



On the superiority of defaultness: Hemispheric perspectives of processing negative and affirmative sarcasm

Rachel Giora, Adi Cholev, Ofer Fein & Orna Peleg

To cite this article: Rachel Giora, Adi Cholev, Ofer Fein & Orna Peleg (2018) On the superiority of defaultness: Hemispheric perspectives of processing negative and affirmative sarcasm, *Metaphor and Symbol*, 33:3, 163-174

To link to this article: <https://doi.org/10.1080/10926488.2018.1481259>



Published online: 01 Aug 2018.



Submit your article to this journal 



View Crossmark data 



On the superiority of defaultness: Hemispheric perspectives of processing negative and affirmative sarcasm

Rachel Giora^a, Adi Choleva^a, Ofer Fein^b, and Orna Peleg^a

^aTel Aviv University; ^bThe Academic College of Tel Aviv-Yaffo

ABSTRACT

Defining defaultness in terms of an unconditional, automatic response to a stimulus allows the Defaultness Hypothesis to predict the speed superiority of default over nondefault counterparts. Here we examined the relative contribution of the cerebral hemispheres to the processing of default versus nondefault interpretations (of Hebrew items). Participants performed a lexical decision task on lateralized probes (*messy*) related to the default/nondefault sarcastic interpretation of their preceding negative/affirmative targets (*He is/He is not the most organized student*). Consistent with the Defaultness Hypothesis, probes were easier to identify in the default-negative than in the nondefault-affirmative condition. However, this superiority was more pronounced in the left hemisphere (LH) than in the right hemisphere. In particular, whereas both hemispheres reflected the superiority of default negative sarcasm over nondefault affirmative sarcasm when both targets were embedded in equally strong sarcastic contexts (Experiment 2), only the LH exhibited this very same superiority when the negatives and affirmatives were presented in isolation (Experiment 1).

Introduction

The Defaultness Hypothesis (Giora, Givoni, & Fein, 2015) posits the speed superiority of default over nondefault responses. Within the framework of the Defaultness Hypothesis, defaultness is defined in terms of an unconditional, automatic response to a stimulus. The focus here is on *interpretations* (i.e., on responses *constructed* on the fly [rather than *accessed* directly from the mental lexicon]).

For such noncoded, constructed responses to be generated *by default*, stimuli should be potentially *ambiguous* between literal and nonliteral interpretations, so that a *preference* is allowed *a priori*. They should, therefore, be (a) novel/noncoded (e.g., Giora, 2003; Mashal & Faust, 2009); (b) free of utterance-internal cues inviting nonliteralness, such as semantic anomaly or internal incongruity (e.g., Beardsley, 1958; Partington, 2011); and (c) free of utterance-external cues inviting non/literalness, such as specific contextual information or explicit marking (e.g., *literally*, *#Sarcasm*, or Hebrew *Staan* [*not really*]; see Campbell & Katz, 2012; Gibbs, 1994; Sulis, Hernandez Farias, Rosso, Patti, & Ruffo, 2016; Ziv, 2013, respectively).

DT to be able to test the anticipated speed superiority of default interpretation over nondefault counterparts, Giora et al. (2015: Experiment 1) first established degree of defaultness by probing items, meeting conditions (a–c), for their preferred (literal vs. sarcastic) interpretation. Results of an offline rating task indicated that, when presented in isolation, the preferred interpretation of the novel, negative items (*He is not the most organized student*) was sarcastic (*messy*); their nondefault nonpreferred interpretation was literal (*ordered*). In contrast, the default preferred interpretation of their equally novel affirmative counterparts (*He is the most organized student*) was literal (*ordered*);

their nondefault nonpreferred interpretation was sarcastic (*messy*). Having established degree of defaultness, Giora et al. (2015: Experiment 2) then weighed the processing speed of default versus nondefault counterparts when embedded in equally strong contexts, supportive of their respective interpretations. Results attested to the speed superiority of default over nondefault interpretations, regardless of negation and equal strength of contextual support. Specifically, **default** Negative Sarcasm (1 in Figure 1) was faster to process than **nondefault** Negative Literalness (3 in Figure 1) and faster yet than **nondefault** Affirmative Sarcasm (2 in Figure 1). Similarly, **default** Affirmative Literalness (4 in Figure 1) was faster to process than **nondefault** Negative Literalness (3 in Figure 1) and also faster than **nondefault** Affirmative Sarcasm (2 in Figure 1).

Giora et al. (2015), then, were able to show that, processing-wise, defaultness prevails, irrespective of factors known to affect processing, such as novelty, negation, non/literalness, or contextual support (see also Filik, Howman, Ralph-Nearman, & Giora, 2018).

The aim of the present study is to examine the relative contribution of the two cerebral hemispheres to the processing of default versus nondefault interpretations. Although it is well established that the left hemisphere (LH) is dominant for language processes, it is now widely acknowledged that *both* hemispheres are engaged in comprehending language, albeit in qualitatively different ways (e.g., Coulson & Williams, 2005; Eviatar & Just, 2006; Faust & Chiarello, 1998; Federmeier & Kutas, 1999; Giora, Zaidel, Soroker, Batori, & Kasher, 2000; Mashal, Faust, & Hendler, 2005; Titone, 1998).

