

20 Salience and Defaultness

20.1 The Graded Salience Hypothesis: Re-defining Salience in Terms of Defaultness

The Graded Salience Hypothesis, introducing the notion of Salience (cf. Giora 1997; 2003), has been recently reviewed and reframed as one of the modules of the Defaultness Hypothesis (cf. Giora/Givoni/Fein 2015a). Whereas the Graded Salience Hypothesis focuses on default, coded and salient *meanings*, the Defaultness Hypothesis also acknowledges default, even if nonsalient, constructed *interpretations*. In both cases, however, defaultness is defined in terms of an automatic response to a stimulus. Given their automaticity, default responses will be evoked unconditionally, initially and directly, regardless of degree of nonliteralness (literal vs nonliteral), contextual support (weak vs strong), negation (negation vs affirmation), and, with regards to interpretations, also novelty (high or less-high). As such, default responses are expected to supersede non-default counterparts.

In this chapter, we focus on default, salient *meanings* while distinguishing them from nondefault, less salient and nonsalient counterparts. We then also discuss the role default meanings play in constructing salience-based often default interpretations (referred to in the literature as ›compositional meanings‹) and weigh them against nondefault nonsalient counterparts. Finally, we present nonsalient yet default *interpretations*.

According to the Graded Salience Hypothesis, now reframed in terms of the Defaultness Hypothesis, default meanings are *salient* meanings – meanings listed in the mental lexicon, ranking high on prominence due to cognitive factors (such as prototypicality, stereotypicality, or individual relevance) as well as usage-based factors (such as degree of exposure, including individual experiential familiarity, conventionality, or frequency of a stimulus); a case in point would be the ›plant‹ meaning of ›tree‹. Still, meanings listed in the mental lexicon, yet ranking lower on prominence, are *less-salient* and hence lower on defaultness; a case in point would be the ›syntactic‹ or ›diagram‹ meaning of ›tree‹. Meanings not listed in the mental lexicon are *nondefault*, *nonsalient*; a case in point would be the brute meaning of ›Yahoo‹, introduced by Jonathan Swift's (1726) *Gulliver's Travels*, which escapes most of today's mental lexicons. To establish meanings' degree of defaultness, their degree of accessibility should be established when presented in isolation. Similarly, to

establish interpretations' degree of defaultness, their degree of accessibility should be established when presented in isolation.

Note, however, that degree of salience is not a constant; rather, it is pliable and receptive to modifications and variabilities, gaining prevalence in a certain community of speakers at a certain time and place, depending, for example, on frequency of use or experiential familiarity; a case in point would be the default salient meaning of ›thick‹, which was primarily used to describe a physical aspect (›dense‹, ›massive‹), but is now denoting a mental aspect (›stupid‹), especially when referring to humans. Degree of salience, then, is not stable but volatile, and should be established empirically before being tested in experimentations on language use. Crucially, though, degree of salience is not dependent on immediate linguistic (or nonlinguistic) context (in so far as it does not involve lexical priming).

Once the automaticity or temporal priority of meanings is established, the Defaultness Hypothesis predicts that default, salient meanings of words and collocations will be activated instantly upon encountering the relevant stimulus, irrespective of factors known to affect processing, such as strength of contextual support, degree of negation, or degree of non/literalness. Less- and nondefault meanings will lag behind and will often depend on contextual information, including explicit cueing (cf. Givoni/Giora/Bergerbest 2013), to allow their activation or construal. Indeed, as predicted, default salient meanings have been shown to be evoked unconditionally, initially and directly, faster than non-default counterparts (cf. Giora 1997; 2003).

According to the Graded Salience Hypothesis, default interpretations are salience-based. Contrary to idioms, which often have both salience-based, here nondefault literal interpretations, and default, idiomatic meanings (the latter is coded and processed faster than the former), when dealing with an utterance's possible interpretations, it's crucial to distinguish between two levels of interpretation. One is *salience-based* interpretation, which relies on the default salient meanings of the words that make up the stimulus; a case in point would be the default, salience-based (here, literal) interpretation (›orderly‹) of an infrequent affirmative utterance – *He is the most organized student*.

A nondefault, nonsalient interpretation of such an utterance will be its sarcastic interpretation (related to ›messy‹), derivable when embedded in a sarcastically biased context. However, in addition to default, salience-based interpretations, some utterances may have a default yet nonsalient interpretation. Even in such

cases, a default interpretation is predicted to be activated faster than a nondefault counterpart; a case in point would be the default (here, sarcastic) interpretation (›messy‹) of an infrequent negative utterance – *He is not the most organized student*; a nondefault interpretation of such an utterance would be its literal interpretation (›he is organized but not the most organized‹), derivable when embedded in a literally biased context (indicating that others are more organized).

