

Singular count NPs in measure constructions*

Hana Filip
Heinrich Heine University

Peter R. Sutton
Heinrich Heine University

Abstract Count Ns like *fence*, *wall* or *twig* notoriously pose problems for the semantic analysis of the mass/count distinction, given that they exhibit grammatical count behavior, thus patterning with bona fide count Ns like *cat*, but unlike the latter, fail to denote quantized predicates (in Krifka’s (1989) original sense); at the same time, they do not denote cumulative predicates, unlike mass Ns, such as *mud* or *water*. This puzzling class of count Ns has another intriguing property, so far largely neglected in contemporary mass/count debates in formal semantics: most of its members felicitously occur in pseudo-partitive (measure) NP constructions. Take, for instance, *wall*, as in *Thick woolen drapes of red and gold covered every inch of wall* (COCA). We argue that count Ns like *fence*, *wall* or *twig* fail to denote quantized predicates, because they admit of multiple overlapping individuation schemas with respect to what counts as ‘one’. In a nutshell, (i) *fence*, like *cat*, but unlike *mud* is quantized at specific counting contexts (and so grammatically count), (ii) *fence*, like *mud*, but unlike *cat* is non-quantized at the null counting context (the union of interpretations across all specific counting contexts), which make them felicitous in pseudo-partitive (measure) NPs.

Keywords: count, mass, mereology, measure (pseudo-partitive) NP, context

1 Main data and question

A large class of count Ns like *fence*, *sequence*, *wall*, *branch*, *arc*, *hedge*, *twig*, *line*, *band*, *table* are notoriously problematic for any semantic analysis of the mass/count distinction. They are count, because they satisfy a hallmark grammatical property of count Ns, such as straightforward compatibility with numerical modifiers, which mass Ns lack:

- (1) three cats_C/fences_C
- (2) three #snows_M/#furnitures_M

* We would like to thank the participants at SALT 27, many of who gave us some helpful comments that contributed a lot to this paper. Thanks also to Scott Grimm, Fred Landman, Beth Levin, and Susan Rothstein and other attendees of the workshop on countability in Düsseldorf in 2016. Also to participants of our 2016 ESSLI class. This research was funded by the German Research Foundation (DFG) CRC991, project C09.

Notionally, there is a striking split among count Ns with respect to how their CRITERION OF INDIVIDUATION, and what is ‘one’ countable entity in their denotation, is specified. Only ‘sortal’ count Ns like *cat* (see Pelletier 1975 for the use of ‘sortal’ in this context) lexically determine their unique criterion of individuation in a way that is context-independent. However, count Ns like *fence* do not. What we take to be ‘one’ in their denotation is not constant through space and time, but is a fundamentally context-sensitive notion.

This notional split among grammatical count Ns was first debated in connection with Krifka’s (1986; 1989) mereological theory of count Ns. One of its basic tenets is that singular count Ns uniformly denote QUANTIZED predicates. For instance, (a) *cat* is quantized, because if it holds of some individual, it cannot also hold of some of its proper parts. However, as Krifka (1989: 87n, due to Partee, p.c.) also observes, there are singular count Ns like *twig* or *sequence* which fail to denote quantized predicates. For instance, if you have something that falls under the description of (a/one) *twig* and break it into two, each part may still be describable by (a/one) *twig*. Depending on our perspective, a fence may be a part of another fence, a sequence (of numbers) may be part of another sequence, a bouquet may have a part that is another bouquet. But this means that quantization *per se* is not a necessary semantic condition for Ns to be grammatically count. This raises the following question:

QUESTION 1: If grammatical counting depends on some notion of individuation, how can it be best characterized, given that it cannot be some context-independent individuation concept inherent in the lexical structure of all count Ns?

In order to address the context-sensitivity of count Ns like *fence*, Rothstein (2010, 2017) proposes that all count Ns have lexical meanings partially specified by context which identifies contextually disjoint sets in their denotation (her ‘semantic atoms’). Acknowledging that this might motivate why count Ns like *cat* and also count Ns like *fence* exhibit the same behavior in count syntax, what still remains puzzling is the observation that only singular count Ns like *fence*, but not those like *cat*, are perfectly natural and felicitous in the measure (pseudo-partitive) NP:

- (3) a. #three pounds of *cat*_C
 b. The crowd lines 100 yards of *fence*_C, three deep. (COCA)

Most agree that the measure (pseudo-partitive) NP disallows singular count Ns, and sanctions mass terms and bare plurals (Krifka 1989; Filip 1992, 2005; Schwarzschild 2002, 2006; Nakanishi 2007; Landman 2016 i.a.):

- (4) a. three inches of *snow*_M
 b. three crates of *furniture*_M
 c. three baskets of *kittens*_C

Although not all singular count Ns that fail to denote quantized predicates are felicitous in the pseudo-partitive (measure) NP, e.g. *#three rows of sequence*, it is highly significant and puzzling that there are *any* singular count Ns that are felicitous in this construction, e.g., *fence*, *hedge*, *wall*, *twig*, *branch* etc., given that, at first blush, this could be taken to mean that *fence*-like Ns are not count Ns at all, but rather dual-life Ns, on a par with *cake* or *stone*. If *fence*-like Ns were dual-life, they ought to be felicitous in all count and mass/plural syntactic environments. However, this prediction is not borne out, because Ns like *fence* and *hedge* cannot occur as bare singulars in all argument positions, unlike mass Ns (5):

- (5) Water/#fence/#hedge/#cat lay between us and the giant alligators.

What the above observations suggest that is that any adequate analysis of count Ns must account not only for the context-sensitivity of grammatically count Ns like *fence*, but also for their seemingly puzzling compatibility with the pseudo-partitive measure NP. Put differently, it must answer the following question:

QUESTION 2: If both singular count Ns like *cat* and *fence* behave grammatically alike in count syntax, as evidenced by (1), and if both are ungrammatical in mass syntax, as evidenced by (5), why are ONLY singular count Ns like *fence* also perfectly natural and felicitous in the measure (pseudo-partitive) NP?