One way to assess hemispheric differences in language processing is by using the divided visual field (DVF) technique. This technique takes advantage of the fact that stimuli presented in the left side of the visual field are initially processed by the right hemisphere (RH) and vice versa. Although information presented in this manner can be later transmitted to both hemispheres, the interpretation of DVF paradigms rests on the assumption that responses to stimuli, presented briefly to one visual field, reflect mainly the processing of that stimulus by the contralateral hemisphere, so that responses to probes displayed in the right visual field (RVF) reflect LH processes; responses to probes displayed in the left visual field (LVF) reflect processes in the RH (for theoretical and

Negatives	Affirmatives
(1) Default Sarcasm During the Communication Department staff meeting, the professors are discussing their students' progress. One of the students has been doing very poorly. Professor A: "Yesterday he handed in an exercise and, once again, I couldn't make any sense of the confused ideas presented in it. The answers were clumsy, unfocused, and the whole paper was hard to follow." Professor B nods in agreement and adds: "Unfortunately, the problem isn't only with his assignments. He is also always late for class, and when it was his turn to present a paper in class he got confused and prepared the wrong essay! I was shocked. What can I say, he is not the most organized student. <i>I'm surprised</i> he didn't learn a lesson from his freshman year experience."	(2) Nondefault Sarcasm During the Communication Department staff meeting, the professors are discussing their students' progress. One of the students has been doing very poorly. Professor A: "Yesterday he handed in an exercise and, once again, I couldn't make any sense of the confused ideas presented in it. The answers were clumsy, unfocused, and the whole thing was hard to follow." Professor B nods in agreement and adds: "Unfortunately, the problem isn't only with his assignments. He is also always late for class, and when it was his turn to present a paper in class he got confused and prepared the wrong essay!" Professor C (chuckles): "In short, it sounds like he really has everything under control." Professor A: "What can I say, he is the most organized student. <i>I'm surprised</i> he didn't learn a lesson from his freshman year experience."
(3) Nondefault Literalness The professors are talking about Omer, one of the department's most excellent students. Professor A: "He is a very efficient lad. Always comes to class on time with all of his papers in order and all his answers are eloquent, exhibiting a clearly structured argumentation. I think that explains his success." Professor B: "Yes, it's true. Omer is simply very consistent and almost never digresses from the heart of the matter. But there are two other students whose argumentation and focus surpass his, so I'd just say that, in comparison to those two, he is not the most organized student. <i>I'm surprised</i> he asked to sit the exam again."	(4) Default Literalness During the Communication Department staff meeting, the professors are discussing their students' progress. One of the students has been doing very well. Professor A: "He is the most committed student in the class. Always on time, always updated on everything. Professor B: "I also enjoy his answers in class. He always insists on a clear argumentation structure and is very eloquent. In his last exam, not only was each answer to the point but also very clear. In my opinion, he is the most organized student. <i>I'm surprised</i> he asked to sit the exam again."

Figure 1. Default and nondefault affirmative and negative items in context. Targets in bold; spillover segments in italics.

electrophysiological support for this assumption, see Banich, 2003; Berardi & Fiorentini, 1997; Coulson, Federmeier, Van Petten, & Kutas, 2005).

Whereas in Giora et al. (2015), items were presented centrally (to both hemispheres), here we utilized the DVF technique in order to examine the hemispheric perspective of processing a minimal pair taken from Giora et al. (2015)—default Negative Sarcasm (*He is not the most organized student*) and nondefault Affirmative Sarcasm (*He is the most organized student*). Items were tested here for Response Time (RT) and Response Accuracy to probes related to the sarcastic interpretation of affirmative and negative targets. Specifically, in Experiment 1, targets were presented in isolation; in Experiment 2 they were embedded in highly supportive contexts (see 1–2 in Figure 1). In both cases, however, they were followed by the same probe-word, related to their possible sarcastic interpretation (*messy*), which instantiated the *default* sarcastic interpretation, in the case of the negatives (see example 1 in Figure 1), and the *nondefault* sarcastic interpretation, in the case of the affirmatives (see example 2 in Figure 1). In both experiments, participants had to make a lexical decision (by pressing a YES or a NO key) as to whether the probe is a word or a nonword. Probes were presented either to the RVF or to the LVF. Using identical probes (*messy*), following both negatives and affirmatives, allowed Experiment 1 to establish the degree of defaultness of the interpretations of the affirmative and negative items.

Looking at the hemispheric division of labor, we aimed to find out which of the hemispheres specializes in processing default, automatic interpretations, and which specializes in processing nondefault (e.g., context-dependent) interpretations. Among other things, the literature on the hemispheric division of labor suggests that familiar, frequent, or conventionalized (i.e., default) meanings will engage the LH (e.g., Giora, 2003; Jung-Beeman, 2005); novel, non-conventionalized (i.e., nondefault) meanings, such as noncoded ironies or metaphors, will engage the RH (e.g., Federmeier, 2007; Giora, 2003; Jung-Beeman, 2005; Mashal & Faust, 2009), which is also involved in making sense of discourse, and even more so, of nonliteral utterances in discourse (e.g., Newman, Just, & Mason, 2003). Indeed, the RH specializes in reinterpretation and integration of linguistic targets in context (see, e.g., Bihrle, Brownell, Powelson, & Gardner, 1986; Brownell, Michel, Powelson, & Gardner, 1983; Brownell, Potterm, Bihrle, & Gardner, 1986; Giora et al., 2000; Zaidel, 1979). We therefore predicted that the superiority of default over nondefault interpretations will be more pronounced in the LH than in the RH.

