use of an overt first person pronoun is re-interpreted as a reported *attitude* interpretation with some type of implicit attitude embedding predicate. Under Schlenker’s account, this surprisingly robust pattern is accounted for in a stipulation in the entry of the first person pronoun that it not occurs in action reports.

This need for stipulation arises because context shifting operators, least as far as they have been proposed in the existing literature, apply to a full (syntactic) clause and (semantic) proposition. By contrast, in the demonstration account there is no embedded subject to stipulate anything about because there is no embedded clause: the demonstrated predicate is the main predicate of the clause. It would seem to be a clear advantage of the analysis proposed here that an additional overt subject under role shift would be entirely unexpected. Furthermore, stipulations would be expected to be language specific: American, French, and Danish sign language may not allow an overt first person pronoun, while another sign language might be expected to allow it (and even potentially, to disallow first person agreement marking). It remains to be seen what occurs beyond these languages, but it would be an advantage to the current proposal’s application to other sign languages if no sign language allows an overt pronoun of any kind in action reports, i.e. if all behave like ASL, and at least that does seem to be the case in French SL (LSF) as well.

In addition to the first person pronoun, Schlenker reports that other indexicals frequently also trigger an attitude interpretation, although judgments were somewhat mixed, with *HERE* more likely to permit an action interpretation than *TOMORROW*, which is more likely to force an attitude interpretation. Although neither the context shifting analysis nor the demonstration analysis make perfect predictions here, it would seem that the flexibility of the demonstration analysis is an advantage for this mixed data, which may depend on how performative versus descriptive the signer wants to be in their use of the indexicals.

A second, related, problem with a context shifting account is that any other non-iconic lexical material *also* forces a reinterpretation as a reported attitude interpretation in LSF. Its status in ASL is more equivocal, so Schlenker (2014) proposes a separate stipulation for both ASL and LSF that under role shift, anything that can be interpreted as iconic must be (“maximize iconicity”), while in LSF, *everything* under role shift should be iconic (“exhaustify iconicity”). In the demonstration analysis, it is of course *part of being a demonstration* to interpret pieces of the language as iconic. Given the variation between LSF and ASL, which is really variation between two or three signers, the demonstration analysis seems to have the advantage in being more flexible than broad stipulations that everything be interpreted iconically. One might want to say that “demonstration” is essentially a renaming of the stipulations to interpret language iconically,

but the accounts do make different predictions. For example, the “maximize iconicity” stipulation fails at predicting the impossibility of a first person pronoun, which is highly iconic, as it is a point to the speaker’s own chest. The demonstration analysis would predict that any subject, even one as iconic as a first person pronoun, is not permitted in the action report because only a *predicate* is demonstrated, given that the demonstration is an event modifier.

Further bolstering the demonstrational view versus clausal context shift, it seems to be the case that most examples of reported action in previous sign language literature involve just the predicate under role shift, not any further arguments and certainly not an entire clause (Engberg-Pedersen 1995; Lillo-Martin 1995, 2012).¹³ Moreover, in discussions of word order in ASL there are reports of exceptions to the typical SVO word order to allow SOV, and these are restricted to cases of action reports (73). This seems to suggest that a preference for only including the predicate in the demonstration may be strong enough to permit noncanonical word order (Liddell 1980), even when there are not information structural reasons for doing so beyond the demonstration.

- (73) a. MAN FORGET NUMBER SVO (noniconic, typical word order)
 b. *MAN NUMBER FORGET *SOV (noniconic, atypical word order)
 ‘The man forgot the number’
- c. MAN SEE MOVIE SVO (noniconic, typical word order)
 d. MAN MOVIE LOOK-UP-AT SOV (iconic predicate, atypical word order)
 ‘The man watched the movie’
- e. WOMAN PIE PUT-IN-OVEN SOV (iconic predicate, atypical word order)
 ‘The woman put the pie into the oven.’
- (Liddell 1980)

Viewed in this way, the fact that the presence of non-iconic lexical material also triggers a reported attitude interpretation in LSF is not at all surprising. In the current demonstration account, Schlenker’s “exhaustify iconicity” requirement that everything in action role shift be iconic is subsumed under the idea that the entire predicate is a demonstration.

