

Signatures of Individuation Across Objects and Events

Sarah Hye-yeon Lee¹, Yue Ji², and Anna Papafragou¹

¹Department of Linguistics, University of Pennsylvania

²Department of English, School of Foreign Languages, Beijing Institute of Technology

The physical world provides humans with continuous streams of experience in both space and time. The human mind, however, can parse and organize this continuous input into discrete, individual units. In the current work, we characterize the representational signatures of basic units of human experience across the spatial (object) and temporal (event) domains. We propose that there are three shared, abstract signatures of individuation underlying the basic units of representation across the two domains. Specifically, individuated entities in both the spatial domain (objects) and temporal domain (bounded events) resist restructuring, have distinct parts, and do not tolerate breaks; unindividuated entities in both the spatial domain (substances) and the temporal domain (unbounded events) lack these features. In three experiments, we confirm these principles and discuss their significance for cognitive and linguistic theories of objects and events.

Public Significance Statement

Humans are able to parse and organize continuous streams of experience into individual mental units, such as objects (entities that extend over space) and events (entities that extend over time). The nature of these basic units of human experience is a foundational topic for the sciences of the mind. This study shows that mental units are connected across the spatial and temporal domains: Both objects and events are represented in terms of their structure.

Keywords: object, event, individuation, ontological categories, boundedness

Supplemental materials: <https://doi.org/10.1037/xge0001581.supp>

The physical world provides humans with continuous streams of experience in both space and time. The human mind, however, can parse and organize this continuous input into discrete, individual units. In the spatial domain, objects are understood as fundamental units for many human cognitive processes from infancy to adulthood (Piaget, 1955; Quine, 1960; Scholl, 2001, Spelke, 1988, 1994). In the temporal domain, events are considered the foundational entities for human perception and cognition across the human lifespan (Baldwin et al., 2001; Zacks & Tversky, 2001).

Despite the rich tradition of research in both object and event cognition, the nature of the basic units within each of these two domains—the mental ontologies within space and time—are not yet well understood: Indeed, there is no single precise definition of what an object or an event is in the respective bodies of literature (e.g., Scholl, 2001; Yates et al., 2023). In the current work, we aim to characterize the representational signatures of basic units of human experience across the spatial and temporal domains. We propose that there are similar, abstract considerations of individuation underlying the basic units of representation across the two domains. Moreover, these considerations allow natural connections to how spatial and temporal entities of different types are encoded in language. We begin by outlining the literature on ontological distinctions within the domains of space and time before we propose an account of the parallels between the two domains. We then turn to a series of experiments that test our account (Experiment 1: No Restructuring, Experiment 2: Distinct Parts, Experiment 3: No Breaks).

Ontological Distinctions in the Spatial Domain

One of the most fundamental ontological distinctions that the human mind makes of the physical world in space is that between objects (e.g., table, ball) and substances (e.g., sand, water). This discussion originates back in Aristotle's analysis of form and matter (Aristotle, 350 B.C.E./1994)—we can say that objects have both matter and form, while substances themselves are matter yet lack an inherent form.

This article was published Online First June 6, 2024.
Agnieszka Konopka served as action editor.

Sarah Hye-yeon Lee  <https://orcid.org/0000-0002-6130-6244>

This work was supported by the National Science Foundation (Grant 2041171, to Anna Papafragou). All the raw data and the code behind the analyses have been made available on the Open Science Framework and can be accessed at <https://osf.io/bwkdv/> (Lee, Ji, & Papafragou, 2023).

Sarah Hye-yeon Lee served as lead for conceptualization, data curation, formal analysis, investigation, methodology, project administration, resources, visualization, writing—original draft, and writing—review and editing. Yue Ji served in a supporting role for conceptualization, resources, supervision, and writing—review and editing. Anna Papafragou served as lead for funding acquisition and supervision and served in a supporting role for conceptualization and writing—review and editing.

Correspondence concerning this article should be addressed to Sarah Hye-yeon Lee, Department of Linguistics, University of Pennsylvania, 3401-C Walnut Street, Suite 300, C Wing, Philadelphia, PA 19104-6228, United States. Email: sarahhl@sas.upenn.edu

Objects and substances differ in terms of infants and adults' ability to track and quantify over them (e.g., Chiang & Wynn, 2000; Hespos et al., 2009; Huntley-Fenner et al., 2002; Rosenberg & Carey, 2006; Soja et al., 1991; van Marle & Scholl, 2003; van Marle & Wynn, 2011). For instance, infants have different expectations for how solid objects and nonsolid substances behave; solid objects often keep their shape over changes in position, but nonsolids, such as water or sand, often deform as they move (e.g., Hespos et al., 2009). Research also shows that quantificational abilities are impaired in substances relative to objects, for both infants and adults (e.g., Huntley-Fenner et al., 2002; van Marle & Scholl, 2003; van Marle & Wynn, 2011). For example, at the same age that infants can successfully represent the precise number of objects that appear on a stage or that have disappeared behind an occluder, they fail to represent the quantity of noncohesive substance (such as sand) poured onto the display in tasks that are otherwise identical (Huntley-Fenner et al., 2002; Rosenberg & Carey, 2006). In all, this literature suggests that discrete objects are critical units of processing for entity tracking and quantification.

Research in vision science also suggests that discrete objects are critical units of processing in the domain of adult attention (object-based attention; see Scholl, 2001, for a review). Within this research tradition, objects are treated as a basic unit of visual attentional selection, giving rise to object-based attention effects; for instance, features belonging to the same object are easier to respond to than features belonging to two different objects (Duncan, 1984). However, the precise definition of an "object" is still under debate (Scholl, 2001).

Similar definitional questions emerge when one considers whether or how the object–substance distinction is mapped onto language. Many authors have proposed that count nouns denote objects and mass nouns denote substances (e.g., Bloom, 1999; Gordon, 1985; Link, 1983). In English, for example, objects are usually denoted with count nouns (e.g., *a table*) and substances are typically denoted with mass nouns (e.g., *wood*). However, the count–mass linguistic distinction does not correlate perfectly with the conceptual object–substance distinction; for example, both the count noun *cows* and the mass noun *cattle* can be used to refer to the same groups of objects. It appears, therefore, that count nouns denote individuals but mass nouns are unspecified for individuation (and could, under certain circumstances, individuate). Indeed, studies of quantity judgments in 4-year-olds and adults demonstrate that some mass nouns (*furniture*) do denote individuals (Barner & Snedeker, 2005). These intuitions about individuation are reflected in the fact that people rely on counting to perform comparisons involving count but not mass nouns: A question such as "Who owns more books?" is answered by simply counting the books but a question such as "Who has more milk?" is answered by checking the volume or weight of milk rather than considering number (Barner et al., 2008).

Ontological Distinctions in the Temporal Domain

In the domain of temporal entities, there has been less focus on what defines the basic unit of experience. Event cognition research has placed emphasis on how people identify event boundaries (Zacks et al., 2007), but less attention has been directed towards the representational units within event boundaries, or the basic mental ontology of events.

A recent proposal has addressed this gap by building on logico-philosophical discussions of events going back to Aristotle (Ji & Papafragou, 2020a). This proposal distinguishes between events that are internally structured in terms of distinct temporal stages and have a well-defined endpoint (*bounded* events; e.g., piling up a deck of cards) and events that are internally unstructured and lack an inherent endpoint (*unbounded* events; e.g., shuffling a deck of cards). This proposed mental ontology is supported by evidence that viewers extract boundedness information when processing naturalistic visual events. In one demonstration (Ji & Papafragou, 2020a), participants watched videos of a character performing everyday actions; some videos were marked by a red frame in a way that corresponded to either the bounded (e.g., dress a teddy bear) or the unbounded event (e.g., pat a teddy bear) category. The participants succeeded in identifying whether the red frame applied to a new set of events. Furthermore, when asked to indicate what kind of event was marked by a red frame, they were likely to mention the structure and organization (or lack thereof) of the events, thereby showing sensitivity to an essential dimension of the bounded–unbounded distinction (cf., also Ji & Papafragou, 2020b). Further work has shown that boundedness affects how people perceive events online, even when not required by the task (Ji & Papafragou, 2022), and is therefore a spontaneous feature of event apprehension. Moreover, boundedness is present in 4-year-olds' nonlinguistic event construals (Ji & Papafragou, 2019, 2020b).

In language, the bounded/unbounded event distinction is mapped broadly speaking onto telicity (Binnick, 2012; Jackendoff, 1991; Mourelatos, 1978; Parsons, 1990; Truswell, 2019; van Hout, 2016; Vendler, 1957). For instance, the telic sentence "Sam entered the house" encodes an experience as a bounded event. By contrast, the atelic sentence "Sam approached the house" encodes an experience as an unbounded event. Several studies have confirmed that telicity conjures different cognitive perspectives on the temporal structure of events during language comprehension (Barner et al., 2008; Malaia et al., 2012; Strickland et al., 2015; Wagner & Carey, 2003; Wellwood et al., 2018). Events denoted by telic phrases can be counted and processed as discrete individuals while those denoted by atelic phrases generally cannot. For instance, to answer a question with a telic predicate such as "Who did more jumping?," people directly count how many jumps occurred but to answer a question with an atelic predicate such as "Who did more running?," people use measurements such as the distance or time duration of running as opposed to number (Barner et al., 2008; Wittenberg & Levy, 2017; on individuation across linguistic and visual stimuli, see also Malaia et al., 2012; Strickland et al., 2015; Wagner & Carey, 2003; Wehry et al., 2019; Wellwood et al., 2018).

A Parallel Between Objects and Events as Mental Individuals

As can be seen by the brief overview above, the object–substance distinction in the object domain has parallels to the bounded–unbounded distinction in the event domain. The idea that events can be likened to objects is far from new (e.g., Bach, 1986; Casati & Varzi, 2008; Davidson, 1970; De Freitas et al., 2014; Ji & Papafragou, 2022; Miller & Johnson-Laird, 1976/2014; Quine, 1985/1996; Wellwood et al., 2018). Nevertheless, the nature of this parallel and the conceptual organizational principles that support

it remain underexplored, despite the foundational nature of this topic for a number of disciplines within cognitive science.

Here we take the position that the crucial distinction between objects and bounded events on one hand and substances and unbounded events on the other is that objects and bounded events qualify as individuals and that substances and unbounded events are nonindividuals. Support for this proposal comes from a study by Papafragou and Ji (2023) showing that there is a strong homology between cognitive representations of events and objects. In that study, after brief training, viewers were able to extend categories of bounded or unbounded events to objects or substances, respectively. For example, they were able to extend a category of bounded events (e.g., dress a teddy bear) to also include novel objects (e.g., a solid, ring-like entity), and a category of unbounded events (e.g., pat a teddy bear) to also include novel substances (e.g., a white nonsolid mass). Importantly, viewers were able to draw such connections between events and objects even in the absence of prior training. Thus, the cognitive representations of bounded/unbounded events and objects/substances seem to be strongly aligned. This is summarized in Table 1.