In Experiment 1, items were the negative and affirmative utterances (*He is/is not the most organized student*), controlled for novelty and degree of defaultness by Giora et al. (2015). They were presented in isolation, followed by a sarcastically related (*messy*), unrelated, or a nonword probe. Participants were asked to make a lexical decision (by pressing a Yes or a No key), as to whether the letter string constituted a word or a nonword. Experiment 2 used the same affirmative and negative targets used in Experiment 1, only here they were embedded in contexts, equally strongly supportive of their targets' sarcastic interpretation. Participants had to make a lexical decision (by pressing a Yes or a No key) as to whether a letter string (*messy*), following the target utterance, was a word or a nonword. Thus, while Experiment 1 is focused mainly on automatic activation processes, Experiment 2 emphasizes controlled integration and selection processes.

In accordance with the Defaultness Hypothesis, we predicted faster and more accurate responses in the default-negative condition than in the nondefault-affirmative condition. However, because processing in the LH tends to be narrowly focused, whereas RH processes tend to be broader in scope (Jung-Beeman, 2005), we predicted that, in the LH, only default interpretations will be activated, whereas, in the RH, both default and nondefault interpretations will be activated. As a result, the superiority of the default-negative condition over the nondefault-affirmative condition will be more pronounced in the LH than in the RH, particularly when targets are presented in isolation (Experiment 1).

In addition to asymmetries in meaning activation, several studies have suggested that the two hemispheres differ in their ability to carry out meaning selection and integration (e.g., Burgess & Simpson, 1988; Faust & Gernsbacher, 1996). According to these studies, the LH activates and selects

only the contextually appropriate meaning, whereas the RH activates and maintains multiple meanings, irrespective of context. However, more recent studies have shown that both hemispheres perform controlled selection and integration processes, albeit in qualitatively different ways (e.g., Federmeier, 2007; Peleg & Eviatar, 2017). In particular, it has been suggested that, in the RH, default interpretations may be easier to select than nondefault interpretations (Kacirik & Chiarello, 2007). We therefore predicted that, in Experiment 2, the RH may also attest to the superiority of default over nondefault interpretations. This is because, even if the RH activates both, default and nondefault interpretations, the contextually relevant (here sarcastic) interpretation will be easier to select in the *default*-negative condition than in the *nondefault*-affirmative condition.

Experiment 1

The aim of Experiment 1 was to test degree of defaultness of negative and affirmative counterparts when presented in isolation, followed by an identical probe word related to their sarcastic interpretation. Response speed and/or response accuracy will establish items' degree of defaultness.

Method

Participants

Twenty undergraduate students (eight males), aged 22–31, participated in the study. They were all healthy, right handed, native speakers of Hebrew, with normal or corrected-to-normal vision. Handedness was assessed with the Edinburgh Handedness Questionnaire (Oldfield, 1971), with 80 as the cutoff point. They were paid 60 NIS for their participation.

Stimuli

The experimental stimuli were the same 12 sentence pairs used in Giora et al. (2015). Each pair consisted of a negative utterance (*He is not the most organized student*) and its affirmative counterpart (*He is the most organized student*). As shown in Giora et al. (2015: Experiment 1), when presented in isolation, the negative constructions were interpreted sarcastically by default; their affirmative counterparts, however, were interpreted literally by default.

As mentioned earlier, for each sentence pair used here, a probe word was selected that was the opposite of what was said. For example, *He is/is not the most organized student* was paired with the probe word (*messy*). Thus, the probe was always related to the sarcastic interpretation of the utterance, even when only optional, as in the case of the affirmatives.

Pretest

To establish the extent to which the selected probes reflect the sarcastic interpretation of the experimental items, a pretest was conducted, with all the items presented in isolation. Thirty judges were asked to rate the relatedness of each probe word to the utterance preceding it, on a 7-point scale, where 1 indicated "not related at all" and 7 indicated "strongly related." In order to avoid repetition, two lists were prepared (each presented to 15 participants), such that probes (*messy*), which followed their negative utterance in the first list (*He is not the most organized student*), followed their affirmative counterparts (*He is the most organized student*) in the second list, and vice versa. Overall, each list comprised 12 experimental items (six in the negative and six in the affirmative), and 10 filler items (five in the negative and 5 in the affirmative). The probes of the filler items were either weakly related or entirely unrelated to their preceding utterance. The mean relatedness score of the negative items was 5.57, whereas the mean score of the affirmative items was

1.46. This suggests that the antonyms selected as probes indeed reflect the sarcastic interpretation, which, as assumed, is the default interpretation of the negative utterances only.

In addition to the 12 experimental items (12 negative–affirmative sentence pairs and their corresponding 12 probe words), 36 filler items were constructed. Like the experimental items, each filler item consisted of a negative–affirmative sentence pair and a letter string probe. First, 12 filler items were constructed in which the probes were not antonymic to the critical concept presented in their preceding utterance. Of these 12 fillers, six were semantically (albeit not antonymically) related to the utterance (e.g., *This shirt suits/does not suit me. Beautiful*), and the other six were completely unrelated (e.g., *He is so Polish/not Polish. Green*). Second, given that the 12 experimental items and the 12 filler items included real words as probes (to which a “YES” response in the lexical decision task was expected), 24 additional fillers were constructed with nonwords as probes (to which a “NO” response in the lexical decision task was expected) (e.g., *We need/don’t need that. Gavtam*). Overall, 48 probes (24 words and 24 nonwords) were included.