2 Background

2.1 Krifka (1989): From cumulativity and quantization to extensive measure functions

Krifka (1989, 1995) proposes that basic lexical mass and count Ns are typically distinguished, but all take their denotation from a SINGLE domain algebraically structured by means of a complete join semilattice, undetermined with respect to atomicity (in departure to the double-domain ontology, atomic and non-atomic, originally proposed by Link (1983)). A defining property of mass Ns is CUMULATIVE reference (Quine 1960). In the simplest terms, if there are two entities to which *P* applies, *P* also applies to their mereological sum:

- (6) CUMULATIVE *P* : $\forall P[\text{CUM}(P) \leftrightarrow \forall x \forall y [P(x) \wedge P(y) \rightarrow P(x \sqcup y)]]$

Mass Ns are analyzed as one-place predicates that lexically specify only a qualitative criterion of application: e.g. $\llbracket \text{water} \rrbracket = \lambda x [\text{WATER}(x)]$. Basic lexical count Ns (*boy*, *apple*) have QUANTIZED reference (a predicate *P* is quantized if and only if whenever it holds of something, it does not hold of any of its proper parts):

- (7) QUANTIZED *P*: $\forall P[\text{QUA}(P) \leftrightarrow \forall x \forall y [P(x) \wedge P(y) \rightarrow \neg(x \sqsubset y)]]$

Count Ns lexically specify not only a qualitative but also a quantitative criterion for their application, which is represented by means of the **NU** function, standing for a ‘natural unit’, e.g., the organism for living beings (Krifka 1989: 84), and ensures quantization of basic lexical count Ns. Count N denotations are analyzed as 2-place relations between entities and numbers: e.g. $\llbracket \text{apple} \rrbracket = \lambda n \lambda x [\text{APPLE}(x) \wedge \text{NU}(\text{APPLE})(x) = n]$. Formally, **NU** is an extensive measure function, and just as other extensive measure functions μ , including standard measures like **OUNCE**, **HOURL**, **LITER**, it is additive (see Krantz, Luce, Suppes & Tversky 1971). An additive measure function tracks the part structure of entities it measures. For instance, adding 2 lbs of apples to 3 lbs of apples yields 5 lbs of apples, which weighs more in pounds than the quantity of any of its proper parts.¹ Krifka (1989) analyzes extensive measure functions as ‘quantizing modifiers’ that derive quantized predicates (e.g., *500 kgs of wool/potatoes*) from non-quantized ones (e.g., *wool, potatoes*):

- (8) QUANTIZING MODIFICATION:
 $\forall P \forall Q [\text{QMOD}(P, Q) \leftrightarrow \neg \text{QUA}(P) \wedge \text{QUA}(Q(P))]$ (Krifka 1989: 82)

As observed at the outset, the quantization property (7) is not a necessary condition for Ns to be grammatically count, as Krifka (1989: 87n, due to Partee, p.c.) also observes. He mentions two examples: *twig* and *sequence*. For instance, the sequence of numbers 1, 2, 3, 4, 5 is a sequence that has a proper subpart, 2, 3, 4 which is also a sequence of numbers, and both fall under *sequence* (Zucchi & White 1996, 2001). But this means that the denotation of *sequence* does not consist of stable ‘natural units’ but can be rather arbitrary.

Krifka’s quantizing modification yields quantized predicates. One of the reasons for this is the ungrammaticality of phrases like **hundred grams of five hundred meters of wool* (Krifka 1998: 202, (12c)), which is predicted on the assumption that the pseudo-partitive (measure) NP *five hundred meters of wool* is quantized. If so, then this straightforwardly precludes the application of *hundred grams* to it, which requires only non-quantized predicates. Nonetheless, the requirement that the quantizing modification yield only unambiguously quantized predicates might be too strong. Pseudo-partitive (measure) NPs like *twenty kilos of potatoes/flour* seem to also admit of what looks like a ‘massy’, $\neg \text{QUA}(P)$, interpretation, under certain narrowly defined conditions (for similar examples see also Rothstein 2011 and Landman (2016)):

- (9) 60 pounds of potatoes is too much even for Sam to lift.

In (9), the predicate has a singular agreement on the verb and it contains *much*, which only selects for mass predicates. This would seem to indicate that the subject measure NP *60 pounds of potatoes* does not have a quantized plural count NP

¹ For a formal definition of an extensive measure function see Champollion & Krifka 2016: §13.21.

interpretation *tout simple*. At the same time, the bare plural term (here *potatoes*) used in the measure (pseudo-partitive) NP to denote what is measured should be best viewed as retaining its individuated, plural count, structure, as [Krifka \(1989\)](#) and [Landman \(2016\)](#), among others, agree (*pace* [Rothstein 2011](#)). If measure NPs like *60 pounds of potatoes* do not always have a quantized interpretation, as (9), for instance, shows, then what still remains puzzling, however, is the ungrammaticality of pseudo-partitive (measure) NPs like **hundred grams of five hundred meters of wool* ([Krifka 1998](#): 202, (12c)). If we were to assume that a pseudo-partitive (measure) NP like *five hundred meters of wool* has either a quantized or a non-quantized interpretation available, depending on context, as [Rothstein \(2011\)](#) and [Landman \(2016\)](#) argue, why is it incompatible with an extensive measure phrase like *hundred grams (of)*, which selects non-quantized predicates?

2.2 Rothstein (2010): Disjointness at a counting context

[Rothstein \(2010\)](#) develops a theory of the mass/count distinction in which context-sensitive count Ns like *fence* or *wall* take the center stage. Their existence renders implausible any uniform analysis of all count Ns in terms of some context-independent individuation concept inherent in their lexical structure, be it an NU function in [Krifka 1989, 1995](#), or ‘atom’ in a join semilattice along the lines of [Chierchia 1998](#). The overall perspective [Rothstein \(2010\)](#) takes is best illustrated with her leading example, fencing around a square field. It may be viewed as one or four non-overlapping fences in the same situation under different criteria of individuation. Consequently, the question *How many fences are there?* has no determinate answer. A perfectly felicitous answer is *Four fences*, but also *One fence*, depending on the individuation schema we choose to apply. According to [Rothstein \(2010\)](#), counting as a grammatical operation depends on atomicity relative to a *counting context*, which is a subset of the domain of entities which count as atomic units for the purpose of counting. Put in more technical terms, grammatical counting relies on counting $\langle \text{entity}, \text{context} \rangle$ pairs, i.e., a denoted entity and the context in which it counts as one. For our example above, and assuming the four sides are *a*, *b*, *c*, and *d*, we get:

$$(10) \quad \text{In context } k_1: |\{\langle a, k_1 \rangle, \langle b, k_1 \rangle, \langle c, k_1 \rangle, \langle d, k_1 \rangle\}| = 4 \quad (\text{Four fences})$$

$$(11) \quad \text{In context } k_2: |\{\langle a \sqcup b \sqcup c \sqcup d, k_2 \rangle\}| = 1 \quad (\text{One fence})$$

The formal implementation of this idea has all lexical Ns associated with a root meaning, which is a subset of a complete atomic Boolean algebra M . Count Ns are typically distinct from mass Ns. Mass N denotations are a subset of root meanings, and of the standard predicative type $\langle e, t \rangle$. Count N denotations are uniformly of the type $\langle e \times k, t \rangle$, i.e., they denote functions from pairs consisting of an individual e and a counting context k , in which that individual counts as ‘one’, to truth values.