Once the automaticity or temporal priority of interpretations is established, the Defaultness Hypothesis predicts that default interpretations of utterances will be activated instantly upon encountering the relevant stimulus, irrespective of all factors known to affect processing, such as strength of contextual support, degree of negation, non/literalness, and even novelty. Nondefault interpretations will lag behind and will often depend on contextual information, often including explicit cueing (cf. Becker/Giora, in press) to allow their activation or construal. Indeed, as predicted, default interpretations have been shown to be evoked unconditionally, initially and directly, faster than nondefault counterparts (cf. Giora/Givoni/Fein 2015a).

Below we review the studies testing the predicted superiority of default meanings (section 2), default salience-based interpretations (section 3), and default yet nonsalient interpretations (section 4) over nondefault counterparts, in equally strongly supportive contexts.

20.2 Default salient meanings

The processing speed superiority of default, salient meanings

Contrary to the Literal-first model (Grice 1975), the Graded Salience Hypothesis, reframed now in terms of the Defaultness Hypothesis, predicts that defaultness, rather than literalness, will be the determining factor affecting processing speed. Indeed, Giora and Fein (1999a) found that participants completed as many fragmented words related to an idiomatic meaning of familiar idioms embedded in literally biasing contexts as they completed words related to their salience-based interpretation. Marlies E. C. Van de Voort and Wietske Vonk (1995) found that familiar idioms were automatically processed idiomatically, and Raymond W. Gibbs (1980) reported that idioms such as *spill the beans* were read faster in an idiomatically than in a literally biasing context. These findings support the view that the idiomatic meanings of familiar idioms are default, salient meanings

(cf. also Libben/Titone 2008). Similarly, Giora and Fein (1999b) found that the sarcastic meaning (›annoying‹) of familiar sarcasm (*Very funny*) was accessible even in contextually incompatible, literally biasing contexts (›amusing‹). Default meanings, then, are accessed automatically and directly, irrespective of contextual information.

What about less-salient meanings? John N. Williams (1992) examined whether the various, literal and metaphorical meanings of polysemous adjectives (*firm*) are functionally independent during processing. He showed that, for ›central‹ (i. e., default, salient) meanings of words (here, the literal meaning ›solid‹), effects were visible as long as 1100 msec following prime onset, even when embedded in contexts *irrelevant* to the salient, ›central‹ meaning. However, no significant priming of targets related to ›non-central‹ (i. e., nondefault, less-salient) meanings (here, the metaphorical ›strict‹ meaning of *firm*) was found in contexts *irrelevant* to the ›non-central‹ meaning.

This happened despite the fact that both types of targets were equally primed when the prime was presented in isolation. Accounting for these results in terms of the Graded Salience and the Defaultness Hypotheses suggests that, while the default, salient meaning of a polysemous concept is always accessed immediately, the nondefault, less-salient meaning is only accessed if explicitly invited, as shown by Givoni et al. (2013), or if the salient meaning fails to meet contextual fit (but cf. Frisson/Pickering 2001 for a different view). (For results from divided visual field and eye-tracking paradigms relating to default (›money‹) and nondefault (›river‹) meanings of homonyms (e. g. *bank*); cf. e. g. Peleg/Eviatar 2008; Frazier/Rayner 1990).

Contrary to the Direct Access View (Gibbs 1994), the Defaultness Hypothesis predicts that defaultness, rather than context, will be the decisive factor with regard to processing speed. Peleg, Giora and Fein (2001) demonstrated that activation of contextually appropriate meanings is not selective. Rather, default, salient, even if inappropriate meanings are activated as well, in spite of contextual misfit.

In Orna Peleg et al. (2001), participants had to make lexical decisions to probes which were related either to the default, salient although contextually inappropriate meaning (›criminals‹) or to the nondefault, nonsalient but contextually appropriate meanings (›kids‹) of the targets (*delinquents*). The probes were displayed immediately following the offset of the targets, which were presented either at the beginning (1) or at the end of the final sentence (2):

- (1) Sarit's sons and mine went on fighting continuously. Sarit said to me: »*These delinquents* won't let us have a moment of peace.*« (Probes displayed at *: Salient – criminals; contextually compatible – kids; Unrelated – painters)
- (2) Sarit's sons and mine went on fighting continuously. Sarit said to me: »*A moment of peace won't let us have these delinquents*.*« (Probes displayed at *: Salient – criminals; contextually compatible – kids; Unrelated – painters)

Results support the Graded Salience and the Defaultness Hypotheses (the former being a part of the latter). They testify to the accessibility of salient meanings in both sentence initial and final positions. Although contextual information availed the appropriate meaning in sentence initial as well as in sentence final position, it did not inhibit salient though inappropriate meanings. Even when they were slower to activate than nonsalient contextually appropriate meanings (as in sentence final position), they were still accessible.