Each of these 48 probes was presented in four different conditions: 2 sentence versions (negative or affirmative) $\times$ 2 visual fields (RVF/LH or LVF/RH). Four lists were prepared such that each probe appeared only once in a list, each time in a different condition. Within each list, 24 probes were presented to the RVF/LH (12 preceded by negative utterances and 12 preceded by affirmative utterances), and the other 24 were presented to the RVF/LH (12 preceded by negative utterances and 12 by affirmative utterances).

Apparatus

The experiment was constructed and run using E-Prime software version 10.242, on an HP Compaq Elite 8300 Microtower desktop computer. Response latencies were collected using a PST Serial Response Box. To ensure central fixation, participants’ eye position was monitored with an iView SMI RED-m- Eye Tracker.

Experimental design and procedures

The experiment used a 2 (Sentence type: negative or affirmative) $\times$ 2 (probe location: RVF/LH or LVF/RH) within participants design. There were 48 experimental permutations for the 12 critical probe words (12 probes $\times$ 2 types of sentences $\times$ 2 visual field [VF] presentations). Four lists were prepared such that all factors were counterbalanced across items and participants. Each list comprised 4 practice trials, 12 experimental trials, and 36 fillers (in all, 52 trials).

Participants were tested individually in a sound-attenuated room. All trials had the same sequence of events. At the start of each trial, participants were presented with a central fixation marker for 500 ms. The offset of the marker was followed by a sentence, presented in the center of the screen for 1,500 ms (identified earlier as comfortable for reading all the sentences presented in the experiment). The offset of the sentence was followed by a central fixation marker, together with a probe string that was presented for 150 ms either to the RVF/LH or to LVF/RH for a lexical decision task. Participants made a Yes or a No decision by pressing the corresponding key on the response box. Response latencies were measured from the onset of the probe presentation until the pressing of the key. Probes were presented such that their innermost boundary, whether to the right or left of the center, was exactly 2° of visual angle from the central fixation marker. Stimuli subtended a maximum of 2.5° of visual angle.

Each participant completed the four lists in two experimental sessions (two lists per session). Trials within each list were presented in a random order, with randomization controlled by the computer. Additionally, the order of the lists was counterbalanced across participants. The sessions were administered with an interval of 2–3 weeks between them. Each session lasted approximately 30 minutes (10 minutes for each list with a 10-minute break between them). Cell means were based on 12 experimental trials per condition per participant.

Results

A 2×2 analysis of variance (ANOVA) was conducted for both RT and Accuracy Rate, with Sentence Type (Negative vs. Affirmative) and Visual Field (RVF/LH vs. LVF/RH) as factors. Response time data were calculated for correct responses only. Response times above three standard deviations (*SD*) from the mean for each participant were trimmed as outliers (less than 1.7% of the data). Responses to filler trials were not analyzed.

The main effect of Visual Field was significant for both RT, $F_{RT}(1,19) = 9.65$, $p < .01$, and Accuracy, $F_{AC}(1,19) = 4.58$, $p < .05$, indicating that probes were responded to more quickly and more accurately when presented to the RVF/LH ($M_{RT} = 749$ msec, $M_{AC} = 0.92$) than to the LVF/RH ($M_{RT} = 805$ msec, $M_{AC} = 0.86$). The effect of Sentence Type, however, was not significant, $F_{RT}(1,19) = 0.99$, *n.s.*, $F_{AC}(1,19) = 0.25$, *n.s.*

Importantly, for RT data (but not for Accuracy), the two-way interaction between Visual Field and Sentence Type was significant: $F_{RT}(1,19) = 4.40$, $p = .05$, $F_{AC}(1,19) = 0.00$, *n.s.* Follow up tests revealed that for RVF/LH probe presentation, responses in the negative condition were significantly faster than in the affirmative condition, $t_{RT}(19) = 2.32$, $p < 0.05$. In contrast, for LVF/RH probe presentation, there was no difference between the two sentence conditions $t_{RT}(19) = 0.35$, *n.s.* (see Figure 2).

Discussion

When presented in isolation, negative constructions (*He is not the most organized student*) and affirmative counterparts (*He is the most organized student*), were responded to similarly accurately in both hemispheres, although more accurately in the LH than in the RH. However, it is in the LH that negative items were processed faster than affirmative counterparts. Being processed faster in the LH than affirmative alternatives establishes the defaultness of the negative items and the nondefaultness of the affirmative counterparts.

Will defaultness supersede contextual strength? Specifically, will default negative items be processed faster and more accurately than nondefault affirmative counterparts when both are embedded