Count N denotations are derived from root N denotations by means of a semantic $COUNT_k(N)$ operation which selects, for a counting context k , a set of disjoint ‘semantic atoms’, in default cases.

Rothstein’s analysis of count N denotations has the right motivation for context-sensitive count Ns like *fence*, for which it was developed, but it requires the same ‘indexical’ (ibid., p.362) analysis also for count Ns like *cat*, for which counting operations are not context-sensitive. This predicts a uniform behavior for all count Ns, both semantic and distributional. But this prediction is not borne out. If singular count Ns like *fence* denote only disjoint ‘semantic atoms’ in default counting contexts, as Rothstein (2010) proposes, then it is unexpected that they are felicitous in the pseudo-partitive measure NP (e.g., *three yards of fence*) and with quantifiers like *much* (e.g., *much fence*), both of which reject singular count predicates.

2.3 Landman (2011, 2016): Context-sensitive non-overlap

Landman’s (2011; 2016) theory of the mass/count distinction is motivated by the idea that counting is a matter of non-overlap, or overlap that cannot be made irrelevant in a given context. All Ns have their denotations built from ‘generators’, which are ‘the things that we would want to count as one’. Grammatical counting amounts to counting elements in a generator set. The denotations of count Ns (*cat*) are built from *non-overlapping generators*. Counting in the count domain succeeds, because count Ns only denote sets of disjoint entities. In contrast, mass Ns are built from *overlapping generators*, which leads to *overspecification* (Landman 2011: 17) with respect to how many countable entities there are in their denotation simultaneously in the same context: There is *no single set* of entities that count as ‘one’ in a given context. There are many such sets simultaneously in the same context, and they overlap; and because there are too many things to count and no single way of counting them, counting in the mass domain ‘goes wrong’.

One of Landman’s (2011) innovations is to identify two cases when counting ‘goes wrong’: (i) MESS mass Ns like *water*, *salt*, *meat*, which have denotations built from overlapping minimal generators, i.e., the smallest things we may count as one; (ii) NEAT mass Ns like *kitchenware*, *furniture*, *silverware*, aka ‘object’ or ‘fake’ mass Ns, which have denotations built from overlapping generators, but their overlap is not located in the minimal generators, but at the level of their sums.

The idea that mass Ns are built from overlapping generators makes for a compelling analysis of NEAT mass Ns like *kitchenware*, *furniture*, *silverware*, which serve as Landman’s main data point. Although they denote sets of atomic entities just like prototypical count Ns, they are grammatically mass, so cannot be used in counting constructions: cp. *#three kitchenware(s)* versus *three cups*. Suppose we have a domain with three kitchenware items: a teacup, a saucer and a teapot. There

is more than one way of partitioning this domain into countable units (*variants* in Landman's terminology). For instance, we can count the minimal generators sitting 'at the bottom' of the Boolean algebra associated with *kitchenware*. On this way of counting, the result is 3. Or, we may count a teacup and a saucer together as one unit of kitchenware, a generator (though not a minimal one), and a teapot as another item, and on this way to count, the result is 2. Alternatively, we can view a teacup, a saucer and a teapot all together as having a joint function qua a single tea set, and from this perspective on counting, the counting result is 1. But this means that there is *no single way of counting*, and, as Landman puts it, 'counting goes wrong'. The key idea here is that there is more than one way of partitioning the denotation of NEAT mass Ns like *kitchenware* with respect to what we may view as single countable units *simultaneously in the same context*. If you insist on counting them, you will always end up counting overlapping units, and arrive at different counting results, because the different possible partitions overlap. There is no principled way of ignoring such overlaps, or making them irrelevant, because there is no principled way of privileging one single partition over others.

However, the idea that mass Ns are built from overlapping generators does not fare well in motivating the mass property of MESS mass Ns like *water*, *salt*, *meat*. Landman (2011) assumes an atomistic domain, just like Chierchia (1998, 2010) and Rothstein (2010), but this presupposes that we know, relative to a context, what the minimal elements in the denotation of MESS mass Ns are, and hence a solution to the notorious MINIMAL PARTS PROBLEM. Therefore, Landman (2016) abandons defining MESS mass Ns in terms of overlapping minimal generators, what he there calls *minimal bases*. Instead, mass Ns are MESS mass if they are not neat, which leaves open the possibility that the bases of MESS mass Ns are atomless. This means that Landman (2016) sides with Krifka (1989); Sutton & Filip (2016b) i.e., in adopting a non-atomic mereology.

3 Empirical evidence: Countability and measure constructions

As observed above, grammatical countability of a given N is not necessarily aligned with that N specifying a unique, context-independent criterion of individuation for its application. Although count Ns like *fence*, *twig*, *sequence* do not lexically specify such a criterion, they pattern with count Ns like *cat*, *boy*, *letter*, which do, in so far as they exhibit grammatical count properties, including the following ones: (i) direct modification by numerical expressions (*three cats*, *three fences*); (ii) straightforward acceptability as arguments of count quantifiers (*each cat*, *each fence*); (iii) pluralization (*cats*, *fences*); (iv) occurrence as bare singulars in argument positions is highly restricted (*There was apple in the salad*, *Kim bought *apple/*fence in the store yesterday*, **Apple/*Fence lay on the kitchen counter*).