Data from ASL, which is more equivocal on this point, instead becomes the puzzle: in a few cases, Schlenker (2014b) reports that non-iconic lexical material seems to be grammatical under action role shift. Older literature on this same phenomenon suggests that this is not grammatical (as above), and all previously reported naturally occurring examples lack any non-iconic lexical material. Why they should be judged grammatical by some signers in controlled contexts, and why in ASL but not LSF, remains a puzzle, but one possibility is that they are grammatical but highly dispreferred, given the flexibility noted above in word order in ASL and the pressures to move non-iconic arguments of iconic verbs to the left in examples like (73). An alternative possibility is that something along the lines of mixed quotation is at work here, as in (14)-(17) above. To complicate an already complicated picture of mixed quotation, even in English some quotations do not allow shifted indexicals if they are simply demonstrating someone else’s use of

¹³ Thanks to Diane Lillo-Martin for this observation.

the word, but rather they require an embedding verb (e.g. say, be like) that requires the words to be part of the demonstration to shift the indexical interpretation (74).

(74). (*Context: Suppose Hugh thinks that Mary is the speaker's sister.*)

#I_i guess I_i should be happy because 'his_i sister' is coming. (sister = Mary)

I_i guess I_i should be happy because my_i 'sister' is coming. (sister = Mary)

Hugh said "his_i sister is coming." (sister = Mary)

(modified from von Stechow 2010)

Something similar to this may also be at work in the difference between the doxastic and proffering predicates discussed by Koulidobrova and Davidson (above), where the doxastics may have an interpretation in which they are demonstrational predicates (i.e. one demonstrates a dream or an imagining), while proffering verbs (*say, claim*) must be language reports, a different beast altogether in which the verb cannot participate in the demonstration but rather introduces it.

The role of iconicity is clearly important both to Schlenker's context shift account and the demonstration account proposed here, and shares with many previous accounts in the literature including cognitive accounts (Liddell 2003) the intuition that iconicity in classifiers, language reports, and action reports are related. All suffer from an under-defined notion of iconicity. Structurally, however, the proposals are quite different: the role shifted predicate is the matrix/main verb in the demonstrational account, but is inside an embedded clause under context shift. This makes quite different predictions with regards to various root clause phenomena that should be checked further, but at least with regards to allowed "subjects" seems to pattern in favor of demonstrational predicate modification. Moreover, they differ on what makes indexicals "shift": in the context shifting operator view there are seemingly no restrictions on how which indexicals can shift to what, beyond that (for some accounts) they shift together. In the demonstration account, on the other hand, indexical expressions should only shift toward more iconic directions. The heart of the matter is to determine what, precisely, is iconic and permits a demonstration, which leads back to the suggestion for more research on what precisely demonstrates iconicity for "demonstration" purposes. My suspicion is that it will be a gradient phenomenon that varies from speaker to speaker, based on the ASL versus LSF variation. What should not vary is the clausal structure, which is consistent with the robust inability of embedded subjects in reported actions.

5. Conclusions on quotation, demonstration, and the pragmatics of iconicity

This paper has proposed a renewed focus on natural (spoken or signed) language data in the study of language reports/quotation, with written language as a special, peripheral case. The focus on language reports in natural language made clear the advantages of an analysis of language reports as a demonstration (Clark and Gerrig 1990). Section 2 proposed that quotative verbs introduce demonstrations of speech (e.g. *say, be like*), and additionally, that *like* can combine with a wide variety of verbs to introduce demonstrations in many more types of non-quotative predicates. Importantly, the notion of demonstration in this view is context dependent, i.e., demonstrations are performed so as to convey whatever aspects of an event are relevant within a given context of speech. Considered in this way, written quotation is merely a special

case of verbatim demonstration in which the context of writing makes the words that are used both the necessary and sufficient demonstratum.