A similar parallel has been discussed in logico-semantic analyses, according to which the count–mass distinction in the semantics of nominals has a counterpart in the notion of telicity in the semantics of verbal predicates (Bach, 1986; Jackendoff, 1991; Taylor, 1977; c.f. Champollion, 2015, 2017; Filip, 2012; Truswell, 2019; Wellwood et al., 2018). As mentioned already, the meaning contents of both count nouns and telic phrases can be counted as individuals and compared by their number but those of both mass nouns and atelic phrases are better measured than counted (Barner et al., 2008; Wittenberg & Levy, 2017).

If the strong alignment between construals of objects/substances and bounded/unbounded events is based on the notion of individuation, what makes a good individual? In an obvious sense, individuation in both domains is supported by the notion of boundaries: While both objects and substances are spatially extended, only objects are delimited by well-defined spatial boundaries. Similarly, while both bounded and unbounded events are temporally extended, only bounded events are delimited by well-defined temporal boundaries. In the object domain, boundaries, or contours, are central to object recognition and perception. For example, it has been proposed that cohesion is one of the most important principles of objecthood: An object must maintain a single bounded contour over time (e.g., Bloom, 1999; Pinker, 1997; Spelke, 1990, 1994). Similarly, in the event domain, it has been shown that identifying boundaries is a core component of perception and has consequences for memory and learning (Zacks & Swallow, 2007; see also: Ji & Papafragou, 2020b, 2022 for event boundedness). Bridging the two domains, a recent study found that nonsigners were systematically more likely to associate sign language signs that end with a gestural boundary with either telic verbs (e.g., *die*, *arrive*) or count nouns (e.g., *ball*, *coin*) compared to atelic verbs or mass nouns (Kuhn et al., 2021).

Table 1
Individuation in the Spatial and the Temporal Domain

Domain	Individuated entity	Unindividuated entity
Spatial domain	Object	Substance
Temporal domain	Bounded event	Unbounded event

Despite its important role, the notion of boundariness does not connect individuation to the content of entities across the domains of space and time. Any individual unit that humans perceive and experience has internal representational content: Objects are not empty shapes and events are not empty segments. We propose that understanding what happens within entity boundaries is also fundamental to human experience. In order to better understand the nature of individuals, or units of experience, we need to shift attention from boundaries to properties of individuals themselves across domains.

Relevant work has addressed the conceptual principles that underpin either “objecthood” or “event boundedness” but typically not both (e.g., Ji & Papafragou, 2020a, 2020b; Prasada et al., 2002; Rips & Hespos, 2015). In one study, cues that indicate nonarbitrariness of structure—regularity of structure, repetition of structure, and the existence of structure-dependent functions—led participants to describe novel entities using count syntax (“There is a blicket”) as opposed to mass syntax (“There is blicket”), and therefore pointed to objecthood (as opposed to substancehood; Prasada et al., 2002). In another study, “naturalness of breaks” was found to encourage individuation across domains: Observers tended to describe images with “natural” spatial breaks and motions with “natural” temporal breaks using plural count nouns and telic verb phrases respectively but chose mass nouns and atelic verb phrases for unnaturally divided images and motions (Wellwood et al., 2018). In the current study, we present and test a set of signatures of individuation across spatial and temporal entities using exclusively nonlinguistic tasks in order to better understand nonlinguistic conceptual principles that support individuation.

Signatures of Individuation Across Domains

We propose that individuated entities (objects and bounded events) possess a well-defined internal structure. On the other hand, unindividuated entities (substances and unbounded events) do not have a structure that is as well-defined (and could be, in some sense, “looser”). Structure refers to the way the parts are organized to make up the whole. Traditional semantics has captured the contrast between objects and telic predicates on one hand versus substances and atelic predicates on the other by mereological (part-whole structure) notions. In our proposal, the nonlinguistic conceptual structure of spatial and temporal entities can also be characterized in terms of their part-whole structure.

Objects have parts with particular spatial configurations. A table, for example, has parts like the tabletop and legs that are arranged in a certain way. Bounded events also have parts with a particular temporal organization.¹ An event of cracking an egg has a beginning, middle, and end. We hypothesize that the internal parts of structured entities are organized in a predefined manner. We therefore posit three principles that should apply to any structured entity across domains.

¹ In viewing the representational structure of individuals in terms of part-whole structure, one needs to consider what counts as a part. In linguistics, a similar problem is known as the “minimal-parts problem”: although *water* is a mass noun, water has minimal parts that are no longer water (Quine, 1960) and although *waltz* is an atelic predicate, waltzing has minimal parts that do not count as waltzing (Dowty, 1979). This issue is outside of the scope of this article and the experiments reported in this work do not run into this problem, but it is important to accurately characterize the granularity of parts.

The first principle is that structured entities resist restructuring. Objects are structured of spatial parts and these parts are organized in a designated manner. That is, a table leg cannot be placed above the tabletop and a shirt sleeve cannot be placed on the neck opening of the shirt. However, sand can be mixed up and still would count as sand and clay can be played around with and still would be clay. What does this mean for events? Bounded events are structured of temporal parts that are organized in a designated temporal manner. That is, the temporal order of the internal parts of bounded events cannot be rearranged—building a house proceeds in steps that cannot arbitrarily be rearranged. On the other hand, this is less of a problem for events like walking, where a third step and the fifth step can be switched and would not be a problem. This is summarized in (1).

1. No restructuring: Individuated entities resist restructuring.
 - a. Objects (but not substances) resist spatial restructuring.
 - b. Bounded (but not unbounded) events resist temporal restructuring.

Possessing internal structure also has consequences for the relationship between subparts of the entity. A bit of sand and another bit of sand are both sand, and therefore they are not understood to be distinct. Similarly, a part of walking and another part of walking are both walking, and therefore would not be understood to be distinct. On the other hand, we cannot say the same thing about two different parts of a table (e.g., the table leg and the tabletop) or two different parts of dressing a teddy bear (e.g., when putting on pants and when putting on a bowtie). This principle is summarized in (2).

2. Distinct parts: Individuated entities have distinct parts.
 - a. Objects (but not substances) have distinct spatial parts.
 - b. Bounded (but not unbounded) events have distinct temporal parts.

The last principle is that the internal organization of structured entities cannot be interrupted. Breaking up (or “dividing”) individuated and unindividuated entities has different consequences: Dividing up individuated entities is more problematic than dividing up unindividuated entities. This is summarized in (3).

3. No breaks: Individuated entities resist breaks.
 - a. Objects (but not substances) resist spatial breaks.
 - b. Bounded (but not unbounded) events resist temporal breaks.

In the current study, we test these three principles in a series of experiments using tasks that do not involve producing or comprehending linguistic descriptions of objects or events. We use everyday, familiar objects and events in order to probe how the mind applies such abstract considerations across the spatial and temporal domains.

Transparency and Openness

We report how we determined the sample size. In all studies, sample sizes were determined a priori, without intermittent data analyses. The studies were not preregistered. The data and analysis scripts are available from the Open Science Framework (<https://osf.io/bwkdv/>). Data were analyzed using R, Version 4.2.1 (R Core Team, 2022).

Experiment 1: No Restructuring

In this experiment, we test the no restructuring principle in the domain of objects (Experiment 1a) and events (Experiment 1b). If individuated entities have a better-defined structural representation than unindividuated entities and are thus subject to this principle, observers should find restructurings to individuated entities (objects, bounded events) more noticeable than those to unindividuated entities (substances, unbounded events). However, if this is not the case, observers should find restructurings to individuated and unindividuated entities to be equally noticeable.

Experiment 1a: Objects

Method

Participants. Data from 22 adult (11 female, 11 male, $M_{\text{age}} = 35.55$, age range = 25–55) native speakers of English residing in the United States were obtained via Prolific. Detailed information about the demographics of participants who participated in all experiments in this study (including race/ethnicity) are reported in the [online supplemental materials](#). A pilot with 10 participants indicated that a sample size of at least 21 was needed to detect an effect of condition with 90% power at significance level $\alpha = .01$ (pilot Cohen’s $d = 0.909$). In order to obtain a sample size of at least 21, we planned to recruit 25 participants on Prolific and exclude any participants who did not complete the experiment. In Experiment 1a, three participants did not complete the experiment, leaving us with 22 participants. Sample size and recruitment procedure in subsequent studies were matched with Experiment 1a. Participants were compensated at an \$8/hr rate for their participation.

Stimuli. We used 16 pairs of images, each depicting a familiar object (e.g., vase) and substance (e.g., clay; see [Table 2](#)). In 10 pairs, the object was the artifact made from the substance counterpart (e.g., vase–clay), and in two pairs, the object was a natural kind and the substance was an artifact made from the object counterpart (e.g., onion–chopped onion). In the remaining four pairs, both the object and the substance were artifacts (e.g., roll of toilet paper–toilet paper). We created spatially restructured versions of each entity by switching the positions of the second and third vertical strips of the image (see [Table 3](#)). Images were edited using the Adobe Photoshop 2022 software. All images were in 400×400 pixel dimensions.

Table 2
List of Image Stimuli in Experiment 1a

No.	Objects (original)	Substances (original)
1	Onigiri	Rice
2	Concrete block	Cement powder
3	Gold bracelet	Gold nuggets
4	T-shirt	Cotton wool
5	Ball of yarn	Yarn threads
6	Vase	Clay
7	Tomato	Chopped tomatoes
8	Key	Metal
9	Wooden bowl	Chopped wood
10	Onion	Chopped onions
11	Roll of toilet paper	Toilet paper
12	Meatball	Ground meat
13	Crystal swan	Crystal
14	Sandcastle	Sand
15	Pot	Terra cotta clay
16	Lotion (bottle)	Lotion (material)

Table 3
Sample Images in Experiment 1a

Entity type	Original	Restructured
Object		
Substance		

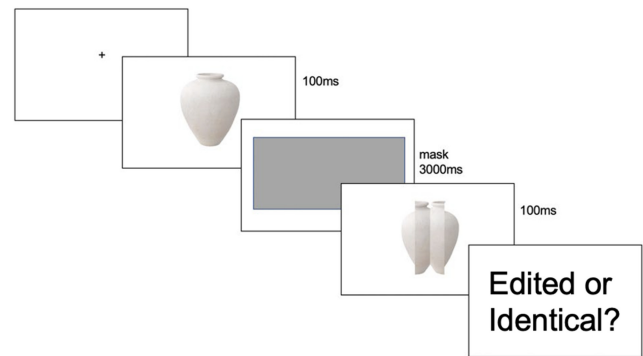
Note. See the online article for the color version of this table.