While the above data clearly speak for the grammatically count status of singular Ns like *fence*, *wall*, *twig*, *branch* or *sequence*, their straightforward acceptability in the measure (pseudo-partitive) NP sets them apart from singular count Ns like *cat*, *baby* or *cabin*. As we see below, *baby* is odd or highly marked in the measure (pseudo-partitive) NP, which may possibly be exploited for special rhetorical effects:

- (12) a. # 6 kilograms of baby
 b. # You can find a heavy piece of baby in the nursery.
- (13) a. Thick woolen drapes of red and gold covered every inch of wall. (COCA)
 b. Thus a cm dry length of twig increased in dry weight by 0.047g. (*Community Ecology of a Coral Cay*, Heatwole et al. p.152)
 c. The cages were 1 foot in diameter and enclosed a 3-foot length of branch. (*California Agriculture*. Mar-Apr, 1989 p.7)
 d. 155 kilometers, or 96 miles, of wall encircled West Berlin (CNN “Berlin wall secrets”)

The above data are puzzling, given that most agree that *no* singular count Ns should be straightforwardly acceptable in the measure (pseudo-partitive) NP (Bach 1981; Krifka 1989; Filip 1992, 2005; Schwarzschild 2002, 2006; Nakanishi 2007; Landman 2016, i.a.). If both *baby* and *fence* exhibit grammatically count behavior, in compliance with the diagnostics in (i)-(iv) above, why are singular count Ns like *fence*, but not like *baby*, perfectly natural and felicitous in the measure (pseudo-partitive) NP? If singular count Ns such as *fence* freely occur in count syntax, which depends on some notion of quantization (Krifka 1989), semantic atomicity (Rothstein 2010) or disjointness (Landman 2011, 2016), for example, how can it be that such Ns are also felicitous in the pseudo-partitive measure NP, which does not welcome singular predicates that are quantized, semantically atomic or disjoint?

4 Formal analysis

Cumulativity (see (6) in Section 2.1) plays a key role in characterizing the application domain of extensive measure phrases, i.e., phrases consisting of a numerical or some other weak quantifier followed by an extensive measure expression like *gram(s)*, *kilometer(s)*. The application domain of extensive measure phrases is commonly assumed to be restricted to mass or plural count predicates, so satisfying the

CUM property (also Krifka 1998, i.a.). We, however, argue that it is better viewed as satisfying a weaker $\neg$ QUA property, as Krifka (1989) (see Section 2.1) originally proposed, albeit for different reasons. Our main motivation is to account for the felicitous use of singular count Ns like *fence* in the pseudo-partitive (measure) NP, given that they do not (necessarily) denote CUM predicates, but $\neg$ QUA

predicates. Non-quantization is weaker than cumulativeness, because $CUM(P)$ asymmetrically entails $\neg QUA(P)$. It is not possible for a predicate to be *both* cumulative and quantized, but it is possible for a predicate to be *both* non-quantized and not cumulative.

Neither would Landman's *non-disjointness*, or *overlap*, be a better candidate for specifying the input condition of measure phrases. There is a subtle, albeit important, difference between Krifka's *quantization* and Landman's *disjointness*. It shows up in both measuring and counting constructions. Take, for instance, a predicate with the set $\{a \sqcup b \sqcup c, c \sqcup d \sqcup e\}$ in its denotation. Such a predicate is quantized in Krifka's sense, but not disjoint in Landman's sense, because its members overlap in the c element. Similarly, two sides of fencing that share a corner post may satisfy the description *two fences* (see also above), as most would agree, despite their overlap in the corner post. (It is not entirely clear whether this overlap can or cannot be made irrelevant, *pace* Landman 2011.) But this means that in order to motivate the application of counting operation over a given domain, requiring its *disjointness* seems too strict, and *quantization* relative to a predicate is the better notion. Moreover, if *two fences* may refer to two sides of fencing overlapping in a corner post, and if the restriction on measuring were *overlap* or *not disjointness*, then this would incorrectly predict that *#two meters of two fences* should be acceptable. Therefore, also when it comes to measuring, *non-quantization* of the predicate denoting what is measured is the better notion than *non-disjointness*.

From the point of view the property of quantization, we may rephrase our main QUESTION 2 as follows: 'How can it be that count Ns such as *fence* can be directly modified by numerals, which presupposes that they denote quantized predicates, and also be felicitous in the pseudo-partitive (measure) NP which presupposes that they denote non-quantized predicates?' Our answer, in broad terms, is that Ns such as *fence* are special in the following way:

The predicates for e.g., single fences are quantized at each specific counting context, but *non-quantized* at the null counting context.

The definitions of *specific counting context* and *null counting context* will be introduced in the next section.

4.1 Common nouns and counting contexts

Sutton & Filip (2016a,b) and Landman (2011, 2016) propose to treat common noun lexical entries as ordered pairs. For Landman (2016), this is $\langle \textit{body}, \textit{base} \rangle$ in which *base* is the counting base predicate, and *body* is a subset of the upward closure of the *base* under mereological sum. For example, the entry for the plural noun *cats* is given as $\langle *CAT, CAT \rangle$. Sutton & Filip (2016a,b) argue the need to explicitly

represent an individuation function (**IND**) and counting contexts c in counting base predicates. Individuation functions apply to number neutral predicates P and return the set of entities that can count as one P . With the addition of a counting context c , (schematically, $c(\mathbf{IND}(P))$) we get the set of entities that count as one P at a counting context c (i.e., under a particular counting perspective).

Here, we combine elements of these two approaches. We assume a predicate for the *extension* of the noun, and a *counting base* predicate. Both the counting base predicate and the extension predicate may include the **IND** function and a counting context. We also add a third projection to our lexical entries to track presuppositions (or, perhaps more neutrally, *preconditions*) for composition, such as Krifka's (1989) restriction that extensive measure phrases compose with only non-quantized predicates. Consequently, lexical entries of common Ns, in our account, are of the form $\langle extension, c_base, preconditions \rangle$.²

Before we provide examples of lexical entries, we give some details about these two important features of our analysis: *the IND function* and *counting contexts*. The **IND** function is of type $\langle \langle e, t \rangle, \langle e, t \rangle \rangle$. It applies to a number neutral predicate P of type $\langle e, t \rangle$ and returns the set of entities that count as one with respect to that predicate. (So $\mathbf{IND}(P)$ is of type $\langle e, t \rangle$.) Whether **IND** is included in a lexical entry is sensitive to the prelinguistic object versus stuff distinction (see Soja, Carey & Spelke 1991, also Barner & Snedeker 2005; Rothstein 2010, i.a.).

For mass Ns like *mud* that denote undifferentiated stuff, there is, intuitively, nothing that clearly counts as 'one'. Therefore, the counting base predicate, lacks the **IND** function, amounting to the number neutral predicate P .