Sign languages provided support for the demonstration view of quotation in three major ways. First, they do not have a commonly accepted writing system, so they remove the tendency to rely on writing to study natural language quotation. Second, iconicity has been an area of focus in the sign language linguistics literature. One example of iconicity formally implemented comes from the analysis of classifiers presented in Section 3 which showed striking similarities to the demonstration analysis of spoken language quotation given in Section 2. The demonstration view of classifier constructions in ASL has the advantages of being both flexible enough to incorporate the gestural aspects of these “depicting verbs” but still formalizable in the framework of event semantics, and finds support in argument alternations and ASL-English bilingual language.

Section 2 and Section 3 focused on two instantiations of demonstration as event modification: quotation (in spoken and sign languages) and classifier predicates (in sign languages). Given that both use the same underlying semantics, it would not be surprising to find that a language would use the same linguistic strategy to instantiate both. I argue that this is precisely the case in *role shift*, which can be used not only for language reports (quotation), but also action reports. In Section 4 I returned to an old intuition by Supalla (1982) that action reports are “body classifier” predicates, which has two major advantages in (i) accounting for the distribution of indexical expressions in action role shift and (ii) predicting their forced iconic interpretation. This analysis removes the special need for monstrous indexicals in action reports, since in this analysis first person marking is a demonstration, not an agreement indexical. By marking the scope of language reports overtly and unifying quotation and classifiers in a single construction, role shift provides a third reason to consider demonstrations from the perspective of sign languages.

Whether the analysis proposed here can be extended to other sign languages remains to be seen, but it seems likely that it could. Zwitserlood (2012) reports that classifiers behaving roughly similarly to the ones in ASL exist in all sign languages to date, with one notable exception, Adamorobe Sign Language (Nyst 2007). As for role shift, Lillo-Martin (2012) notes that other sign languages that have been reported to follow the same basic pattern of language reports and action reports under role shift (at least, the fact that role shift is used and that sometimes the indexicals shift) include British, Catalan, German, Nicaragua, Quebec, and Swedish sign languages.

This paper follows the spirit of Schlenker et al. (2013) in attempting to bring together both the formal and iconic properties of sign languages into a single proposal by allowing the flexibility of an iconic component within a formalized system. However, the current analysis purposely offloads much of the iconic work (done through the notion of a demonstration) to the pragmatic component of the grammar, and by doing so is able to provide a unified treatment of both sign language and spoken language data. More generally, the sign language literature has often been described as having a tension between those who formalize the structure within existing frameworks in order to treat it like spoken language versus those who resist a formalization in order to account for sign language specific phenomena. A similar situation may be found in mathematical proofs, where symbolic (arbitrary) formalism has been opposed to (iconic) diagrams in an effort to increase accuracy, while having the effect of sacrificing often crucial

information and insight provided through the iconicity (Shin 2012). Through the notion of a demonstration, entered into the structure as an arbitrary symbol but calculated through a pragmatic component based on iconicity, we are able here to model both arbitrary and non-arbitrary aspects of a number of puzzles including role shift, quotation, and classifier predicates.

I will conclude with some thoughts about this pragmatic component of iconicity, and in particular the “faithfulness” or verbatim representation of quotations and demonstrations. As already discussed, quotation in written language is frequently interpreted as verbatim repetition from the original context (modulo disfluencies and language of utterance) but this requirement is much looser in spoken language. Taking them one step further to the demonstrative *be like* examples, we even saw whines, the behavior of cats, etc., being able to be “quoted” in written language and reported in spoken language speech reports. When it comes to classifier predicates, this flexibility allowed by the demonstrational verbs is also welcome, since different aspects of the meaning of space for classifiers are relevant at different times. It seems that pragmatics does quite a lot of work.

Consider two examples of the pragmatic component using classifier predicates. In the first example, (75a), the hands representing the man and the women, respectively, move together in a straight line toward the center of the signing space. Here, the signer may want to convey that the people did, in fact, walk in a straight line toward each other, but it could also be the case that they wobbled quite a bit, or even that the path was not straight at all but that the path is not relevant to the story. In contrast, the same example but with a different movement (75b) show the figures making a notable zig-zagging path through space. The interpretation here is that the man and women did, in fact, not move toward each other in a straight path; if they did, denial of the statement is licensed.