The original stimuli came from a pool of images that were normed in a manner similar to Li et al.'s (2009) Experiment 3, where participants were asked to rate the entities in their original (not restructured) form on a scale of 1–7, with 1 being a *good object* and 7 being a *good substance*. The stimuli were rated by 15 naïve native English speakers who did not participate in any of the other experiments reported in this study. Items categorized as objects had a mean rating of 2.62 ($SD = 2.25$), and items categorized as substances had a mean rating of 4.81 ($SD = 2.31$), with people reasonably rating substances higher than objects on our response scale, $t(14) = -7.1$, $p < .001$. Following the rating scales in Li et al. (2009), we additionally tested the stimuli on several features that have been associated with objecthood or lack thereof (see Ontological Distinctions in the Spatial Domain section): the complexity of their overall shape and outline (1 = *not at all complex* to 7 = *extremely complex*), the degree to which their function depended on their overall shape and outline (1 = *not at all dependent* to 7 = *extremely dependent*), as well as their cohesiveness/solidity (1 = *not at all cohesive/solid* to 7 = *extremely cohesive/solid*). Complexity ratings were similar across items categorized as objects ($M = 4.12$, $SD = 1.95$) and substances ($M = 3.34$, $SD = 1.90$), $t(14) = 1.8$, $p = .1$. Object stimuli ($M = 5.22$, $SD = 1.84$) were rated higher than substance stimuli ($M = 3.68$, $SD = 2.02$) in terms of shape-dependent function, $t(14) = 3.6$, $p < .01$. Similarly, object stimuli ($M = 5.04$, $SD = 1.86$) were rated higher than substance stimuli ($M = 3.81$, $SD = 2.03$) in terms of cohesiveness, $t(14) = 5.5$, $p < .001$. These results are consistent with Li et al.'s findings and confirm our choice of spatial entities.

In all experiments reported in this study, two lists were created and presented to participants using a standard Latin Square design: Each participant was presented with only one condition of each target item and each of the two conditions appeared the same number of times on any of the two lists. Both lists contained the same set of 12 filler items, pseudorandomly distributed throughout the list. Six of the filler items were identical image pairs and the other six filler items were image pairs that differed in color only, where the color difference was very subtle.

Procedure. All experiments reported in this study were hosted online on PennController IBEX (Zehr & Schwarz, 2018; <https://www.pciibex.net/>), and participants completed them remotely via the internet. At the beginning of each trial, a fixation cross was displayed for 1,000 ms. After the fixation cross, the original image was displayed for 100 ms. Then, the screen was masked for 3,000 ms.

Figure 1
Trial Structure in Experiment 1a



Note. + = fixation cross. See the online article for the color version of this figure.

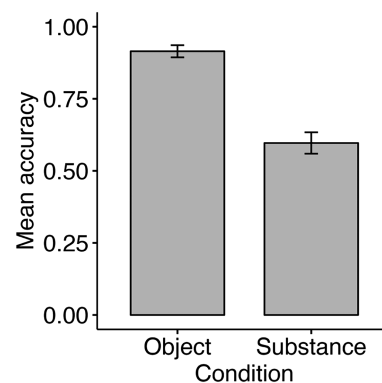
Afterwards, the restructured image was displayed for 100 ms. There was no postmask after the restructured image. At the end of each trial, participants were asked to identify whether the two items they saw were identical (“edited or identical?”; see Figure 1). The experiment took approximately 4 min to complete.

Results and Discussion

Results from Experiment 1a are shown in Figure 2. The accuracy of the participants' responses was analyzed using generalized linear mixed effects models (*glmer*). We coded condition (object vs. substance) using centered contrasts (object = 0.5, substance = -0.5) and included it as the fixed effect. As random effects, we entered intercepts for subjects and items, as well as by-subject and by-item random slopes for the effects of condition. They were then reduced (starting with by-item effects) via model comparison, wherein only random effects that contributed significantly to the model ($p < .05$) were included (Baayen et al., 2008). The reported model included by-subject random slopes and intercepts, and by-item intercepts.

There was a main effect of condition ($\beta = 2.45$, $SE = 0.55$, $z = 4.44$, $p < .001$), such that participants were more likely to accurately judge that

Figure 2
Mean Accuracy by Condition in Experiment 1a



Note. Error bars represent $\pm SE$.

the original and the restructured images were different when presented with objects ($M = 91.5\%$, $SD = 0.28$) than with substances ($M = 59.7\%$, $SD = 0.49$). (For filler items that were identical pairs, 29.5% [$SD = 0.46$] of the participants responded that they were different. The average accuracy for filler items that differed slightly in color was 65.9% [$SD = 0.48$].) We take these results as support for the no restructuring principle. Observers are better at detecting structural changes to objects because objects possess a better-defined internal structure and thus are more likely to resist restructuring. However, structural changes to substances are less noticeable because substances do not possess as well-defined an internal structure and are less likely to resist restructuring.

Experiment 1b: Events

Method

Participants. Twenty-three (16 female, seven male, $M_{\text{age}} = 38.05$, age range = 21–59) new adult native speakers of English residing in the United States participated via Prolific. Participants were compensated at an \$8/hr rate for their participation.

Stimuli. For target items, we used 16 pairs of videos from Ji and Papafragou (2020a). All videos involved the same girl doing a familiar everyday action in a lab room. Paired videos had the same duration and showed a bounded and an unbounded event (see Table 4). For the 16 pairs used as target stimuli in Experiment 1b, we created temporally restructured versions of each event by dividing each video into four temporal segments of equal duration and switching the second and third segments. This mirrors the stimuli design in Experiment 1a, where each

Table 4

List of Video Stimuli in Experiment 1b

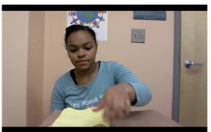
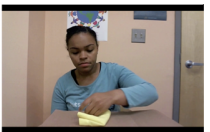
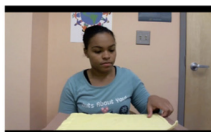
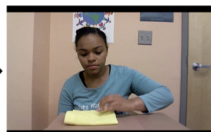
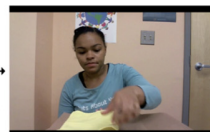
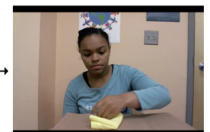
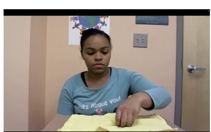
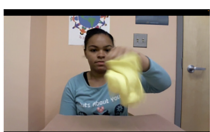
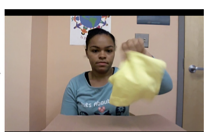
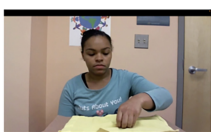
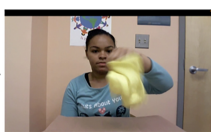
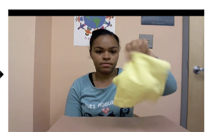
No.	Bounded events (original)	Unbounded events (original)
1	Scoop up yogurt	Stir yogurt
2	Fold a handkerchief	Wave a handkerchief
3	Pile up a deck of cards	Shuffle a deck of cards
4	Group pawns based on color	Mix pawns of two colors
5	Roll up a towel	Twist a towel
6	Dress a teddy bear	Pat a teddy bear
7	Put up one's hair	Scratch one's hair
8	Close a fan	Use a fan
9	Eat a pretzel	Eat cheerios
10	Peel a banana	Crack peanuts
11	Stick a sticker	Stick stickers
12	Tie a knot	Tie knots
13	Tear up a tissue	Tear slices off tissues
14	Cut a ribbon in half	Cut pieces from a roll
15	Flip a postcard	Flip pages
16	Blow a balloon	Blow bubbles

entity was divided into four segments of equal widths and the second and third segments were switched. Videos were edited using the Adobe Premiere Pro 2022 software (for an example, see Table 5).

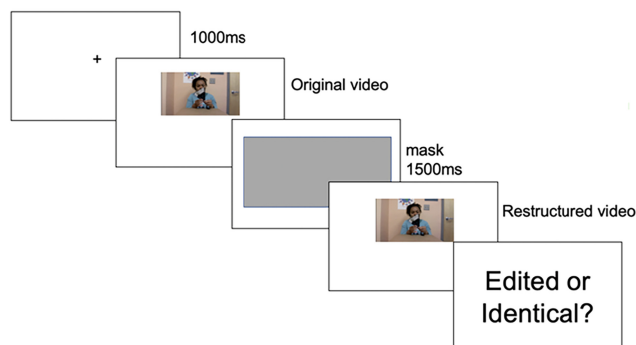
The original versions of these videos were drawn from a pool of 20 pairs of bounded–unbounded videos in the earlier study (duration range = 4.5–13 s, $M = 7.98$ s). That set had been normed to ensure that all video stimuli would illustrate the contrast in boundedness (Ji & Papafragou, 2022): Participants ($n = 40$) judged videos of bounded events as “something with a beginning, midpoint and specific endpoint” 87% of the time but said the same for videos of

Table 5

Sample Videos in Experiment 1b

Event type	Video type	Stimuli			
Bounded event (fold a handkerchief)	Original				
	Restructured				
Unbounded event (wave a handkerchief)	Original				
	Restructured				

Note. See the online article for the color version of this table.

Figure 3*Trial Structure in Experiment 1b*

Note. + = fixation cross. See the online article for the color version of this figure.

unbounded events only 21.5% of the time, a significant difference, $t(39) = 20.33, p < .001$. Additional norming confirmed that the videos within each class were equivalent in other dimensions, including intentionality (Ji & Papafragou, 2020a): On a scale from 1 = *totally unintentional* to 7 = *intentional*, participants ($n = 20$) rated bounded ($M = 5.67$) and unbounded events ($M = 5.62$) similarly, $t(19) = 1.34, p = .195$. It is of note that norming for event stimuli is parallel in some respects to the norming previously reported for object stimuli in Experiment 1a. For instance, the question about event timepoints is related to the cohesiveness/solidity question in Experiment 1a: Possessing a “beginning, midpoint and specific endpoint” in the event domain is similar to being a cohesive entity (which has each of the parts in a specific place) in the object domain. Similarly, the intentionality question about events is related in some sense to the shape-dependent function question about objects, since intentionality in the event domain can be likened to function in the object domain—they both make reference to what needs to be achieved in the real world.