For Ns which denote Spelke objects or collections thereof, the counting base predicate is derived by means of the **IND** function, and it is introduced into the lexical entries, including those denoted by 'fake' or 'object' mass Ns like *furniture*, and *jewelry*. So applying **IND** to a number neutral predicate P yields a set of entities that either (i) is always maximally disjoint, or (ii) includes overlapping variants (maximally disjoint sets). Examples of case (i) include Ns that are 'naturally atomic' in that they denote clearly individuable entities, such as Ns denoting natural kinds (*cat*, *apple*). Examples of case (ii) include Ns that are collective artifact Ns (*furniture*, *jewelry*) and also Ns which refer to objects, but which rely on context to determine what is taken to be one entity in their denotation (*fence*, *hedge*). For example, a pair of earrings can count as one item of jewelry or as two (Landman 2011), and fencing around a field can count as one fence, or as four fences. (For a fuller discussion of these noun classes, see, amongst others, Sutton & Filip 2016a,b.)

Building on some independent suggestions in Rothstein 2010 and Landman 2011,

² In order to modify these different projections in composition with other expressions, we use the projection functions π_1 , π_2 , and π_3 such that:

if $X = \langle \phi, \psi, \chi \rangle_{\langle a \times b \times c \rangle}$, then $\pi_1(X) = \phi_a$, $\pi_2(X) = \psi_b$ and, $\pi_3(X) = \chi_c$

we assume a set of SPECIFIC COUNTING CONTEXTS $\mathcal{C} = \{c_1, \dots, c_n\}$ which apply to possibly overlapping sets and yield maximally disjoint subsets. The latter, maximally disjoint subsets, roughly correspond to Landman’s (2011) *variants*, or *counting contexts* (or ‘counting perspectives’) in Rothstein’s sense (for more detail see Sutton & Filip 2016b). Specific counting contexts are functions of type $\langle\langle e, t \rangle, \langle e, t \rangle\rangle$ (which is in a slight departure from Sutton & Filip 2016b. (So $c(\mathbf{IND}(P))$ is of type $\langle e, t \rangle$.) For example, for the overlapping set in (14), there are two possible counting contexts c_1 and c_2 , given in (15) and (16), and each contains only non-overlapping entities that count as ‘one’.

$$(14) \quad \mathbf{IND}(P) = \{\text{cup}, \text{saucer}, \text{cup} \sqcup \text{saucer}\}$$

$$(15) \quad c_1(\mathbf{IND}(P)) = \{\text{cup}, \text{saucer}\}$$

$$(16) \quad c_2(\mathbf{IND}(P)) = \{\text{cup} \sqcup \text{saucer}\}$$

Our null counting context c_0 is defined as a function on sets that returns the the union of interpretations of that set across all specific counting contexts:

$$(17) \quad c_0(X) = \bigcup_{c_i \in \mathcal{C}} (c_i(X))$$

This also implies that the null counting context is the identity function on sets (which is why it is ‘null’):

$$(18) \quad c_0(X) = X$$

As our simple examples in (14-16) illustrate, depending on the specific counting context, we may count the singularities, the cup and the saucer each individually as ‘one’, but also their sum as ‘one’; and these two counting variants overlap. Our null counting context thus captures Landman’s (2011) notion of *overlapping variants, simultaneously in the same context* (see above).

Ns are assigned interpretations relative to either a specific counting context $c_i \in \mathcal{C}$ or the null counting context c_0 . Together with the semantic properties of their counting bases, which may be maximally disjoint sets (e.g., *cat*) or possibly overlapping sets (e.g., *jewelry*), this allows us, for a large class of N concepts, to make predictions about their encoding as grammatically count or mass Ns, and also when we can(not) expect to find mass/count variation cross- and intralinguistically, as we show in Sutton & Filip 2016a,b.

Let us now sketch in some detail how these lexical assumptions allow us to specify lexical entries of mass and count Ns starting with ‘sortal’ count Ns like *cat*, *apple*, *boy*, and analogous terms cross-linguistically. Based on our world knowledge, knowing what *cat*, means includes knowing what is one whole cat in any context,

and what it is remains stable across all contexts. Such count Ns, therefore, have **IND** sets that are maximally disjoint (non-overlapping) when interpreted at any specific counting context $c_i \in \mathcal{C}$, and consequently, also when interpreted at the null counting context c_0 . Their counting base predicates are quantized, which motivates that such Ns are grammatically count, and stably so cross-linguistically.

In contrast, Ns like *mud* denote stuff which contains nothing that clearly and plausibly counts as ‘one’, even at specific counting contexts (barring possibly highly specialized scientific or technical contexts). As Ns like *mud* provide no intuitive individuation schema, there is no **IND** function in their lexical entry, as we see in (19). Such predicates are non-quantized, and yield mass concepts. Given that these properties do not vary with particular context, the context of evaluation for their counting base predicates is fixed as the null counting context. This is done by having their lexical entries saturated with the null counting context.³

$$(19) \quad \begin{aligned} \llbracket \text{mud} \rrbracket^{c_i} &= \lambda c. \lambda x. \langle c_0(\text{MUD})(x), \lambda y. c_0(\text{MUD})(y), \emptyset \rangle (c_i) \\ &= \lambda x. \langle c_0(\text{MUD})(x), \lambda y. c_0(\text{MUD})(y), \emptyset \rangle \end{aligned}$$

As already mentioned, Spelke object-denoting Ns which introduce a counting base predicate derived by means of the **IND** function, may still depend on context for specifying what is ‘one’ in their denotation, because it may concern not only singularities, but also pluralities of objects (see (15) and (16) above). Such Ns include collective artifact mass Ns (*furniture, jewelry*) and also count Ns such as *fence*, which are the focus of this paper. Given that what counts as one fence is relative to a specific counting context $c_i \in \mathcal{C}$ and may vary in size from context to context, at the c_0 , its counting base **IND** predicate will not be quantized. At specific counting contexts $c_i \in \mathcal{C}$, in which we ‘carve out’ maximally disjoint fence chunks for the purposes of counting, the counting base **IND** predicate will be lexicalized by the count N *fence*, as we see in (20). This also predicts that when interpreted at the null counting context c_0 , the counting base **IND** predicate will yield a mass noun concept, which in English is lexicalized by *fencing*, as we see in (21).