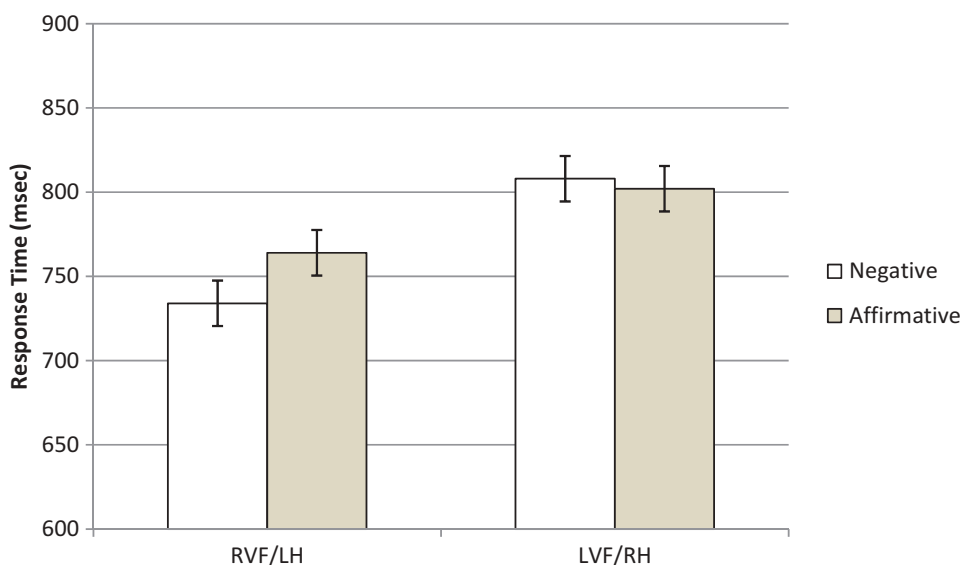


Figure 2. Response time of correct responses to probe words as a function of Sentence Type (Negative vs. Affirmative) and Visual Field (RVF/LH vs. LVF/RH)—Experiment 1. Error bars represent standard errors. Standard errors were calculated according to Loftus and Masson (1994) recommendations for within-subjects designs.

(1) Default Negative Sarcasm	(2) Nondefault Affirmative Sarcasm
During the Communication Department staff meeting, the professors are discussing their students' progress. One of the students has been doing very poorly. Professor A: "Yesterday he handed in an exercise and, once again, I couldn't make any sense of the confused ideas presented in it. The answers were clumsy, unfocused, and the whole paper was hard to follow." Professor B nods in agreement and adds: "Unfortunately, the problem isn't only with his assignments. He is also always late for class, and when it was his turn to present a paper in class he got confused and prepared the wrong essay! I was shocked. What can I say, he is not the most organized student. " <i>Messy</i>	During the Communication Department staff meeting, the professors are discussing their students' progress. One of the students has been doing very poorly. Professor A: "Yesterday he handed in an exercise and, once again, I couldn't make any sense of the confused ideas presented in it. The answers were clumsy, unfocused, and the whole thing was hard to follow." Professor B nods in agreement and adds: "Unfortunately, the problem isn't only with his assignments. He is also always late for class, and when it was his turn to present a paper in class he got confused and prepared the wrong essay!" Professor C (chuckles): "In short, it sounds like he really has everything under control." Professor A: "What can I say, he is the most organized student. " <i>Messy</i>

Figure 3. Default and nondefault affirmative and negative items in context. Targets in bold; probes in italics.

in equally strongly supportive contexts (see Figure 3)? How will the cerebral hemispheres reflect these differences?

Experiment 2

The aim of Experiment 2 was to attest to the superiority of defaultness over nondefaultness even when items are embedded in equally strongly biased contexts. We thus expected negative targets to be more accurate and speedier than affirmative counterparts.

Method

Participants

Twenty undergraduate students (seven males), aged 20–35, participated in the study. They were all healthy, right handed, native speakers of Hebrew, with normal or corrected-to-normal vision. Handedness was assessed with the Edinburgh Handedness Questionnaire (Oldfield, 1971), with 80 as the cutoff point. They were paid 60 NIS for their participation.

Stimuli

The targets were identical to those used in Experiment 1. Only, here, they were preceded by contexts (short of the spillover segments), taken from Giora et al. (2015), which were equally highly supportive of their sarcastic interpretation (see 1–2 in Figure 2). Here too, each probe word (*messy*) was paired with two sarcastic items, one ending in a negative utterance (*He is not the most organized student*) and one in its affirmative equivalent (*He is the most organized student*).

Pretest

To ensure that the experimental probes (*messy*) indeed reflect the sarcastic interpretation of both the negative and affirmative versions, a pretest was conducted, in which 30 judges were asked to rate the relatedness of the probe to its preceding context on a 7-point relatedness scale, where 1 was “not related at all” and 7 was “strongly related.” In order to avoid repetition, two lists were prepared (each involving 15 participants), such that probes, presented in the negative version in the first list, were presented in the affirmative version in the second list, and vice versa. Overall, each list included 12 experimental sarcastic items and their corresponding sarcastic probes (six contexts ending in a negative utterance and six in an affirmative utterance), and 10 filler items (five ending in a negative

utterance and five in an affirmative utterance). The probes in the filler items were either weakly related or unrelated to their preceding utterance. The mean relatedness score for the experimental probes was 5.57 in the negative version and 6.02 in the affirmative version, with no significant difference between these two conditions. This suggests that the antonyms selected as probes indeed reflect the sarcastic interpretation of the final target utterance, which is the default interpretation, in the case of the negative versions, and the nondefault interpretation, in the case of the affirmative counterparts.

Contexts were also constructed for the 36 filler items (as well as for the four practice trials). Thus, overall, the experiment included the same 48 items as in Experiment 1 (12 experimental items and 36 filler items), all, however, embedded in a highly supportive context. Each probe word was presented in four different conditions: 2 Sentence Type (negative or affirmative) $\times$ 2 visual fields (RVF/LH or LVF/RH). Four lists were prepared such that each item appeared only once in a list, each time in a different condition. Within each list, 24 probes were presented to the RVF/LH (12 with negative utterances and 12 with affirmative utterances), and the other 24 to the LVF/RH (12 with negative utterances and 12 with affirmative utterances).

Apparatus

The experiment was constructed and run using E-Prime software version 10.242, on an HP Compaq Elite 8300 Microtower desktop computer. Response latencies were collected using a PST Serial Response Box. To ensure central fixation, participants' eye position was monitored with an iView SMI RED-m- Eye Tracker.