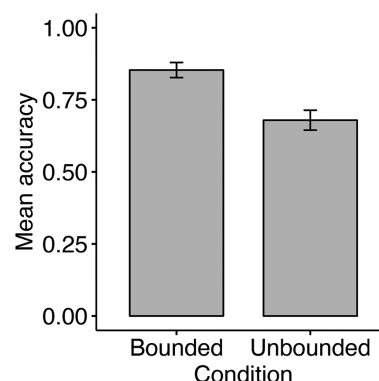
In addition to the target items, we included twelve filler items. Six of the filler items were identical video pairs, and the other six of the filler items were video pairs where the edited version involved temporally reversing segments of the video.

Procedure. At the beginning of each trial, a fixation cross was displayed for 1,000 ms. Once the fixation cross disappeared, participants watched the original video. Then, the screen was masked for 1,500 ms. Afterwards, participants watched the restructured video. At the end of each trial, participants were asked to identify whether the two videos they watched were identical (“edited or identical?”; see Figure 3). The experiment took approximately 13 min to complete.

Results and Discussion

Results from Experiment 1b are shown in Figure 4. The accuracy of the participants’ responses was analyzed in the same way as in Experiment 1a, with condition (bounded vs. unbounded; contrast-coded as bounded = 0.5 and unbounded = −0.5) as a fixed effect. The reported model included by-subject and by-item random slopes and intercepts.

There was an effect of condition ($\beta = 1.38, SE = 0.72, z = 1.92, p = .05$), such that participants were more likely to accurately judge the original video and the restructured video as different in the

Figure 4*Mean Accuracy by Condition in Experiment 1b*

Note. Error bars represent $\pm SE$.

bounded event condition ($M = 85.33\%$, $SD = 0.36$) than in the unbounded event condition ($M = 67.93\%$, $SD = 0.47$). (For filler items that were identical pairs of videos, 78.3% [$SD = 0.41$] of the participants responded that they were different. For filler items that were pairs of edited videos, 93.5% [$SD = 0.25$] of the participants accurately responded that they were different.) As expected, then, participants had more difficulty detecting structural changes to unbounded events than to bounded events. We take these results to again provide support for the no restructuring principle, now in the event domain. Observers are better at detecting structural changes to bounded events because bounded events possess a better-defined internal structure and thus are more likely to resist restructuring.

Experiment 2: Distinct Parts

In Experiment 2, we test the principle of distinct parts in the domain of objects (Experiment 2a) and events (Experiment 2b). If this principle applies to individuated entities, observers should be better at detecting the difference between two random subparts of individuated entities (objects, bounded events) than the difference between two random subparts of unindividuated entities (substances, unbounded events). However, if there is no difference in the structural representation of individuated and unindividuated entities, observers’ ability to detect differences between subparts of individuated and unindividuated entities would not differ.

Experiment 2a: Objects

Method

Participants. Twenty-four (15 female, nine male, $M_{age} = 40.33$, age range = 25–70) new adult native speakers of English residing in the United States participated via Prolific. Participants were compensated at an \$8/hr rate for their participation.

Stimuli. Using the 16 original images used in Experiment 1, we extracted two different 80×80 pixel parts from each image, using the Figma editor. One part was extracted from the center of each entity (middle part), and another part from the top right corner of each entity (edge part). See Table 6 for examples. We selected the center and edge parts of each entity to make the parts maximally

Table 6
Sample Image Stimuli in Experiment 2a

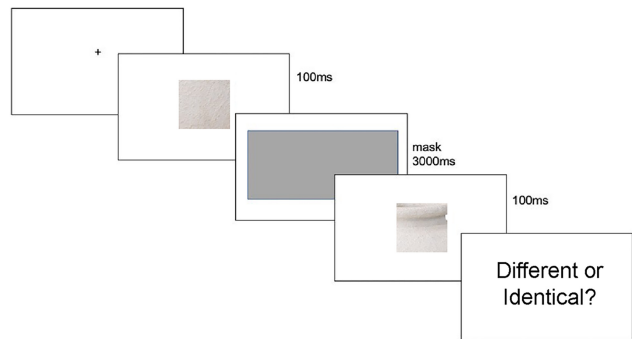
Entity type	Middle part	Edge part
Object		
Substance		

Note. From Adobe Stock, 2024 (<https://stock.adobe.com/in>). In the public domain. See the online article for the color version of this table.

distinct from each other, for both objects and substances. Moreover, our selection of center and edge parts of spatial entities mirrors our selection of middle and end segments of temporal entities in Experiment 2b. In addition to target stimuli, we included twelve filler items. Six of the filler items were identical image pairs. The other six were middle and edge parts of filler images such as a car or a bowl of colorful marbles.

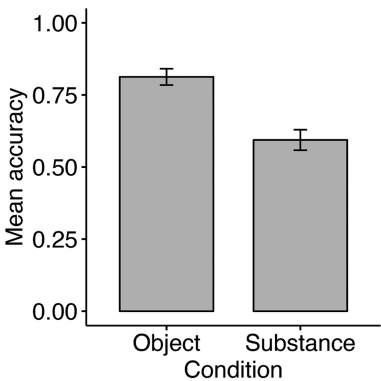
Procedure. The trial structure of Experiment 2a was similar to that of Experiment 1a. At the beginning of each trial, a fixation cross was displayed for 1,000 ms. After the fixation cross, one part of an entity was displayed for 100 ms. Then, the screen was masked for 3,000 ms. Afterwards, the other part of the entity was displayed for 100 ms. There was no postmask after the other part of the entity. At the end of each trial, participants were simply asked, “Different or Identical?” (see Figure 5). The ordering of the segments was counterbalanced so that in one half of the trials, participants saw the middle segment first, and in the other half, they saw the edge segment first. The experiment took approximately 4 min to complete.

Figure 5
Trial Structure in Experiment 2a



Note. + = fixation cross. See the online article for the color version of this figure.

Figure 6
Mean Accuracy by Condition in Experiment 2a



Note. Error bars represent $\pm SE$.

Results and Discussion

Results from Experiment 2a are shown in Figure 6. Participants’ responses were analyzed in the same way as in Experiment 1a. The reported model included by-subject and by-item random slopes and intercepts.

There was a main effect of condition ($\beta = 1.89$, $SE = 0.57$, $z = 3.29$, $p < .001$), such that participants were more likely to accurately identify the two parts as distinct in the object condition ($M = 81.25\%$, $SD = 0.39$) than in the substance condition ($M = 59.38\%$, $SD = 0.49$). (For filler items that were identical pairs of images, 79.2% [$SD = 0.41$] of the participants responded that they were different. For filler items that were distinct parts of an entity, 84.7% [$SD = 0.36$] of the participants accurately identified them as different.) As the principle of distinct parts predicted, objects are more likely to have distinguishable parts than substances.

Experiment 2b: Events

Method

Participants. Twenty-one (11 female, 10 male, $M_{age} = 39.24$, age range = 20–79) new adult native speakers of English residing in the United States participated via Prolific. Participants were compensated at an \$8/hr rate for their participation.

Stimuli. We segmented each original video from Experiment 1b into nine temporal segments, and used the fifth (middle) and the eighth segments (see Table 7). This mirrors the stimuli design in Experiment 2a, where the middle and edge parts of spatial entities were selected.

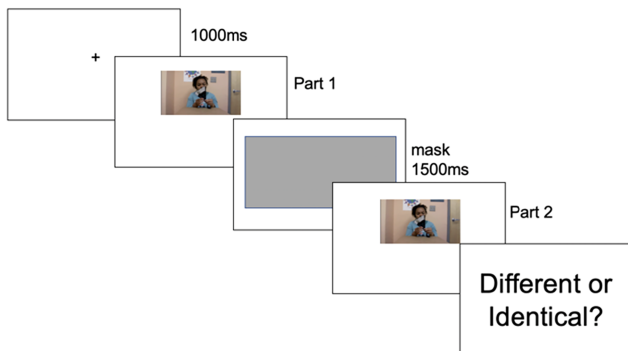
Procedure. The trial structure of Experiment 2b was similar to that of Experiment 1b. At the beginning of each trial, a fixation cross was displayed for 1,000 ms. Once the fixation cross disappeared, participants watched a video segment. Then, the screen was masked for 1,500 ms. Afterwards, participants watched the other video segment. At the end of each trial, participants were simply asked, “different or identical?” (see Figure 7). The ordering of the segments was counterbalanced so that in half of the trials, participants saw the middle segment first, and in the other half, they saw the end segment first. The experiment took approximately 9 min to complete.

Table 7
Sample Video Stimuli in Experiment 2b

Event type	Middle segment	End segment
Bounded event		
Unbounded event		

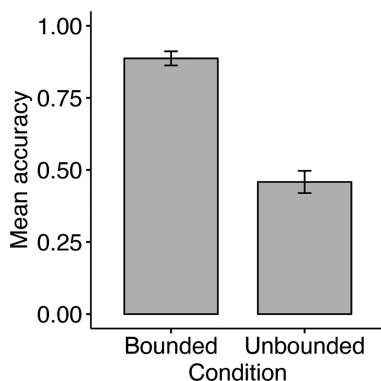
Note. See the online article for the color version of this table.

Figure 7
Trial Structure in Experiment 2b



Note. + = fixation cross. See the online article for the color version of this figure.

Figure 8
Mean Accuracy by Condition in Experiment 2b



Note. Error bars represent $\pm SE$.

Results and Discussion

Results from Experiment 2b are shown in Figure 8. Participants' responses were analyzed in the same way as in Experiment 1b. The reported model included by-subject random slopes and intercepts, and by-item intercepts.

There was a main effect of condition ($\beta = 3.70$, $SE = 0.85$, $z = 4.36$, $p < .001$), such that participants were more likely to accurately identify the two segments as distinct for bounded ($M = 88.69\%$, $SD = 0.32$) than for unbounded events ($M = 45.83\%$, $SD = 0.50$). (For filler items that were identical pairs of videos, 78.6% [$SD = 0.41$] of the participants responded that they were different. For filler items that were different pairs of videos, 94.4% [$SD = 0.23$] of the participants accurately responded that they were different.) Again, as the principle of distinct parts would predict, bounded events are more likely than unbounded events to have distinguishable parts.

Experiment 3: No Breaks

In the present experiment, we test the no breaks principle in the domain of objects (Experiment 3a) and events (Experiment 3b). If individuated entities have a better-defined internal structural representation and are thus subject to this principle, observers should find breaks to individuated entities (objects, bounded events) more problematic than those to unindividuated entities (substances, unbounded events). However, if there is no difference in the structural representation of individuated and unindividuated entities, observers would find breaks to individuated and unindividuated entities equally problematic. Notice that Experiment 3 asks participants directly to assess the gravity of structural disruptions to real-world entities corresponding to the contents of images or videos (as opposed to assessing the images and videos themselves, and thus by inference their contents, as in Experiments 1 and 2).