$$(20) \quad \begin{aligned} \llbracket \text{fence} \rrbracket^{c_i} &= \lambda c. \lambda x. \langle c(\text{IND}(\text{FENCE}))(x), \lambda y. c(\text{IND}(\text{FENCE}))(y), \emptyset \rangle (c_i) \\ &= \lambda x. \langle c_i(\text{IND}(\text{FENCE}))(x), \lambda y. c_i(\text{IND}(\text{FENCE}))(y), \emptyset \rangle \end{aligned}$$

$$(21) \quad \begin{aligned} \llbracket \text{fencing} \rrbracket^{c_i} &= \lambda c. \lambda x. \langle c_0(\text{IND}(\text{FENCE}))(x), \lambda y. c_0(\text{IND}(\text{FENCE}))(y), \emptyset \rangle (c_i) \\ &= \lambda x. \langle c_0(\text{IND}(\text{FENCE}))(x), \lambda y. c_0(\text{IND}(\text{FENCE}))(y), \emptyset \rangle \end{aligned}$$

4.2 Direct numerical attachment to common nouns

Numerical expressions in English (*one, two*, etc.) denote numerals: $\llbracket \text{one} \rrbracket = 1$, $\llbracket \text{two} \rrbracket = 2$ etc. For number marking languages, such as English, there is a generally

³ We use ‘ $\emptyset$ ’ for the empty proposition.

available type-shifting operation (MOD) which shifts numerical expressions denoting numerals into numerical modifiers, which can then compose with common Ns.

$$\text{MOD} = \lambda n. \lambda P. \lambda x. \langle \pi_1(P(x)), \mu_{card}(x, \pi_2(P(x)) = n, \text{QUA}(\pi_2(P(x))) \rangle$$

$\mu_{card}(x, Q)$ is a cardinality function that maps an entity x and a property Q onto a numeral n (the number of Q s that x is). Numerical determiners also come with a precondition that the property Q is quantized *QUA*.

For number marking languages, such as English, numerical expressions which denote numerals can be shifted using MOD and then compose with common Ns. For example, for *three fences*, the result is a set of entities that have fence properties, to the amount of three relative to the counting of fences in context c_i , and presupposes that the property $c_i(\text{IND}(\text{FENCE}))$ is quantized. This is shown in (22a-22d). The *-operator indicates upwards closure under mereological sum.

$$(22a) \llbracket \text{three} \rrbracket^{c_i} = 3$$

$$(22b) \text{MOD}(\llbracket \text{three} \rrbracket^{c_i}) = \lambda P. \lambda x. \langle \pi_1(P(x)), \mu_{card}(x, \pi_2(P(x)) = 3, \text{QUA}(\pi_2(P(x))) \rangle$$

$$(22c) \llbracket \text{fences} \rrbracket^{c_i} = \lambda x. \langle {}^*c_i(\text{IND}(\text{FENCE}))(x), \lambda y. c_i(\text{IND}(\text{FENCE}))(y), \emptyset \rangle$$

$$(22d) \llbracket \text{three fences} \rrbracket^{c_i} = \text{MOD}(\llbracket \text{three} \rrbracket^{c_i})(\llbracket \text{fence(s)} \rrbracket^{c_i}) = \\ \lambda x. \langle {}^*c_i(\text{IND}(\text{FENCE}))(x), \mu_{card}(x, \lambda y. c_i(\text{IND}(\text{FENCE}))(y)) = 3, \\ \text{QUA}(\lambda y. c_i(\text{IND}(\text{FENCE}))(y)) \rangle$$

However, were one to try to apply *three* to *fencing* this *quantized* precondition would be false, since $c_0(\text{IND}(\text{FENCE}))$ is non-quantized (it denotes entities that are proper parts of each other), and so the result would be infelicitous.

4.3 Pseudo-partitive measure constructions

Any satisfactory analysis of the pseudo-partitive (measure) NP must correctly predict that it is felicitous with all mass and plural terms, assuming the measure is appropriate (*three pounds of mud* versus *#three seconds of mud*), but not with most singular count Ns (*#three kilos of cat*), while sanctioning a certain sizable class of singular count Ns (*300 meters of fence*). In a nutshell, we combine the distinction between specific counting contexts and the null counting context (Sutton & Filip 2016a,b) with Krifka's (1989) analysis of the pseudo-partitive (measure) NP as requiring an extension predicate that is non-quantized. Singular count Ns like *fence* are felicitous in the pseudo-partitive (measure) NP, because, if interpreted at the null counting context, they have non-quantized extension predicates. In this respect, singular count Ns like *fence* pattern with mass Ns like *mud*, and differ from singular count Ns like *cat*. The key differences between these three classes of Ns are summarised in

	Quantized counting base predicate in the lexicon	Quantized extension at c_0	Measure NP accept- ability
<i>cat</i>	Yes	Yes	No
<i>cats</i>	Yes	No	Yes
<i>fence</i>	Yes	No	Yes
<i>mud</i>	No	No	Yes

Table 1 Measure NP acceptability and quantization property at c_0

Table 1. Both *cat* and *fence*, but not *mud*, have counting base predicates which are quantized as parts of their lexical entries. As we saw in Section 4.2, this allows the direct attachment of numerical expressions to *cat* and *fence*, but not to *mud*. In contrast, at the null counting context c_0 , the extensions of *fence* and *mud*, but also plural Ns like *cats*, are non-quantized, whereas the counting base predicate of the singular N *cat* is quantized. This, we contend, is the property that allows the singular count N *fence*, the plural N *cats* and the mass N *mud* to felicitously appear in the pseudo-partitive (measure) NP and explains why the singular count N *cat* is not felicitous in this construction.

We propose that words for extensive measure phrases (like *meter*, *kilo*) introduce the presupposition on nominal arguments with which they compose to form a pseudo-partitive (measure) NP have a non-quantized extension predicate when this predicate is interpreted at the null counting context c_0 . This will, in line with data, filter out singular count Ns such as *cat*, but allow mass Ns such as *mud*, plural count Ns such as *cats*, but also singular count Ns such as *fence*. Notice that although the lexical entry for *fence* indicates the specific counting context of utterance $c_i \in \mathcal{C}$ as the context of evaluation, what matters for the felicitous application of an extensive measure phrase like *three meters (of)* is that its extension predicate is non-quantized at the at the null counting context c_0 . Part of the job of extensive measure phrases can be seen as applying the null counting context to its nominal argument before it is evaluated at the context of utterance, then check, as a presupposition of composition, that its extension is non-quantized.