- (75)
- a. Right hand: CL-1(straight movement toward center)
Left hand: CL-1(straight movement toward center)
‘The two people walk toward each other.’
#Response: ‘No, they went toward each other in a zigzag pattern!’
#Response: ‘And actually, they went toward each other in a zigzag pattern!’
 - b. Right hand: CL-1(zigzag movement toward center)
Left hand: CL-1(zigzag movement toward center)
‘The individuals walked toward each other in a zigzag pattern.’
Response: ‘No, they went straight for each other!’

A fruitful way to view what is happening in (75a) versus (75b) is through Grice’s Maxim of Quantity: say just as much as it needed, and no more. In (75a) the straight path is the easiest way to articulate bringing the two hands together and so may not be intending to demonstrate any manner of path at all. In contrast, the zigzagging path (75b) takes more effort to produce and so should be interpreted as part of the demonstration. Precisely the same thing happens in demonstrations involving direct language reports, in which the default language (the one that takes no additional effort to produce given our conversational context) conveys nothing about what language John used but the marked/effortful language does ((25)-(26) repeated below in (76)).

- (76) a. We met under the Eiffel tower, and the first thing he said was “my name is John.”
 b. We met under the Eiffel tower, and the first thing he said was “je m’appelle John.”

Interestingly, when it comes to (76), a rule similar to “maximize iconicity” would be too permissive: it cannot account for cases in which there is no difference in iconicity (e.g. the languages in (76), etc.) and yet one is interpreted as part of the demonstration (i.e. as iconic) and the other is not.

There is also at least one other way in which a rule like “maximize iconicity” is not (at least, as currently stated or easily extended) permissive enough: when the iconic aspect doesn’t map onto a single dimension. Consider Schlenker’s (2014b) example (77a), in which the distance of the two hands in the production of the verb GROW correlates positively with the size of the growing event, and the speed of the movement of the hands correlates positively with the speed of the group’s growth. The iconicity function can encode these correlates into the semantics, but faces more complexities in the slight modification in (77b).

- (77) a. MY GROUP GROW
 [where GROW is signed with endpoints not far from each other, in a fast motion]
 ‘My group grew to a small size at a fast pace’
 b. MY GROUP GROW
 [where GROW is signed in a “jerky” motion alternating slow and fast growth with wide endpoints]
 ‘My group grew at an uneven rate but eventually became large’

In (77b), were “maximize iconicity” to apply, it would only make claims about the speed and final location of the handshapes and cannot account for the starts and stops. Of course, the function could be amended to include acceleration as well, but at some point it seems like such a specific function falls short. A seeming advantage of the specific iconic functions over the demonstration view is that they make a more specific prediction than the demonstration theory, but it is precisely because they encode the individual dimensions of iconicity that they may continue to fall short, because any new dimension that is added to a demonstration will be able to be interpreted, as long as it is understood to be relevant to the demonstration to each of the interlocutors.

That the extent of the demonstration should be pragmatically determined seems to be supported by experimental data. When signers were asked to describe videos designed to elicit action reports/body shift, they often alternated between action reports and classifier predicates to describe the same event (Quinto-Pozos 2007), an alternation also reported in Aarons and Morgan (2003). As Quinto-Pozos notes, it seems to be (close to) obligatory to have some type of demonstration when describing the videos that signers saw, but the details of how the demonstration is introduced (via classifier or body shift, or how detailed) is left to the syntax/semantic structure, and the interpretation of the demonstration to the pragmatics.

An important idea in this paper is that at least two domains –classifiers and role shift– which look superficially quite different from English (at least, very different from written English, but arguable not from spoken English) are analyzed in a framework of demonstration that is already

present in the spoken language in expressions of quotatives and event modifiers like “like”. I think it is important to mention that not *every* use of space in sign languages involves a demonstration: examples of spatial language in sign languages that are not demonstrations are (i) the assignment of discourse referents to arbitrary loci in the signing space for anaphoric reference and the associated verb agreement; (ii) the iconic -but not demonstrational - relationship of plural loci with each other (Schlenker et al. 2013); and (iii) the use of height in noun phrases to indicate specificity (Barberà 2012) or quantifier domain size (Davidson and Gagne 2014), among others. While three-dimensional space *can* be used for demonstrations, this is only one of many uses within the sign linguistic system.