Experiment 3a: Objects

Method

Participants. Twenty-one (six female, 15 male, $M_{\text{age}} = 39.62$, age range = 20–67) new adult native speakers of English residing in the United States participated via Prolific. Participants were compensated at an \$8/hr rate for their participation.

Stimuli. Using the 16 pairs of object and substance images from Experiment 1a, we created “disrupted” (broken/discontinuous) versions of each entity by inserting blank (white) vertical gaps in between four vertical strips of equal widths. As in Experiment 1a, we used four segments (see Table 8). In addition to the target items, we also included twelve filler items. In six filler items, the entities shown before and after the monster’s break-in were identical. In the other six items, the entities shown before and after the monster’s break-in differed in color.

Procedure. Participants were provided with the following story:

A monster broke into Alessandra’s house and messed up her possessions! We will first show you what the item originally looked like. Next, we will show you what it looks like after the monster broke in. Your task is to help Alessandra assess the damage done to her possessions! For each item, how much of a problem is what the monster did?

Table 8

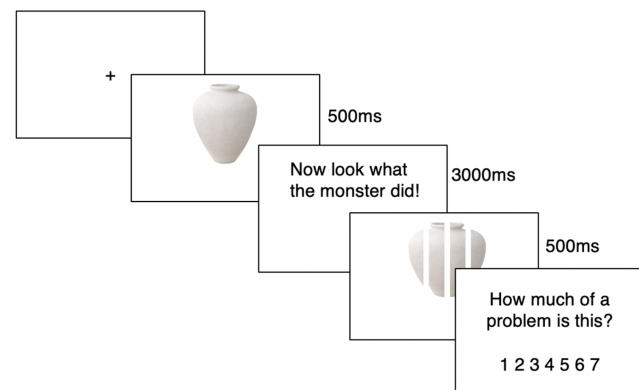
Sample “Disrupted” Entities in Experiment 3a

Object	Substance
	

Note. From Adobe Stock, 2024 (<https://stock.adobe.com/in>). In the public domain. See the online article for the color version of this table.

Figure 9

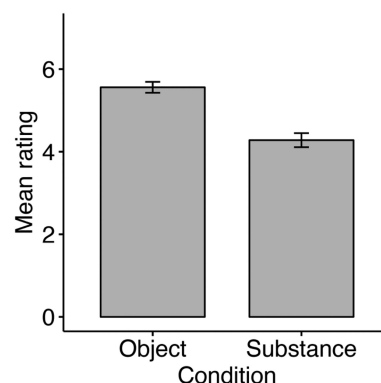
Trial Structure in Experiment 3a



Note. From Adobe Stock, 2024 (<https://stock.adobe.com/in>). In the public domain. + = fixation cross. See the online article for the color version of this figure.

Figure 10

Mean Rating by Condition in Experiment 3a



Note. Error bars represent $\pm SE$.

(Here, participants saw a picture of a monster, but not of Alessandra.) At the beginning of each trial, a fixation cross was displayed for 1,000 ms. After the fixation cross, the original image was displayed for 500 ms. Then, for 3,000 ms, participants saw a message saying, “Now look what the monster did!” Afterwards, the disrupted version of the entity was displayed for 500 ms. Then, participants were asked, “How much of a problem is this?”, and were given a 7-point scale, where 1 indicated *no problem at all* and 7 indicated *a serious problem* (Figure 9). The experiment took approximately 4 min to complete.

Results and Discussion

Results from Experiment 3a are shown in Figure 10. Participants’ responses were analyzed using linear mixed effects models (*lmer*). We coded condition (object vs. substance) using centered contrasts (0.5, −0.5) and included it as the fixed effect. As random effects, we entered intercepts for subjects and items, as well as by-subject and by-item random slopes for the effects of condition. Model comparison and selection were conducted in the same way as in Experiments 1 and 2. The reported model included by-subject and by-item random slopes and intercepts.

There was a main effect of condition ($\beta = 1.26$, $SE = 0.37$, $t = 3.38$, $p < .01$), such that participants judged structural breaks to objects as more problematic ($M = 5.56$, $SD = 1.71$) than structural breaks to substances ($M = 4.28$, $SD = 2.20$). (The mean response to filler items that were identical pairs was 1.58 [$SD = 1.34$], and the mean response to filler items that differed in color was 2.57 [$SD = 1.89$].) As expected by the no breaks principle, structural breaks do, in fact, affect how one evaluates real-world spatial entities (and not just images that depict these entities).

Experiment 3b: Events

Method

Participants. Twenty-five (12 female, 13 male, $M_{\text{age}} = 36.2$, age range = 21–75) new adult native speakers of English residing in the United States participated via Prolific. Participants were compensated at an \$8/hr rate for their participation.

Table 9
Sample Video Stimuli in Experiment 3b

Event type	Stimuli					
Bounded event						
	Monster interrupts girl.		Monster interrupts girl.		Monster interrupts girl.	
Unbounded event						
	Monster interrupts girl.		Monster interrupts girl.		Monster interrupts girl.	

Note. See the online article for the color version of this table.

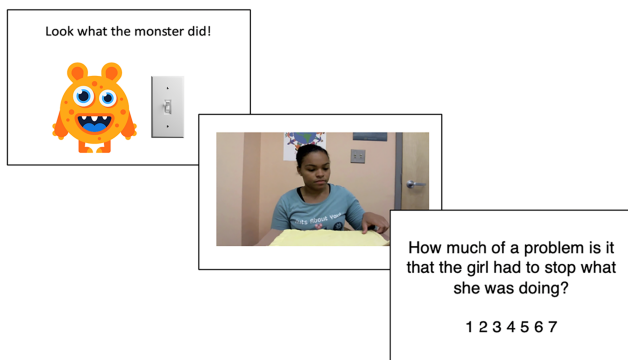
Stimuli. Using the original 16 videos from Experiment 1b, we created videos in which breaks were inserted between each quarter of the video. As in Experiment 1b, we used four segments. Between each of the four segments, we inserted breaks (three breaks in total) that interrupted the action itself (as opposed to editorial breaks in the video only). During each break, the screen turned entirely black and participants heard sounds of the light switch turning off, the girl walking over to the monster and telling them off (“What are you doing!,” “Stop!,” “Stop it!”), the monster screeching back at the girl, the girl coming back to the chair, and the light switch turning back on (Table 9).

Procedure. Participants were provided with the following context:

A monster breaks into Alessandra’s room and messes up what she’s doing by switching the lights on and off! Alessandra can’t see what she’s doing in the dark so she stops her actions when the lights are off. Your task is to determine how much of a problem is what the monster did to Alessandra’s actions.

On each trial, participants saw a screen showing a monster and a light switch, along with the message, “Look what the monster did! He switched the lights on and off and messed up Alessandra’s actions!”. They then watched the video with breaks.² Afterwards, they were asked, “How much of a problem is it that the girl had to stop what she was doing?”

Figure 11
Trial Structure in Experiment 3b



Note. From Adobe Stock, 2024 (<https://stock.adobe.com/in>). In the public domain. See the online article for the color version of this figure.

what she was doing?”, and were asked to respond on a 7-point scale, where 1 = *no problem at all* and 7 = *a serious problem* (Figure 11). The experiment took approximately 16 min to complete.

Results and Discussion

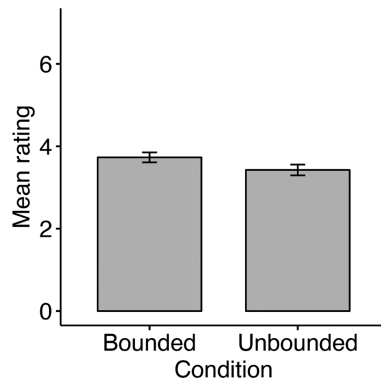
Results from Experiment 3b are shown in Figure 12. Participants’ responses were analyzed in the same way as in Experiment 3a, with condition (bounded vs. unbounded; contrast-coded as 0.5 and –0.5) as a fixed effect. The reported model included by-subject and by-item random intercepts.

There was a main effect of condition ($\beta = 0.31$, $SE = 0.12$, $t = 2.45$, $p = .01$), such that participants judged temporal breaks to bounded events as more problematic ($M = 3.73$, $SD = 1.71$) than temporal breaks to unbounded events ($M = 3.43$, $SD = 1.86$). These results are consistent with the no breaks principle in the domain of events. Again, these results support the claim that structural principles affect how one evaluates real-world events (and not just videos that depict events.)

We do, however, note that the numerical difference is not as great as in Experiment 3a. This may be due to the fact that participants saw that the girl was still able to go through and finish the action even in the bounded event condition, and therefore found such disruptions less problematic than disruptions to spatial entities that would be less likely to function in the same way once they were “broken.” Thus, temporal breaks to events may be understood as less problematic than spatial breaks to objects. Indeed, our everyday experience always consists of events and actions that get interrupted. Our naps get interrupted by construction noise and our meetings get interrupted by phone calls. We switch back and forth between different events like baking a cake and doing laundry. Despite such interruptions, we still manage to understand and track events as individual units. It would be interesting to ask whether a stronger

² Unlike Experiment 3a, participants only watched the video with breaks (and not the “original” video) because the provided contexts were different. In Experiment 3a, a monster caused damages to spatial entities in a way that it was possible to show the “before/intact” and “after/damaged” versions of the entity. In Experiment 3b, however, the temporal nature of events did not allow for this type of design: the monster interrupted the action in a way that there were no “before/intact” and “after/interrupted” versions of the events. Notice also that, unlike Experiment 2a, the form of the question is more specific so as to clarify that participants were being asked to evaluate how problematic the presence of the breaks was for the actual action.

Figure 12
Mean Rating by Condition in
Experiment 3b



Note. Error bars represent $\pm SE$.

asymmetry would have been observed if the interrupted event had not terminated, that is, if the disruption had more lasting effects in ways comparable to the spatial disruptions in Experiment 3a.

Summarizing all of our findings, viewers process perceptual changes (restructuring, part extraction, breaks) across both images and videos in systematically different ways depending on whether the stimuli are taken to depict an individual or not.

General Discussion

What are the basic units of human experience across space and time? And are there connections among these mental units across the two domains? This is a foundational topic for cognitive science, with clear implications for the way humans organize their representation of objects and events and encode these mental particulars in language. In this paper, we investigated the units of human cognition across the spatial and the temporal domain, with an eye towards cross-domain similarities.