We give a derivation for *three meters of fence* in (23a-23d):

$$(23a) \llbracket \text{three} \rrbracket^{c_i} = 3$$

$$(23b) \llbracket \text{meters of} \rrbracket^{c_i} = \lambda n. \lambda P. \lambda d. \lambda x.$$

$$\langle \pi_1(P(c_0)(x)), \mu_m(x, d) = n, \neg \text{QUA}(\lambda y. \pi_1(P(c_0)(y))) \rangle$$

$$(23c) \llbracket \text{fence} \rrbracket = \lambda c. \lambda x. \langle c(\mathbf{IND}(\text{FENCE}))(x), \lambda y. c(\mathbf{IND}(\text{FENCE}))(y), \emptyset \rangle$$

$$(23d) \llbracket \text{three meters of fence} \rrbracket^{c_i} = \llbracket \text{meters of} \rrbracket^{c_i}(\llbracket \text{three} \rrbracket^{c_i})(\llbracket \text{fence} \rrbracket) = \lambda d. \lambda x. \\ \langle c_0(\mathbf{IND}(\text{FENCE}))(x), \mu_m(x, d) = 3, \neg \text{QUA}(\lambda y. c_0(\mathbf{IND}(\text{FENCE}))(y)) \rangle$$

Measure words that denote extensive measure functions (Krifka 1989) (e.g. *meter*, *kilo*) take a numeral and a noun denotation as arguments. Unlike in the direct numerical modification case, numerical words (*three*) do not need to be shifted into numerical modifiers by MOD, but instead denote numerals: $\llbracket \text{three} \rrbracket = 3$. The combination of a measure phrase and numeral with a noun denotation has the following effect (where (i)-(iii) tally with the first, second and third projection of the resulting tuple): (i) The base extension is the extension of the noun at the null counting context. (ii) The extension is restricted to entities that measure n with respect to the extensive measure function. For words denoting extensive measure functions like *meter*, we also assume a contextually provided dimension, d , since, for example, *three meters of fence was before us* could refer to three meters of fence in height or in length. (iii) There is a presupposition that the extension of the argument noun denotation is non-quantized when evaluated at the null counting context. In (23d), this presupposition is satisfied, so the measure phrase is felicitous.

If we try to compose e.g. *three kilos of* with the singular count N *cat* as in (24a-24c), the semantics of *kilo* introduces the presupposition that the extension predicate of *cat* is non-quantized at the null counting context. However, **IND**(CAT) (the set of single cats) is quantized. When the *non-quantized* presupposition is false, our analysis predicts that the measure NP will not be felicitous. In contrast (24d), the extension predicate of the plural *cats* is non-quantized at c_0 , thus *three kilos of cats* is felicitous.

$$(24a) \llbracket \text{kilos of} \rrbracket = \lambda n. \lambda P. \lambda x.$$

$$\langle \pi_1(P(c_0)(x)), \mu_{kg}(x) = n, \neg \text{QUA}(\lambda y. \pi_1(P(c_0)(y))) \rangle$$

$$(24b) \llbracket \text{cat(s)} \rrbracket = \lambda c. \lambda x. \langle {}^{(*)}c(\mathbf{IND}(\text{CAT}))(x), \lambda y. c(\mathbf{IND}(\text{CAT}))(y), \emptyset \rangle$$

$$(24c) \llbracket \# \text{three kilos of cat} \rrbracket^{c_i} =$$

$$\lambda x. \langle c_0(\mathbf{IND}(\text{CAT}))(x), \mu_{kilo}(x) = 3, \neg \text{QUA}(\lambda y. c_0(\mathbf{IND}(\text{CAT}))(y)) \rangle$$

$$(24d) \llbracket \text{three kilos of cats} \rrbracket^{c_i} =$$

$$\lambda x. \langle {}^*c_0(\mathbf{IND}(\text{CAT}))(x), \mu_{kilo}(x) = 3, \neg \text{QUA}({}^*\lambda y. c_0(\mathbf{IND}(\text{CAT}))(y)) \rangle$$

5 Conclusion

We have identified a distinctive lexical property of count Ns like *fence* that explains what differentiates them from count Ns like *cat* and also from mass Ns like *mud*. Intuitively, *fence*-like Ns admit of multiple alternative individuation or counting schemas simultaneously at any context, while bona fide count Ns like *cat* are associated with only one in all contexts. Ns like *mud* provide no intuitive individuation

schema for the stuff they denote, no ‘instruction’ about what could plausibly and reliably be identified as ‘one’ in their denotation.

On our analysis, we propose that Ns like *fence* and *cat* denote quantized counting base predicates at specific counting contexts, and are therefore grammatically count. At the same time, at the null counting context c_0 , which is the union of interpretations across all specific counting contexts, the extension predicate of *fence*-like Ns is non-quantized, due to the overlap among members across alternate individuation/counting schemas. The extension predicates of plural count nouns are non-quantized, simply because sums of individuals are included in their denotations. The extension predicates of mass Ns like *mud* are also non-quantized, but due simply to the lack of any individuation/counting schema. This, we contend, is the property that allows singular count Ns like *fence*, plural count nouns like *cats*, and mass Ns like *mud* to felicitously appear in the pseudo-partitive measure NP. In contrast, singular count Ns like *cat* are not felicitous in the measure NP, because their extension predicates are quantized at the null counting context, as there is only one individuation/counting scheme that reliably applies to any entities across all specific counting contexts.

We are still left with some puzzles, however. Not all count Ns that fail to be quantized are felicitous in the pseudo-partitive (measure) NP. While *fence*-like Ns are perfectly acceptable in this context, mathematical concepts (*sequence*, *line*, *arc*) and classifier-like Ns like *piece*, *quantity* or *bouquet* are not. We suspect that the infelicity of singular count Ns denoting mathematical concepts (e.g., *# three rows of sequence*) may be connected to the abstract, mathematical nature of these concepts. The infelicity of classifier-like Ns (e.g., *# three lbs of piece of ice*) might have to do with the bar on measuring the same thing twice (Bach 1981). Finally, we also observe intriguing differences among the *fence*-like count Ns. For instance, only some are also compatible with quantifiers like *much* whose application is restricted to mass Ns, at least in some contexts: cf. *How much fence do you need?*, *‘The question is how much wall do you need?’ Goodlatte said. ‘You won’t build a 30-foot high wall along 2000 miles of southern border.’ (The News Virginian, October 18, 2017, “Goodlatte sees drug tunnels and potential fence on border visit.”)*