Unlike analyses involving context shift operators for role shift as separate from quotation, the analysis proposed in this paper analysis does not require any modality-specific machinery and in fact suggests that there are no differences between sign and spoken languages in the way that they blend demonstrations with other parts of meaning. What sign languages can do, so can spoken languages; one reason this has seemed to not be the case in previous literature is the focus on written instead of natural spoken or signed languages. In this paper, sign languages provided a window into how direct language reports work *without* the confines of a writing system, but *with* a clear marking of the scope of quotation. In doing so, we turn the tables on the relationship between sign languages and spoken languages by using tools and understanding developed for sign language iconicity to better understand spoken, and more generally all natural, language.

References

- Aarons, D. & Morgan, R. (2003). Classifier predicates and the creation of multiple perspectives in South African Sign Language. *Sign Language Studies*, 3 (2), 125–156
- Abbott, B. (2005). Some notes on quotation. *Belgian Journal of Linguistics* 17, 13–26.
- Aikhenvald, A. (2000). *Classifiers: A Typology of Noun Categorization Devices*. Oxford: Oxford University Press.
- Allan, K. (1977). Classifiers. *Language* 53, 285–311.
- Anand, P. & Nevins, A (2004). Shifty Indexicals in Changing Contexts. *Proceedings of SALT 14*, CLC Publication.
- Barberà, G. (2012). *The meaning of space in Catalan Sign Language (LSC). Reference, specificity and structure in signed discourse*. PhD Thesis, Universitat Pompeu Fabra.
- Benedicto, E., & Brentari, D. (2004). Where did all the arguments go?: Argument-Changing Properties of Classifiers in ASL. *Natural Language and Linguistic Theory*, 22, 743–810.
- Clark, H. H., & Gerrig, R. J. (1990). Quotations as Demonstrations. *Language*, 66(4), 764–805.
- Cogill-Koez, D. (2000). A model of signed language 'classifier predicates' as templated visual representation. *Sign Language & Linguistics* 3:2, 209–236.
- Cormier, K., Smith, S. & Sevcikova, Z. (In press.) Rethinking constructed action. *Sign Language and Linguistics*.
- Davidson, D. (1967). The logical form of action sentences. In Nicholas Rescher (ed.), *The Logic of Decision and Action*, pp. 81–95. University of Pittsburgh Press, Pittsburgh. Republished in Donald Davidson, 1980. *Essays on Actions and Events*. Oxford University Press, Oxford.
- Davidson, D. (1979) Quotation. *Theory and Decision* 11, 27–40.
- Davidson, K. & Gagne, D. (2014) Vertical Representations of Quantifier Domains. in *Proceedings of Sinn und Bedeutung* 18.
- DeMatteo, A. (1977). Visual Imagery and Visual Analogues in American Sign Language. In L. Friedman (Ed.), *On The Other Hand*, Academic Press, New York.
- Dik, S. C. (1975). The semantic representation of manner adverbials. In *Linguistics in the Netherlands 1972-1973*, ed. A. Kraak, 96-121. Assen: Van Gorcum.
- Dudis, P. (2004). Body Partitioning and Real-Space Blends. *Cognitive Linguistics* 15(2), 223–238
- Emmorey, K., Borinstein, H.B., Thompson, R., & Gollan, T.H. (2008). Bimodal bilingualism. *Bilingualism: Language and Cognition*, 11(1), 43–61.
- Emmorey, K. & Herzig, M. (2003). Categorical versus gradient properties of classifier constructions in ASL. In K. Emmorey (Ed.), *Perspectives on Classifier Constructions in Signed Languages*, Lawrence Erlbaum Associates, Mahwah NJ, pp. 222–246
- Emmorey, K. & Reilly, J. (1998). The development of quotation and reported action: conveying perspective in ASL. In E.V. Clark (Ed.), *The Proceedings of the Twenty-Ninth Annual Child Research Forum*, Center for the Study of Language and Information, Stanford, CA, pp. 81–90