We began with the assumption that ontological distinctions within each domain align in terms of an underlying notion of individuation (Papafragou & Ji, 2023; see also Bach, 1986; Davidson, 1970; Ji & Papafragou, 2022; Kuhn et al., 2021; Quine, 1985/1996; Wagner & Carey, 2003; Wellwood et al., 2018). Specifically, within the domain of spatial entities, objects qualify as individuals while substances do not; similarly, in the event domain, bounded events count as individuals while unbounded events do not. We proposed that the common abstract signature of individuation that crosscuts the two domains is the presence of a well-defined internal structure—that is, a principled, nonarbitrary way that parts are organized to make up a (spatial or temporal) whole. In a series of experiments, we tested and confirmed three principles that follow from this proposal. First, individuated entities across the object and event domains obey the no restructuring principle: Objects and bounded events are more likely to resist restructuring of their parts but substances and unbounded events are less likely to do so. Second, individuated entities submit to the distinct parts principle: Objects and bounded events are more likely to be understood as possessing distinct parts but substances and unbounded events are less likely to. Third, individuated entities follow the no breaks principle: Objects and bounded events are more

likely to resist spatial and temporal breaks, respectively, but substances and unbounded events are less likely to. Together, these data support the idea that there is a conceptual difference between mental individuals and nonindividuals in both space and time (see Papafragou & Ji, 2023), and offer the first evidence in support of common nonlinguistic signatures of conceptual individuation across the domains of objects and events.

Implications for Theories of Object and Event Cognition

In the literature, it is widely assumed that individuation across the spatial and temporal domains is supported by the presence of boundaries. Within the spatial domain, while both objects and substances are spatially extended, only objects are clearly spatially delimited (Bloom, 1999; Kuhn et al., 2021; Pinker, 1997; Spelke, 1990, 1994). Similarly, within the temporal domain, while both bounded and unbounded events are temporally extended, only bounded events are clearly temporally delimited (Ji & Papafragou, 2020b, 2022; Kuhn et al., 2021; Strickland et al., 2015; Zacks & Swallow, 2007). Despite the uncontroversial role of boundaries in the processing of objects and events, here we have argued that the deeper internal structure of entities, rather than a mere notion of boundaries or contours, is fundamental for characterizing conceptual individuation across objects and events. For example, results from Experiment 2 (distinct parts), where external boundaries were not at all relevant, would not be predicted by an account that solely focuses on boundaries. Our proposal is consistent with the possibility that an entity's boundaries are a secondary property that follows from how the entity is internally structured. Thus, even though we have been using the terms bounded and unbounded events, the relevant categorization does not just concern event boundaries themselves but the internal structure (or the lack thereof) of temporal entities. Our perspective requires shifting the theoretical emphasis away from a sole focus on segmentation or boundariness and calls for theories of object and event cognition that properly account for the internal architecture of individual units.³

In the spatial domain, the idea that objects have structure and distinct parts is not new (Biederman, 1987; Blum, 1973; Hoffman & Richards, 1984; Hummel & Biederman, 1992; Marr & Nishihara, 1978). For example, in Biederman's (1987) theory of human visual object recognition, objects are recognized as an arrangement of simple geometric components, such as blocks, cylinders, wedges, and cones. However, this view has mainly focused on how objects are recognized in the visual world and has not generalized the idea of partonomic structure to the domain-independent notion of a

³ One way of conceptualizing viewers' commitments to the internal structure of units across the object and event domains is reminiscent of Aristotle's analysis of form and matter that we have already alluded to (Aristotle, 350 B.C.E./1994). Objects have both matter and form, while substances themselves are matter yet lack an inherent form. For example, objects such as wooden chairs have both matter (wood) and form (chair form); substances such as wood are themselves matter and lack an inherent form. In a similar way, we have proposed that unbounded events such as walking are the "event matter" themselves yet lack an inherent event form (temporal structure and resulting boundaries). Bounded events such as a walk to the store, on the other hand, have both "event matter" and "event form." We stress that, on our account, substances and unbounded events have independent status as entities and are not simply properties of other (individuated) entities.

cognitive individual. Similarly, in the temporal domain, the idea that events can have parts has been around for some time but has been used so far in limited ways. According to the best-known version of this idea, finer grained events make up coarser grained, larger events (Zacks & Tversky, 2001); for example, an event of making coffee has subparts such as grinding coffee beans, picking up the coffee cup, and pouring coffee into the cup. Our study suggests that it is also crucial to study how smaller scale events are internally structured: Picking up the coffee cup has a well-defined internal structure to it—including a beginning, middle, and end—but not all events have such structure. A unique feature of our account is the suggestion that a proper theory of events should be able to account for the part-whole structure of events in a manner similar to the part-whole structure of objects. By contrast, on the above theories of objects and events, there is no expectation that objects and events should behave identically, and therefore no way of explaining the line of results in the present paper without additional theoretical machinery.⁴

The present perspective highlights the notion of structure as the crucial conceptual signature of individuals. What, then, does it specifically mean for an entity to possess a well-defined structure? Put differently: What cues lead us to think of an entity as having a well-defined structure? Recall that, in an earlier study, regularity of structure, repetition of structure, and the existence of structure-dependent functions led participants to describe novel entities using count syntax (“There is an X”) as opposed to mass syntax (“There is X”), and therefore pointed to objecthood as opposed to substancehood (Prasada et al., 2002). For instance, entities that had “regular” shapes (straight edges and curves with constant or smoothly changing curvature) were more likely to be described with count nouns. In this sense, Prasada et al.’s (2002) notion of “regularity of structure” refers to “regularity of boundaries/contours” as opposed to an entity’s internal structure. Prasada et al. (2002) also showed that possessing (nonarbitrary) structure implicates existence of a structure-generating process and a structure-dependent function. While the function proposal makes sense for artifacts like tables or keys, it is less straightforward for natural kinds such as apples and giraffes. It is also unclear whether and how such cues would generalize to the event domain.

Relatedly, naturalness of spatial or temporal breaks has been found to encourage linguistic individuation across the static and dynamic domains (Wellwood et al., 2018). Even though our stimuli did not probe the same idea, it is clear that entities that have a well-defined internal structure are likely to have more natural divisions. For example, a wooden ladder comes with very regular, natural divisions and patterns between each wooden step, whereas a pile of firewood is arranged in a random, arbitrary way such that the “divisions” between each wood piece are not regular. Similarly, in the domain of events, closing a fan comes with clear steps that have a beginning, middle, and end, whereas using a fan can involve more random movements that can be arranged in an arbitrary way.

Summarizing, the current study asked: Once an entity is construed as having a well-defined structure, what predictions follow? We found that the presence of well-defined structure has real cognitive consequences for entity perception and processing. Naturally, the cues that signal individuation and the structural consequences of individuation are two sides of the same coin. As discussed above, entities that follow our three principles are

likely to have regular structure, a significant function, a designated shape/boundary, and natural breaks, among other things, although our tasks do not directly probe these properties. Our study was the first to reveal that a fundamental set of structural principles underlies object and event construals in a series of non-linguistic tasks.

Implications for the Language–Cognition Interface

Understanding object and event structure in terms of individuation allows for meaningful links to linguistic semantic theories. As mentioned already, such theories have long noted the contrast between objects and telic verb phrases on one hand versus substances and atelic verb phrases on the other (Bach, 1986; Jackendoff, 1991; Taylor, 1977; cf. Barner et al., 2008; Champollion, 2015, 2017; Filip, 2012; Kuhn et al., 2021; Truswell, 2019; Wellwood et al., 2018; Wittenberg & Levy, 2017). Our study shows that the nonlinguistic structure of spatial and temporal entities can be characterized in terms of similar individuation profiles. This in turn supports the conclusion that an analysis of natural language can reveal meaning distinctions that characterize conceptual systems beyond language; furthermore, it suggests that the explanatory scope of linguistic theory should adjust accordingly so that it is not called upon to explain phenomena that could be in part explained by broader cognitive architecture.

Specifically, our proposal that underscores objects’ and events’ part-whole structure obviously converges in part with linguistic analyses of the semantics of nouns and predicates in terms of their mereological (part-whole) structure. In semantic theories, mass nouns and atelic predicates are theorized to have cumulative reference: If two things are wood, then their sum is also wood, and if two things are walking, then their sum is also walking (e.g., Bach, 1986; Link, 1983). They also have divisive reference: Any part of anything that is wood is also wood, and the same applies to walking. We can understand cumulativity as looking upward from part to whole, and divisiveness as looking downward from whole to part. What semantic theories do not explicitly discuss is looking sideways. That is, linguistic representations do not specify the way the parts are arranged to form a whole. However, our results—in particular results from Experiment 1—strongly suggest that the manner in which the parts are arranged to form a whole is crucial in humans’ conceptual representations of individuated entities. This suggests that the cognitive system is attuned to structural aspects of physical entities that language and linguistic theory do not make reference to. This is not unreasonable, given that language does not name parts of individuals in the same way that the cognitive system can perceive and experience them.

For purposes of the current study, we focused on structure-related principles that apply once an entity is construed as an

⁴ The current proposal nicely links to other properties of objecthood that have been discussed in the literature. For instance, cohesion (and similarly, solidity) can be understood as a feature that follows from a certain type of relationship that holds among the parts of an object. Because the parts of objects, but not of substances, are spatially arranged in a designated way, they need to cohere in a certain way. If the parts do not cohere with each other and are thus susceptible to restructuring, the whole entity would violate the no restructuring principle. Indeed, the object stimuli used in our experiment were rated higher than substances in terms of cohesion (see norming data of Experiment 1; cf. also Li et al., 2009).

individual, and indeed, we intentionally used image and video stimuli that are likely to be construed as individuals or not (see the norming reported in Experiment 1). However, both spatial and temporal stimuli are often multiply interpretable. The same static entity in the physical world can either be construed as an object (e.g., a meatball) or a substance (e.g., meat). Similarly, one may construe the same perceptual happening as an individuated bounded event (e.g., running across the finish line) or as an unbounded event (e.g., running). That is, in the same way that substances provide the material for objects, unbounded events provide the “material” for bounded events (see also Footnote 3). How one labels the very same entity can have consequences for whether the viewer individuates it or not. For example, an entity referred to as *a meatball* can be thought of as an individual but the same entity labeled as *meat* might be represented differently; similarly, an eventuality referred to as *drawing a balloon* can be thought of as an individual but the same eventuality described as *drawing* might lose its individuated nature (Barner et al., 2008; Vurgun et al., 2022; Wellwood et al., 2018). Measure phrases and temporally delimiting phrases can have similar consequences for construal. While an entity referred to as sand is likely to be construed as a non-individual, adding a measure phrase to the noun as in a pile of sand signals an individuated entity; similarly, a drawing eventuality referred to with a bounded temporal phrase as in drawing for an hour can be thought of as an individual.⁵ It is an interesting question whether count/mass syntax (for objects) or telic/atelic syntax (for events) would cross-cut the nonlinguistic construal of the same entity in our probes, and would do so across languages that encode such nominal and verbal distinctions differently (as in classifier languages; Li et al., 2009).