Experimental design and procedures

The experimental design and procedures were exactly the same as in Experiment 1, except for the target sentences that, here, were embedded in equally strong, sarcastically biasing contexts (pretested for equal strength of contextual support here and in Giora et al., 2015). All trials had the same sequence of events. At the start of each trial, participants were presented with a central fixation marker for 500 ms. The offset of the marker was followed by the context (without the final target sentence). Participants self-paced their reading of the context, which was displayed segment by segment (making up part of a sentence or a full sentence). They advanced the text by pressing the spacebar. After reading the context, the target sentence was presented in the center of the screen for 1,500 ms. The offset of the sentence was followed by a central fixation marker, together with a probe string that was presented for 150 ms either to the LVF/RH or to the RVF/LH for a lexical decision response.

Each participant completed the four lists in two experimental sessions (two lists per session). Trials within each list were presented in a random order, with randomization controlled by the computer and the order of lists counterbalanced across participants. The sessions were administered with an interval of 2–3 weeks between them. Each testing session lasted approximately 60 minutes (20 minutes for each list, including a 10-minute break between them). Cell means were based on 12 experimental trials per condition per participant.

Results

As in Experiment 1, two 2×2 ANOVAs were conducted (one for Response Time and one for Accuracy), with Sentence Type and Visual Field as factors.

The main effect of Visual Field was significant for RT, $F_{RT}(1,19) = 10.10$, $p < .01$, but not for Accuracy, $F_{AC}(1,19) = 2.34$, $p = .14$, indicating that probes were responded to more quickly when presented to the RVF/LH ($M_{RT} = 840$ msec) than to the LVF/RH ($M_{RT} = 931$ msec).

The main effect for sentence type was significant for both RT, $F_{RT}(1,19) = 4.85$, $p < .05$, and Accuracy, $F_{AC}(1,19) = 6.73$, $p < .05$, indicating that probes were responded to more quickly and

more accurately in the negative ($M_{RT} = 867$ msec, $M_{AC} = 0.89$) than in the affirmative version ($M_{RT} = 904$ msec, $M_{AC} = 0.85$). However, the two-way interaction between Visual Field and Sentence Type was not significant $F_{RT}(1,19) = 2.34$, $p = .14$, $F_{AC}(1,19) = 1.16$, $n.s.$, indicating that both hemispheres were sensitive to the negative–affirmative manipulation.

Discussion

Results of Experiment 2 support the Defaultness Hypothesis. They show that, despite being longer (in terms of word number) than affirmative counterparts, default negative sarcasm fared better than nondefault affirmative sarcasm, in terms of response accuracy, response speed, or both. The superiority of defaultness over nondefaultness, then, overrides degree of negation, novelty, nonliteralness, or contextual strength.

General discussion

In this study we examined the sensitivity of the cerebral hemispheres to degree of defaultness of interpretations, derived when stimuli were presented in and out of context. Items were default negative sarcasm and nondefault affirmative sarcasm (*He is/He is not the most organized student*), established as such by Giora et al. (2015). Hence, when embedded in strong contexts, equally supportive of their respective interpretations, negative items were processed faster than affirmative counterparts, attesting to the speed superiority of defaultness (Giora et al., 2015).

Here we examined the hemispheric perspectives of default and nondefault interpretations. First, to establish their degree of defaultness, both affirmative and negative targets were presented in isolation, followed by a probe word (*messy*), reflecting their sarcastic interpretation, which was their default interpretation in the negative condition and their nondefault interpretation in the affirmative condition. The same probe words also followed the targets, when embedded in highly strong contexts supportive of their sarcastic interpretations (controlled for biasing strength here and in Giora et al., 2015). Measures were response speed and response accuracy to probes.

Results indeed replicate those of Giora et al. (2015). They attest to the superiority of defaultness over nondefaultness, both in terms of response accuracy and response speed. When presented in isolation (Experiment 1), default negative items were processed faster than nondefault affirmative counterparts in the LH. When tested in supportive contexts (Experiment 2), default negative sarcasm was responded to more accurately and faster than nondefault affirmative sarcasm, irrespective of VF presentation. In all, default negative sarcasm exhibited its superiority over nondefault affirmative sarcasm via response speed and response accuracy.

Although both hemispheres exhibited the expected superiority of defaultness, this was more pronounced in the LH than in the RH. In particular, whereas both hemispheres reflected the superiority of default negative sarcasm over nondefault affirmative sarcasm when the negative and affirmative targets were embedded in a sarcastic context (Experiment 2), only the LH exhibited this very same superiority when the negative and affirmative targets were presented in isolation (Experiment 1).

The results of Experiment 1 can be explained within the framework of the fine-coarse coding model (Jung-Beeman, 2005), which postulates that the cerebral hemispheres differ in their breadth of semantic activation—narrow, focused semantic fields in the LH versus broader, diffused semantic fields in the RH. Specifically, according to this model, meaning activation in the LH may include only default interpretations, whereas meaning activation in the RH may include both default and nondefault interpretations. Thus, when negative targets were displayed (*He is not the most organized student*), only the default sarcastic interpretation was activated in the LH, while both the default sarcastic and possibly the nondefault literal interpretations (see 3 in Figure 1) were activated in the RH. Similarly, when affirmative targets were displayed (*He is the most organized student*), only the default literal interpretation (see 4 in Figure 1) was activated in the LH, while both the default literal and the nondefault sarcastic interpretations were activated in the RH. As a result, in the LH, the

sarcastic probe (*messy*) was easier to respond to in the negative than in the affirmative condition, while in the RH, no difference was found between the two conditions.