References

- Bach, Emmon. 1981. On time, tense, and aspects: An essay in english metaphysics. In Peter Cole (ed.), *Radical Pragmatics*, 63–81. Academic Press.
- Barner, David & Jesse Snedeker. 2005. Quantity judgments and individuation: evidence that mass nouns count. *Cognition* 97(1). 41–66. doi:10.1016/j.cognition.2004.06.009.
- Champollion, Lucas & Manfred Krifka. 2016. Mereology. In Maria Aloni (ed.), *The*

- Cambridge Handbook of Formal Semantics*, 513–541. Cambridge University Press. doi:[10.1017/cbo9781139236157.014](https://doi.org/10.1017/cbo9781139236157.014).
- Chierchia, Gennaro. 1998. Plurality of mass nouns and the notion of “semantic parameter”. In Susan Rothstein (ed.), *Events and Grammar: Studies in Linguistics and Philosophy Vol. 7*, 53–103. Kluwer. doi:[10.1007/978-94-011-3969-4_4](https://doi.org/10.1007/978-94-011-3969-4_4).
- Chierchia, Gennaro. 2010. Mass nouns, vagueness and semantic variation. *Synthese* 174. 99–149. doi:[10.1007/s11229-009-9686-6](https://doi.org/10.1007/s11229-009-9686-6).
- Filip, Hana. 1992. Aspect and interpretation of nominal arguments. In C.P. Canakis, G.P. Chan & J.M. Denton (eds.), *Twenty-Eighth Meeting of the Chicago Linguistic Society*, 139–158. Chicago: University of Chicago.
- Filip, Hana. 2005. Measures and indefinites. In Gregory N. Carlson & Francis Jeffry Pelletier (eds.), *References and Quantification: The Partee Effect*, 229–288. Stanford: CSLI.
- Krantz, H., David, R. Duncan Luce, Patrick Suppes & Amos Tversky. 1971. *Foundations of measurement: Additive and polynomial representations*. Academic Press. doi:[10.1016/c2013-0-11004-5](https://doi.org/10.1016/c2013-0-11004-5).
- Krifka, Manfred. 1986. *Nominalreferenz und Zeitkonstitution. Zur Semantik von Massentermen, Pluraltermen und Aspektklassen*. Universität München PhD dissertation. Doctoral Dissertation.
- Krifka, Manfred. 1989. Nominal reference, temporal constitution and quantification in event semantics. In Renate Bartsch, J. F. A. K. van Benthem & P. van Emde Boas (eds.), *Semantics and Contextual Expression*, 75–115. Foris Publications.
- Krifka, Manfred. 1995. Common nouns: A contrastive analysis of English and Chinese. In G.N. Carlson & F. J. Pelletier (eds.), *The Generic Book*, 398–411. Chicago University Press.
- Krifka, Manfred. 1998. The origins of telicity. In Susan Rothstein (ed.), *Events and grammar*, 197–235. Dordrecht: Kluwer.
- Landman, Fred. 2011. Count Nouns–Mass Nouns–Neat Nouns–Mess nouns. *The Baltic International Yearbook of Cognition, Logic and Communication* 6. 1–67.
- Landman, Fred. 2016. Iceberg semantics for count nouns and mass nouns: The evidence from portions. *The Baltic International Yearbook of Cognition Logic and Communication* 11. 1–48.
- Link, Godehard. 1983. The logical analysis of plurals and mass terms: A lattice-theoretic approach. In Rainer Bäuerle, Urs Egli & Arnim von Stechow (eds.), *Meaning, Use and the Interpretation of Language*, 303–323. Berlin: de Gruyter.
- Nakanishi, Kimiko. 2007. Measurement in the nominal and verbal domains. *Linguistics and Philosophy* 30(2). 235–276. doi:[10.1007/s10988-007-9016-8](https://doi.org/10.1007/s10988-007-9016-8).
- Pelletier, Francis Jeffry. 1975. Non-singular reference: Some preliminaries. *Philosophica* 5(4). 451–465.

- Quine, W. V. 1960. *Word and Object*. The Mit Press.
- Rothstein, Susan. 2010. Counting and the mass/count distinction. *Journal of Semantics* 27(3). 343–397. doi:[10.1093/jos/ffq007](https://doi.org/10.1093/jos/ffq007).
- Rothstein, Susan. 2011. Counting, measuring and the semantics of classifiers. *The Baltic International Yearbook of Cognition, Logic and Communication* 6. 1–41.
- Rothstein, Susan. 2017. *Semantics for Counting and Measuring*. Key Topics in Semantics and Pragmatics. Cambridge University Press. doi:[10.1017/9780511734830](https://doi.org/10.1017/9780511734830).
- Schwarzschild, Roger. 2002. The grammar of measurement. *Semantics and Linguistic Theory (SALT)* 12. 225–245. doi:[10.3765/salt.v0i0.2870](https://doi.org/10.3765/salt.v0i0.2870).
- Schwarzschild, Roger. 2006. The role of dimensions in the syntax of noun phrases. *Syntax* 9. 67–110. doi:[10.1111/j.1467-9612.2006.00083.x](https://doi.org/10.1111/j.1467-9612.2006.00083.x).
- Soja, Nancy, Susan Carey & Elizabeth Spelke. 1991. Ontological categories guide young children's inductions of word meaning: Object terms and substance terms. *Cognition* 38. 179–211. doi:[10.1016/0010-0277\(91\)90051-5](https://doi.org/10.1016/0010-0277(91)90051-5).
- Sutton, Peter & Hana Filip. 2016a. Counting in context: Count/mass variation and restrictions on coercion in collective artifact nouns. *Semantics and Linguistic Theory (SALT)* 26. 350–370. doi:[10.3765/salt.v26i0.3796](https://doi.org/10.3765/salt.v26i0.3796).
- Sutton, Peter R. & Hana Filip. 2016b. Mass/count variation, a mereological, two-dimensional semantics. *The Baltic International Yearbook of Cognition Logic and Communication* 11. 1–45.
- Zucchi, Sandro & Michael White. 1996. Twigs, sequences and the temporal constitution of predicates. *Semantics and Linguistic Theory (SALT)* 6. 223–270. doi:[10.3765/salt.v0i0.2774](https://doi.org/10.3765/salt.v0i0.2774).
- Zucchi, Sandro & Michael White. 2001. Twigs, sequences and the temporal constitution of predicates. *Linguistics and Philosophy* 24(2). 223–270. <https://doi.org/10.1023%2Fa%3A1005690022190>.