The issue of construal becomes even trickier by the fact that what counts as an individual has been shown to differ across adults and children: When children are asked to count labeled entities such as “forks,” they count even a detached part of a fork as a separate entity, unlike adults who count only whole forks (Shipley & Shepperson, 1990). Furthermore, beyond considerations already discussed earlier, individuation can be subject to contextual factors for both children (e.g., Srinivasan et al., 2013) and adults (Syrett & Aravind, 2022). For example, in some cases, even adults might include partial forks in their count of “forks” (Syrett & Aravind, 2022). It remains to be seen whether these contextual factors affect the signatures of individuation discussed in our study, especially since our data show that individuals otherwise resist restructuring and breaks. It also remains to be seen whether such contextual factors extend beyond partial objects to partial events (see Mathis & Papafragou, 2022).

Extensions (Constraints on Generality)

Our main findings can be extended in several ways. First, it is important to test our proposed account on a greater variety of spatial and temporal entities. While we used stimuli that were controlled for complexity, it is important to note that even within canonical individuals (e.g., objects), there are degrees to the complexity of structure: Some entities have a well-defined and complex structure (e.g., consisting of multiple distinct parts) and others have a well-defined but simpler or coarser structure. A wooden block, or a rock, for example, seemingly has a simpler structure than a wooden figurine or a statue. Nevertheless, blocks and rocks are still governed by structural principles whereby their parts are held together in a certain manner, and such entities

would be considered to have been destroyed if this arrangement was lost. Simple structure in objects is confusable with a less well-defined structure, which could lead to borderline effects (rock, rope, and string are such simple-structure cases that readily accept either count or mass syntax in English). Similarly, in the event domain, events such as clapping (encoded by semelfactives) involve a simple, binary change scale (a before and after frame that defines the event endpoint) and thus a minimal temporal structure. This minimal temporal structure is nevertheless still distinct from the lack of a well-defined structure. Future work may test how the degree of complexity contributes to structural representations. We expect our account to generalize to the cognitive profile of a broader set of entities across a wider range of complexity.

Second, even though our experimental stimuli were presented in an isolated form (objects/substances were each introduced against a white background and events were already cleanly parsed out into short video clips), our account raises the question of how they would get segmented from the continuous stream of input during our natural perception of the world. One could hypothesize that the internal structure of entities affects the viewer’s predictions about the input (in ways that bear on prediction-driven theories of event segmentation; Zacks et al., 2007). We know that boundedness can be computed online: Viewers’ attention allocation during bounded and unbounded events differs even when boundedness is not relevant to the task (Ji & Papafragou, 2022). Specifically, attention to endpoints soars for bounded but not for unbounded events. Similarly, we know from the object literature that substances behave differently from objects in tracking tasks (van Marle & Scholl, 2003; van Marle & Wynn, 2011). One could therefore extend our current findings by embedding our current stimuli into a continuous stream of activity using explicit event segmentation tasks or by probing indirect behavioral or neurophysiological markers.

Finally, our results indicate that there is a conceptual difference between individuals and nonindividuals across the object and event domains, and that this difference is supported by similar structural properties that characterize individuation broadly construed. It remains to be seen how individuation across domains interfaces with various types of cognitive processes. As mentioned already, individuals such as objects are privileged in certain cognitive operations such as counting and tracking (e.g., Chiang & Wynn, 2000; Hespos et al., 2009; Huntley-Fenner et al., 2002; Rosenberg & Carey, 2006; Soja et al., 1991; van Marle & Scholl, 2003; van Marle & Wynn, 2011). Nevertheless, nonindividuals alongside individuals can be identified, named, listed, remembered, measured (timed), and entered into causal relations: We can pick out and talk about sand at the beach; we can list substances such as wood, metal, and concrete as stuff needed to build a house; we can remember a baby crying on the airplane; and we can blame the baby’s crying for the postflight fatigue. Future work can usefully link signatures of individuation across domains to how spatial and temporal entities are recognized, tracked, named, remembered, counted, and connected through a variety of relations to other mental entities.

⁵ Overriding the default temporal interpretation of an event in language is known as aspectual coercion, a phenomenon associated with processing costs (Brennan & Pykkänen, 2008; Husband et al., 2006; Piñango et al., 1999; Todorova et al., 2000; see, e.g., Bott, 2010 for other types of coercion). We believe that language reveals shifts in entity construal because it probably reflects preexisting flexibility in underlying conceptualization.

References

- Aristotle. (1994). *Metaphysics books Z and H* (D. Bostock, Trans. & Ed.). Oxford University Press. (Original work published 350 B.C.E.).
- Baayen, R. H., Davidson, D. J., & Bates, D. M. (2008). Mixed-effects modeling with crossed random effects for subjects and items. *Journal of Memory and Language*, 59(4), 390–412. <https://doi.org/10.1016/j.jml.2007.12.005>
- Bach, E. (1986). The algebra of events. *Linguistics and Philosophy*, 9(1), 5–16. <https://doi.org/10.1007/BF00627432>
- Baldwin, D. A., Baird, J. A., Saylor, M. M., & Clark, M. A. (2001). Infants detect structure in human action: A first step toward understanding others' intentions? *Child Development*, 72(3), 708–717. <https://doi.org/10.1111/1467-8624.00310>
- Barner, D., & Snedeker, J. (2005). Quantity judgments and individuation: Evidence that mass nouns count. *Cognition*, 97(1), 41–66. <https://doi.org/10.1016/j.cognition.2004.06.009>
- Barner, D., Wagner, L., & Snedeker, J. (2008). Events and the ontology of individuals: Verbs as a source of individuating mass and count nouns. *Cognition*, 106(2), 805–832. <https://doi.org/10.1016/j.cognition.2007.05.001>
- Biederman, I. (1987). Recognition-by-components: A theory of human image understanding. *Psychological Review*, 94(2), 115–147. <https://doi.org/10.1037/0033-295X.94.2.115>
- Binnick, R. I. (Ed.). (2012). *The Oxford handbook of tense and aspect*. Oxford University Press.
- Bloom, P. (1999). The role of semantics in solving the bootstrapping problem. In R. Jackendoff, P. Bloom, & K. Wynn (Eds.), *Language, logic, and concepts: Essays in memory of John Macnamara* (pp. 285–310). MIT Press.
- Blum, H. (1973). Biological shape and visual science (Part I). *Journal of Theoretical Biology*, 38(2), 205–287. [https://doi.org/10.1016/0022-5193\(73\)90175-6](https://doi.org/10.1016/0022-5193(73)90175-6)
- Bott, O. (2010). *The processing of events*. John Benjamins Publishing.
- Brennan, J., & Pykkänen, L. (2008). Processing events: Behavioral and neuromagnetic correlates of aspectual coercion. *Brain and Language*, 106(2), 132–143. <https://doi.org/10.1016/j.bandl.2008.04.003>
- Casati, R., & Varzi, A. C. (2008). Event concepts. In T. F. Shipley & J. M. Zacks (Eds.), *Understanding events: From perception to action* (pp. 31–53). Oxford Academic.
- Champollion, L. (2015). Stratified reference: the common core of distributivity, aspect, and measurement. *Theoretical Linguistics*, 41(3–4), 109–149. <https://doi.org/10.1515/tl-2015-0008>
- Champollion, L. (2017). *Parts of a whole: Distributivity as a bridge between aspect and measurement*. Oxford University Press.
- Chiang, W. C., & Wynn, K. (2000). Infants' tracking of objects and collections. *Cognition*, 77(3), 169–195. [https://doi.org/10.1016/S0010-0277\(00\)00091-3](https://doi.org/10.1016/S0010-0277(00)00091-3)
- Davidson, D. (1970). Events and particulars. *Noûs*, 4(1), 25–32. <https://doi.org/10.2307/2214289>
- De Freitas, J., Liverence, B. M., & Scholl, B. J. (2014). Attentional rhythm: A temporal analogue of object-based attention. *Journal of Experimental Psychology: General*, 143(1), 71–76. <https://doi.org/10.1037/a0032296>
- Dowty, D. (1979). *Word meaning and Montague grammar*. Kluwer.
- Duncan, J. (1984). Selective attention and the organization of visual information. *Journal of Experimental Psychology: General*, 113(4), 501–517. <https://doi.org/10.1037/0096-3445.113.4.501>
- Filip, H. (2012). Lexical aspect. In R. I. Binnick (Ed.), *The Oxford handbook of tense and aspect* (pp. 721–751). Oxford University Press.
- Gordon, P. (1985). Evaluating the semantic categories hypothesis: The case of the count/mass distinction. *Cognition*, 20(3), 209–242. [https://doi.org/10.1016/0010-0277\(85\)90009-5](https://doi.org/10.1016/0010-0277(85)90009-5)
- Hespos, S. J., Ferry, A. L., & Rips, L. J. (2009). Five-month-old infants have different expectations for solids and liquids. *Psychological Science*, 20(5), 603–611. <https://doi.org/10.1111/j.1467-9280.2009.02331.x>
- Hoffman, D. D., & Richards, W. A. (1984). Parts of recognition. *Cognition*, 18(1–3), 65–96. [https://doi.org/10.1016/0010-0277\(84\)90022-2](https://doi.org/10.1016/0010-0277(84)90022-2)
- Hummel, J. E., & Biederman, I. (1992). Dynamic binding in a neural network for shape recognition. *Psychological Review*, 99(3), 480–517. <https://doi.org/10.1037/0033-295X.99.3.480>
- Huntley-Fenner, G., Carey, S., & Solimando, A. (2002). Objects are individuals but stuff doesn't count: Perceived rigidity and cohesiveness influence infants' representations of small groups of distinct entities. *Cognition*, 85(3), 203–221. [https://doi.org/10.1016/S0010-0277\(02\)00088-4](https://doi.org/10.1016/S0010-0277(02)00088-4)
- Husband, E. M., Beretta, A., & Stockall, L. (2006). *Aspectual computation: Evidence for immediate commitment*. In Talk given at Architectures and Mechanisms of Language Processing 12th Annual Conference, Nijmegen, The Netherlands.
- Jackendoff, R. (1991). Parts and boundaries. *Cognition*, 41(1–3), 9–45. [https://doi.org/10.1016/0010-0277\(91\)90031-X](https://doi.org/10.1016/0010-0277(91)90031-X)
- Ji, Y., & Papafragou, A. (2019), November 7–10. *Children are sensitive to the internal temporal profiles of events*. Talk at the 44th Boston University Conference on Language Development.
- Ji, Y., & Papafragou, A. (2020a). Is there an end in sight? Viewers' sensitivity to abstract event structure. *Cognition*, 197, Article 104197. <https://doi.org/10.1016/j.cognition.2020.104197>
- Ji, Y., & Papafragou, A. (2020b). Midpoints, endpoints and the cognitive structure of events. *Language, Cognition and Neuroscience*, 35(10), 1465–1479. <https://doi.org/10.1080/23273798.2020.1797839>
- Ji, Y., & Papafragou, A. (2022). Boundedness in event cognition: Viewers spontaneously represent the temporal texture of events. *Journal of Memory and Language*, 127, Article 104353. <https://doi.org/10.1016/j.jml.2022.104353>
- Kuhn, J., Geraci, C., Schlenker, P., & Strickland, B. (2021). Boundaries in space and time: Iconic biases across modalities. *Cognition*, 210, Article 104596. <https://doi.org/10.1016/j.cognition.2021.104596>
- Lee, S. H., Ji, Y., & Papafragou, A. (2023, May 4). *Signatures of individuation across objects and events*. <https://osf.io/bwkdv>
- Li, P., Dunham, Y., & Carey, S. (2009). Of substance: The nature of language effects on entity construal. *Cognitive Psychology*, 58(4), 487–524. <https://doi.org/10.1016/j.cogpsych.2008.12.001>
- Link, G. (1983). The logical analysis of plurals and mass terms: A lattice-theoretic approach. In R. Bauerle, C. Schwarze, & A. von Stechow (Eds.), *Meaning, use, and interpretation of language* (pp. 302–323). Walter de Gruyter.
- Malaia, E., Ranaweera, R., Wilbur, R., & Talavage, T. (2012). Event segmentation in a visual language: Neural bases of processing American Sign Language predicates. *Neuroimage*, 59(4), 4094–4101. <https://doi.org/10.1016/j.neuroimage.2011.10.034>
- Marr, D., & Nishihara, H. K. (1978). Representation and recognition of the spatial organization of three-dimensional shapes. *Proceedings of the Royal Society of London. Series B. Biological Sciences*, 200(1140), 269–294. <https://doi.org/10.1098/rspb.1978.0020>
- Mathis, A., & Papafragou, A. (2022). Agents' goals affect construal of event endpoints. *Journal of Memory and Language*, 127, Article 104373. <https://doi.org/10.1016/j.jml.2022.104373>
- Miller, G. A., & Johnson-Laird, P. N. (1976/2014). *Language and perception*. Harvard University Press.
- Mourelatos, A. P. (1978). Events, processes, and states. *Linguistics and Philosophy*, 2(3), 415–434. <https://doi.org/10.1007/BF00149015>
- Papafragou, A., & Ji, Y. (2023). Events and objects are similar cognitive entities. *Cognitive Psychology*, 143, Article 104353. <https://doi.org/10.1016/j.cogpsych.2023.101573>
- Parsons, T. (1990). *Events in the semantics of English: A study in subatomic semantics*. MIT Press.
- Piaget, J. (1955). *The child's construction of reality*. Oxford University Press.
- Piñango, M. M., Zurif, E., & Jackendoff, R. (1999). Real-time processing implications of enriched composition at the syntax-semantics interface.