- Engberg-Pedersen, E. (1995). Point of View Expressed through Shifters. Emmorey, K. and J. Reilly, *Language, Gesture, and Space*. Hillsdale, NJ. Lawrence Erlbaum Associates, 133-154.
- von Fintel, K. (2010). How Multi-Dimensional is Quotation? note at the Harvard-MIT-UConn Workshop on Indexicals, Speech Acts, and Logophors. retrieved at <http://philpapers.org/rec/VONHMI>
- Gehrke, Berit & Castroviejo, E. (2015). Manner and degree: An introduction. *Natural Language and Linguistic Theory*, published online 29 April 2015.
- Haynie, H. Bower, C., & LaPalombara, H. (2014) Sound symbolism in the Languages of Australia. *PLOS ONE* DOI: 10.1371/journal.pone.0092852
- Heim, I. (1982). *The Semantics of Definite and Indefinite Noun Phrases*. Doctoral dissertation, University of Massachusetts Amherst
- Hermann, A. & Steinbach, M. (2012). Quotation in Sign Languages – A Visible Context Shift. In I. van Alphen and I. Buchstaller, *Quotatives: Cross-linguistic and Cross Disciplinary Perspectives*. Amsterdam: Benjamins, 203-228.
- Hinton, L., Nichols, J., & Ohala, J. (Eds.) (2006) *Sound Symbolism*. Cambridge: Cambridge University Press.
- Huebl, A. (2012). Role shift, indexicals, and beyond- new evidence from German Sign Language. Presentation at *Texas Linguistic Society* conference, Austin, TX, 23-24 June.
- Janis, W. (1992). *Morphosyntax of the ASL Verb Phrase*. PhD Thesis, State University of New York, Buffalo.
- Kaplan, D. (1989). Demonstratives: An Essay on the Semantics, Logic, Metaphysics, and Epistemology of Demonstratives and Other Indexicals. In Themes from Kaplan, ed. Joseph Almog, John Perry, and Howard Wettstein. 481–614. New York: Oxford University Press.
- Kasimir, E. (2008). Prosodic correlates of subclausal quotation marks, *ZAS Papers in Linguistics* 49, 67-77.
- Kegl, J. (1990). Predicate Argument Structure and Verb-Class Organization in the ASL Lexicon. In C. Lucas, *Sign Language Research: Theoretical Issues*, Gallaudet University Press, Washington DC, 149-175.
- Klima, E. & U. Bellugi (1979). *The Signs of Language*. Harvard University Press.
- Koulidobrova, E. & Davidson, K. (2014) Check your attitude: Role-shift and embedding in ASL. Presentation at Formal and Experimental Advances in Sign language Theory, Venice, June 2014.
- Kratzer, A. (1998). More structural analogies between pronouns and tenses. In *Proceedings of Semantics and Linguistic Theory (SALT) 8*. CLC Publications.
- Landman, M. & Morzycki, M. (2003). Event-Kinds and the Representation of Manner. In N. Mae Antrim, G. Goodall, M. Schulte-Nafeh, and V. Samiian, *Proceedings of the Western Conference in Linguistics (WECOL) 11*.
- Liddell, S. (1980). *American Sign Language Syntax*. Mouton, the Netherlands.