If, in the absence of context, the RH does not distinguish between default and nondefault interpretations (Experiment 1), why do we see the superiority of default negative sarcasm over nondefault affirmative sarcasm, when a sarcastically biasing context is provided (Experiment 2)? We suggest that, in the RH, selection rather than activation processes are affected. First, although the RH exhaustively activates both default and nondefault interpretations, when a biasing context is provided, it is capable of selecting the contextually appropriate interpretation (e.g., Peleg & Eviatar, 2008, 2009, 2017; Peleg, Markus, & Eviatar, 2012). Second, following Kacirik and Chiarello (2007), it is also assumed that, in the RH, default interpretations may be easier to select than nondefault interpretations. In other words, it is easier to disinvest from (or reject) the *nondefault* literal interpretation, in the negative condition, than the *default* literal interpretation, in the affirmative condition. As a result, sarcastic probes were faster to respond to in the negative condition than in the affirmative condition.

Although the LH showed the same pattern of results, the underlying processes that led to these results may be different. In the LH, only default interpretations are automatically activated (as shown in Experiment 1, in which early activation processes already reflect the superiority of defaultness). In the negative condition in the LH, then, the default sarcastic interpretation is compatible with its prior context and can be easily and directly integrated with the preceding context. However, in the affirmative condition in this hemisphere, the default literal interpretation is in conflict with the preceding context, thereby requiring further processes, which are costly. As a result, sarcastic probes are easier to respond to in the negative condition than in the affirmative condition.

In sum, results from two experiments support the Defaultness Hypothesis (Giora et al., 2015). While Experiment 1 replicates the superiority of negative sarcasm as a default interpretation when in isolation, Experiment 2 tests its superiority when in a supportive context. These minimal affirmative and negative pairs, identical in terms of their sarcastic interpretation and strength of contextual support, differ, however, in terms of defaultness. Therefore, they were processed differently in the brain, while attesting to the superiority of defaultness. As shown by Experiment 2, default negative sarcasm superseded nondefault affirmative sarcasm in terms of response accuracy and response speed. It is defaultness rather than degree of nonliteralness, negation, novelty, or strength of contextual support that matters.

Acknowledgments

We are very grateful to Nira Mashal and Dominic Thompson for the enlightening comments on the article, to Efrat Levant and Tal Norman for their help with running the experiments, and to Andrey Markus for programming them.

Funding

This research was supported by The Israel Science Foundation (grant no. 436/12) to Rachel Giora.

References

- Banich, M. T. (2003). Interaction between the hemispheres and its implications for the processing capacity of the brain. In R. Davidson & K. Hugdahl (Eds.), *Brain asymmetry* (Vol. 2, pp. 261–302). Cambridge, MA: MIT Press.
- Beardsley, C. M. (1958). *Aesthetics*. New York, NY: Harcourt, Brace and World.
- Berardi, N., & Fiorentini, A. (1997). Interhemispheric transfer of spatial and temporal frequency information. In S. Christman (Ed.), *Cerebral asymmetries in sensory and perceptual processing* (pp. 55–79). New York, NY: Elsevier Science.
- Bihrlé, A. M., Brownell, H. H., Powelson, J. A., & Gardner, H. (1986). Comprehension of humorous and non-humorous materials by left and right brain-damaged patients. *Brain and Cognition*, 5, 399–411. doi:10.1016/0278-2626(86)90042-4