- Journal of Psycholinguistic Research*, 28(4), 395–414. <https://doi.org/10.1023/A:1023241115818>
- Pinker, S. (1997). Words and rules in the human brain. *Nature*, 387(6633), 547–548. <https://doi.org/10.1038/42347>
- Prasada, S., Ferenz, K., & Haskell, T. (2002). Conceiving of entities as objects and as stuff. *Cognition*, 83(2), 141–165. [https://doi.org/10.1016/S0010-0277\(01\)00173-1](https://doi.org/10.1016/S0010-0277(01)00173-1)
- Quine, W. V. (1996). Events and reification. In R. Casati & A. C. Varzi (Eds.), *Events*. (pp. 107–116). Dartmouth. (Original work published 1985).
- Quine, W. V. O. (1960). *Word and object*. MIT Press.
- R Core Team. (2022). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. <https://www.R-project.org/>
- Rips, L. J., & Hespos, S. J. (2015). Divisions of the physical world: Concepts of objects and substances. *Psychological Bulletin*, 141(4), 786–811. <https://doi.org/10.1037/bul0000011>
- Rosenberg, R. D., & Carey, S. (2006). Infants' indexing of objects vs. non-cohesive substances. *Journal of Vision*, 6(6), 611. <https://doi.org/10.1167/6.6.611>
- Scholl, B. J. (2001). Objects and attention: The state of the art. *Cognition*, 80(1–2), 1–46. [https://doi.org/10.1016/S0010-0277\(00\)00152-9](https://doi.org/10.1016/S0010-0277(00)00152-9)
- Shipley, E. F., & Shepperson, B. (1990). Countable entities: Developmental changes. *Cognition*, 34(2), 109–136. [https://doi.org/10.1016/0010-0277\(90\)90041-H](https://doi.org/10.1016/0010-0277(90)90041-H)
- Soja, N. N., Carey, S., & Spelke, E. S. (1991). Ontological categories guide young children's inductions of word meaning: Object terms and substance terms. *Cognition*, 38(2), 179–211. [https://doi.org/10.1016/0010-0277\(91\)90051-5](https://doi.org/10.1016/0010-0277(91)90051-5)
- Spelke, E. S. (1988). The origins of physical knowledge. In L. Weiskrantz (Ed.), *Thought without language* (pp. 168–184). Clarendon Press/Oxford University Press.
- Spelke, E. S. (1990). Principles of object perception. *Cognitive Science*, 14(1), 29–56. https://doi.org/10.1207/s15516709cog1401_3
- Spelke, E. S. (1994). Initial knowledge: Six suggestions. *Cognition*, 50(1–3), 431–445. [https://doi.org/10.1016/0010-0277\(94\)90039-6](https://doi.org/10.1016/0010-0277(94)90039-6)
- Srinivasan, M., Chestnut, E., Li, P., & Barner, D. (2013). Sortal concepts and pragmatic inference in children's early quantification of objects. *Cognitive Psychology*, 66(3), 302–326. <https://doi.org/10.1016/j.cogpsych.2013.01.003>
- Strickland, B., Geraci, C., Chemla, E., Schlenker, P., Kelepir, M., & Pfau, R. (2015). Event representations constrain the structure of language: Sign language as a window into universally accessible linguistic biases. *Proceedings of the National Academy of Sciences of the United States of America*, 112(19), 5968–5973. <https://doi.org/10.1073/pnas.1423080112>
- Syrett, K., & Aravind, A. (2022). Context sensitivity and the semantics of count nouns in the evaluation of partial objects by children and adults. *Journal of Child Language*, 49(2), 239–265. <https://doi.org/10.1017/S0305000921000027>
- Taylor, B. (1977). Tense and continuity. *Linguistics and Philosophy*, 1(2), 199–220. <https://doi.org/10.1007/BF00351103>
- Todorova, M., Straub, K., Badecker, W., & Frank, R. (2000). *Aspectual coercion and the on-line computation of sentential aspect*. In Proceedings of the twenty-second annual conference of the Cognitive Science Society, Philadelphia, PA.
- Truswell, R. (Ed.). (2019). *The Oxford handbook of event structure*. Oxford University Press.
- van Hout, A. (2016). Lexical and grammatical aspect. In J. Lidz, W. Snyder, & J. Pater (Eds.), *The Oxford handbook of developmental linguistics* (pp. 587–610). Oxford University Press.
- van Marle, K., & Scholl, B. J. (2003). Attentive tracking of objects vs. substances. *Psychological Science*, 14(5), 498–504. <https://doi.org/10.1111/1467-9280.03451>
- van Marle, K., & Wynn, K. (2011). Tracking and quantifying objects and non-cohesive substances. *Developmental Science*, 14(3), 502–515. <https://doi.org/10.1111/j.1467-7687.2010.00998.x>
- Vendler, Z. (1957). Verbs and times. *The Philosophical Review*, 66(2), 143–160. <https://doi.org/10.2307/2182371>
- Vurgun, U., Ji, Y., & Papafragou, A. (2022). *Language affects event representations*. In Talk at the 2nd Conference on Experiments in Linguistic Meaning (ELM), University of Pennsylvania, May 18–21.
- Wagner, L., & Carey, S. (2003). Individuation of objects and events: A developmental study. *Cognition*, 90(2), 163–191. [https://doi.org/10.1016/S0010-0277\(03\)00143-4](https://doi.org/10.1016/S0010-0277(03)00143-4)
- Wehry, J., Hafri, A., & Trueswell, J. C. (2019). *The end's in plain sight: Implicit association of visual and conceptual boundedness*. In Proceedings of the twenty-second annual conference of the Cognitive Science Society (pp. 1185–1191).
- Wellwood, A., Hespos, S. J., & Rips, L. (2018). The object: Substance: Event: Process analogy. In T. Lombrozo, J. Knobe, & S. Nichols (Eds.), *Oxford Studies in experimental philosophy* (Vol. 2, pp. 183–212). Oxford University Press.
- Wittenberg, E., & Levy, R. (2017). If you want a quick kiss, make it count: How choice of syntactic construction affects event construal. *Journal of Memory and Language*, 94, 254–271. <https://doi.org/10.1016/j.jml.2016.12.001>
- Yates, T., Sherman, B., & Yousif, S. (2023). More than a moment: What does it mean to call something an 'event'? *Psychonomic Bulletin & Review*, 30(6), 2067–2083. <https://doi.org/10.3758/s13423-023-02311-4>
- Zacks, J. M., Speer, N. K., Swallow, K. M., Braver, T. S., & Reynolds, J. R. (2007). Event perception: A mind–brain perspective. *Psychological Bulletin*, 133(2), 273–293. <https://doi.org/10.1037/0033-2909.133.2.273>
- Zacks, J. M., & Swallow, K. M. (2007). Event segmentation. *Current Directions in Psychological Science*, 16(2), 80–84. <https://doi.org/10.1111/j.1467-8721.2007.00480.x>
- Zacks, J. M., & Tversky, B. (2001). Event structure in perception and conception. *Psychological Bulletin*, 127(1), 3–21. <https://doi.org/10.1037/0033-2909.127.1.3>
- Zehr, J., & Schwarz, F. (2018). *PennController for Internet Based Experiments (IBEX)*. <https://doi.org/10.17605/OSF.IO/MD832>

Received May 4, 2023

Revision received February 12, 2024

Accepted February 21, 2024 ■