- Liddell, S. (2003). *Grammar, Gesture, and Meaning in American Sign Language*. Cambridge: Cambridge University Press.
- Lillo-Martin D. (1995). The Point of View Predicate in American Sign Language. Emmorey, K. and J. Reilly, *Language, Gesture, and Space*. Hillsdale, NJ. Lawrence Erlbaum Associates, 155-170.
- Lillo-Martin, D. (2012). Utterance reports and constructed action. In R. Pfau, M. Steinbach, and B. Woll, *Sign Language: An International Handbook*. De Gruyter Mouton, 365-387.
- Lillo-Martin, D., & Chen Pichler, D. (2008). Development of sign language acquisition corpora. Proceedings of the 3rd Workshop on the Representation and Processing of Sign Languages, 6th Language Resources and Evaluation Conference, 129-133.
- Macken, E. Perry, J. & Haas, C. (1993). Richly Grounded Symbols in American Sign Language. *Sign Language Studies* 81, 375-394.
- Maier, E. (2014). Mixed quotation: The grammar of apparently transparent opacity. *Semantics & Pragmatics* 7(7). 1-67.
- Meir, I, Padden, C., Aronoff, M. & Sandler, W. (2007). Body as subject. *Journal of Linguistics* 43(3), 531 - 563.
- Morzycki, M. (2014). *Modification*. Manuscript for *Key Topics in Semantics and Pragmatics* series, Cambridge University Press.
- Nyst, V. (2007). *A Descriptive Analysis of Adamorobe Sign Language (Ghana)*. PhD Dissertation, University of Amsterdam, Utrecht: LOT.
- Petroj, V., Guerrero, K., & Davidson, K. (2014) ASL dominant code-blending in the whispering of bimodal bilingual children. In *Proceedings of the Boston University Conference on Language Development* 39.
- Piñón, C. (2007). Manner adverbs and manners, paper presented at the conference 'Ereignisse semantik', Tübingen, December 2007.
- Potts, C. (2004). The dimensions of quotation. In Barker and Jacobson (Ed.s) *Direct Compositionality*. 405-431.
- Quer, J. 2005 Context Shift and Indexical Variables in Sign Languages. In Georgala, E. and Howell, J (eds). *Proceedings from Semantics and Linguistic Theory 15*, Ithaca NY 152-168.
- Quine, W. (1940). *Mathematical logic*. Vol. 4. Cambridge: Harvard University Press
- Quinto-Pozos, D. (2007). Can constructed action be considered obligatory? *Lingua* 117, 1285-1314
- Schick, B. (1987). *The Acquisition of Classifier Predicates in American Sign Language*. PhD Thesis, Purdue University.
- Schick, B. (1990). Classifier Predicates in American Sign Language. *International Journal of Sign Linguistics* 1, 15-40.
- Schlenker, P. (2003). A Plea for Monsters. *Linguistics and Philosophy* 26(1), 29-120
- Schlenker, P., Lamberton, J. and M. Santoro (2013). Iconic Variables. *Linguistics and Philosophy* 36(2), 91-149

- Schlenker, P. (2014a). Super Monsters I : Attitude and Action Role Shift in Sign Language, manuscript.
- Schlenker, P. (2014b). Super Monsters II : Role Shift, Iconicity and Quotation in Sign Language, manuscript.
- Shan, C. (2010). The character of quotation. *Linguistics and Philosophy* 33(5), 417-443.
- Shin, S.-J. (2012). The forgotten individual: diagrammatic reasoning in mathematics. *Synthese*. 186:149-168.
- Stokoe, W. (1960). Sign Language Structure: An Outline of the Visual Communication Systems of the American Deaf. *Studies in linguistics: Occasional papers (No. 8)*. Buffalo: Dept. of Anthropology and Linguistics.
- Supalla, T. (1982). *Structure and Acquisition of Verbs of Motion and Location in American Sign Language*. PhD Thesis, University of California, Ssan Diego.
- Tannen, D. (1989). *Talking Voices: Repetition, Dialogue, and Imagery in Conversational Discourse*. Cambridge: Cambridge University Press.
- Zucchi, S. (2004). Monsters in the Visual Mode. Manuscript, Università degli Studi di Milano.
- Zucchi, S., Cecchetto C., & Geraci, C. (2012). Event Descriptions and Classifier Predicates in Sign Languages. Presentation FEAST in Venice, June 21, 2011.
- Zwitserslood, I. (1996). Who'll Handle the Object. MA Thesis, Utrecht University.
- Zwitserslood, I. (2012). Classifiers. In R. Pfau, M. Steinbach, and B. Woll, *Sign Language: An International Handbook*. De Gruyter Mouton, 158-185.