- Brownell, H. H., Michel, D., Powelson, J., & Gardner, H. (1983). Surprise but not coherence: Sensitivity to verbal humor in right-hemisphere patients. *Brain and Language*, 18, 20–27. doi:10.1016/0093-934X(83)90002-0
- Brownell, H. H., Potterm, H. H., Bihrl, A. M., & Gardner, H. (1986). Inference deficits in right brain-damaged patients. *Brain and Language*, 27, 310–321. doi:10.1016/0093-934X(86)90022-2
- Burgess, C., & Simpson, G. B. (1988). Cerebral hemispheric mechanisms in the retrieval of ambiguous word meanings. *Brain and Language*, 33(1), 86–103. doi:10.1016/0093-934X(88)90056-9
- Campbell, J. D., & Katz, A. N. (2012). Are there necessary conditions for inducing a sense of sarcastic irony? *Discourse Processes*, 49, 459–480. doi:10.1080/0163853X.2012.687863
- Coulson, S., Federmeier, K. D., Van Petten, C., & Kutas, M. (2005). Right hemisphere sensitivity to word- and sentence-level context: Evidence from event-related brain potentials. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 31(1), 129–147.
- Coulson, S., & Williams, R. W. (2005). Hemispheric asymmetries and joke comprehension. *Neuropsychologia*, 43, 128–141. doi:10.1016/j.neuropsychologia.2004.03.015
- Eviatar, Z., & Just, M. A. (2006). Brain correlates of discourse processing: An fMRI investigation of irony and metaphor comprehension. *Neuropsychologia*, 44, 2348–2359. doi:10.1016/j.neuropsychologia.2006.05.007
- Faust, M., & Chiarello, C. (1998). Sentence context and lexical ambiguity resolution by the two hemispheres. *Neuropsychologia*, 36, 827–835. doi:10.1016/S0028-3932(98)00042-6
- Faust, M. E., & Gernsbacher, M. A. (1996). Cerebral mechanisms for suppression of inappropriate information during sentence comprehension. *Brain and Language*, 53(2), 234–259. doi:10.1006/brln.1996.0046
- Federmeier, K. D. (2007). Thinking ahead: The role and roots of prediction in language comprehension. *Psychophysiology*, 44(4), 491–505. doi:10.1111/psyp.2007.44.issue-4
- Federmeier, K. D., & Kutas, M. (1999). Right words and left words: Electrophysiological evidence for hemispheric differences in meaning processing. *Cognitive Brain Research*, 8, 373–392. doi:10.1016/S0926-6410(99)00036-1
- Filik, R., Howman, H., Ralph-Nearman, C., & Giora, R. (In progress). *The role of defaultness in sarcasm interpretation: Evidence from eye-tracking during reading. This volume.*
- Gibbs, R. W., Jr. (1994). *The poetics of mind: Figurative thought, language, and understanding*. New York, NY: Cambridge University Press.
- Giora, R. (2003). *On our mind: Salience, context, and figurative language*. New York: Oxford University Press.
- Giora, R., Givoni, S., & Fein, O. (2015). Defaultness reigns: The case of sarcasm. *Metaphor and Symbol*, 30/4, 290–313. doi:10.1080/10926488.2015.1074804
- Giora, R., Zaidel, E., Soroker, N., Batori, G., & Kasher, A. (2000). Differential effect of right and left hemispheric damage on understanding sarcasm and metaphor. *Metaphor and Symbol*, 15, 63–83. doi:10.1080/10926488.2000.9678865
- Jung-Beeman, M. (2005). Bilateral brain processes for comprehending natural language. *Trends in Cognitive Sciences*, 9 (11), 512–518. doi:10.1016/j.tics.2005.09.009
- Kacirik, N. A., & Chiarello, C. (2007). Understanding metaphoric language: Is the right hemisphere uniquely involved? *Brain and Language*, 100, 188–207. doi:10.1016/j.bandl.2005.10.010
- Loftus, G. R., & Masson, M. E. J. (1994). Using confidence intervals in within subject designs. *Psychonomic Bulletin & Review*, 1/4, 476–490. doi:10.3758/BF03210951
- Mashal, N., & Faust, M. (2009). Conventionalization of novel metaphors: A shift in hemispheric asymmetry. *Laterality*, 14(6), 573–589. doi:10.1080/13576500902734645
- Mashal, N., Faust, M., & Hendler, T. (2005). The role of the right hemisphere in processing nonsalient metaphorical meanings: Application of principal components analysis to fMRI data. *Neuropsychologia*, 43(14), 2084–2100. doi:10.1016/j.neuropsychologia.2005.03.019
- Newman, S. D., Just, M. A., & Mason, R. A. (2003). Understanding text with the right side of the brain: What functional neuroimaging has to say. In L. M. B. Tomitch & C. Rodrigues (Eds.), *Ensaio sobre a linguagem e o cérebro humano: Contribuições multidisciplinares* (pp. 71–84). Porto Alegre, Brazil: Artmed.
- Oldfield, R. C. (1971). The assessment and analysis of handedness: The Edinburgh inventory. *Neuropsychologia*, 9, 97–113. doi:10.1016/0028-3932(71)90067-4
- Partington, A. (2011). Phrasal irony: Its form, function and exploitation. *Journal of Pragmatics*, 43, 1786–1800. doi:10.1016/j.pragma.2010.11.001
- Peleg, O., & Eviatar, Z. (2008). Hemispheric sensitivities to lexical and contextual constraints: Evidence from ambiguity resolution. *Brain and Language*, 105(2), 71–82. doi:10.1016/j.bandl.2007.09.004
- Peleg, O., & Eviatar, Z. (2009). Semantic asymmetries are modulated by phonological asymmetries: Evidence from the disambiguation of heterophonic versus homophonic homographs. *Brain and Cognition*, 70, 154–162. doi:10.1016/j.bandc.2009.01.007
- Peleg, O., & Eviatar, Z. (2017). Controlled semantic processes within and between the two cerebral hemispheres. *Laterality*, 22(1), 1–16. doi:10.1080/1357650X.2015.1092547
- Peleg, O., Markus, A., & Eviatar, Z. (2012). Hemispheric asymmetries in meaning selection: Evidence from the disambiguation of homophonic vs. heterophonic homographs. *Brain and Cognition*, 80, 328–337. doi:10.1016/j.bandc.2012.08.005

- Sulis, E., Hernandez Farias, D. I., Rosso, P., Patti, V., & Ruffo, G. (2016). Figurative messages and affect in Twitter: Differences between #irony, #sarcasm and #not. *Knowledge-Based Systems*, 132–143. doi:[10.1016/j.knosys.2016.05.035](https://doi.org/10.1016/j.knosys.2016.05.035)
- Titone, D. A. (1998). Hemispheric differences in context sensitivity during lexical ambiguity resolution. *Brain and Language*, 65, 361–394. doi:[10.1006/brln.1998.1998](https://doi.org/10.1006/brln.1998.1998)
- Zaidel, E. (1979). Performance on the ITPA following cerebral commissurotomy and hemispherectomy. *Neuropsychologia*, 17, 259–280. doi:[10.1016/0028-3932\(79\)90073-3](https://doi.org/10.1016/0028-3932(79)90073-3)
- Ziv, Y. (2013). *Staam*: Maintaining consistency in discourse. In M. Florentin (Ed.), *Collection of articles on language* (pp. 151–159). Jerusalem, Israel: Hebrew Academy (In Hebrew).