Such findings demonstrate that, as predicted by the Graded Salience and the Defaultness Hypotheses, in addition to contextual mechanisms, modular, bottom-up mechanisms are at work as well: contextual mechanisms do not interact initially with lexical processes and, therefore, do not block default, salient yet inappropriate meanings.

Similarly, Colston and Gibbs (2002) found that default, salient metaphorical meanings of an utterance's constituents (*sharp*) took less time to read in metaphorically (3) than in sarcastically biasing contexts (4). Their results attest to the superiority of defaultness (here, the salient, metaphorical meanings) over nondefaultness (here, context-based sarcastic interpretations):

- (3) You are a teacher at an elementary school. You are discussing a new student with your assistant teacher. The student did extremely well on her entrance examinations. You say to your assistant, »*This one's really sharp.*«
- (4) You are a teacher at an elementary school. You are gathering teaching supplies with your assistant teacher. Some of the scissors you have are in really bad shape. You find one pair that won't cut anything. You say to your assistant, »*This one's really sharp.*«

Contrary to the Suppression Hypothesis (cf. Gernsbacher 1990), results from several studies show that

even negation does not block default, salient meanings initially (e. g. Macdonald/Just 1989; for a review cf. Giora 2006). Using a lexical-decision task, Uri Hasson and Sam Glucksberg (2006) found that »fast« was facilitated at 150 ms and also at 500 ms ISI following both affirmative (*The train to Boston was a rocket*) and the negative (*The train to Boston was no rocket*) metaphorical utterances (cf. also Kaup 2001 and Becker 2016).

Moreover, and contrary to the view that negation slows down processing (e. g. Clark/Clark 1977; Mayo/Schul/Burnstein 2004), default, salient idiomatic negatives (*the apple doesn't fall far from the tree*) are faster to read, despite being longer, compared to their non-default, noncoded, affirmative counterparts (*the apple falls far from the tree*), when embedded in identical neutral contexts (cf. Giora/Meytes/Tamir et al. 2017b).

The role of default salient meanings in affecting pleasure

Still, while salience facilitates processing, it interferes with deriving nonsalience, slowing it down in the process. Often, though, it affects pleasure, especially when, following its automatic activation, it is retained and interacts with nonsalience (cf. Giora/Fein/Kronrod et al. 2004; Giora/Fein/Kotler et al. 2015c; Giora/Givoni/Heruti et al. 2017a).

According to the Graded Salience and the Defaultness Hypotheses but contrary to the Literal First Model, pleasing effects are not related to nonliteralness. Instead, they are an end-product of optimal innovations, the latter involving deautomatization of either default salient meanings or default interpretations. Giora et al. (2004) discuss nondefault, nonsalient responses (*A peace of paper*) invoking default salient alternatives (*A piece of paper*), allowing their similarities and differences to be assessable. Indeed, in Giora et al. (2004), participants rated optimal innovations (*A peace of paper*) as more pleasing than familiar expressions (*A piece of paper*) whose salient meanings are listed in the mental lexicon, and more pleasing than noncoded nonsalient pure innovations (*A pill of pepper*).

How can we tell that an optimal innovation activates the salient, default response while deautomatizing it? Giora et al.'s (2004) study shows that familiar stimuli, whose meaning is salient, take less time to process following optimal innovations relative to pure innovations, indicating that optimal innova-

tions activate and prime their default salient meanings. (On default salience-based literals activating their default salient metaphorical meanings; cf. Giora et al. 2015c).

Giora, Gazal, Goldstein et al. (2012) tested these predictions with adults with Asperger's syndrome (AS). They found that like controls, AS adults were more sensitive to degree of defaultness/salience than to degree of literalness. Materials (e. g. familiar metaphors: *flower bed*; novel metaphors: *dying star*; familiar literals: *wooden table*; and novel literals: *Tverian horse*) were presented in and out of context. Both groups took longer to read optimally innovative items compared to familiar ones, regardless of degree of non/literalness. Innovation, then, comes with a processing cost, yet it does not distinguish literal from nonliteral innovation. (For similar patterns of behavior from neurological studies; cf. also Gold/Faust 2012; Colich/Wang/Rudie et al. 2012).

20.3 Default salience-based interpretations: The processing speed superiority of default, salience-based interpretations

Default salience-based interpretations are the product of the processing mechanisms of utterances' non-coded interpretations which are compositional, relying heavily on the default, coded meanings of the utterances' components. For example, processing *she's radiant* involves accessing the coded, default ›happy‹ meaning of *radiant*. The Graded Salience Hypothesis predicts that interpretations based on the default, coded, salient meanings of their utterances' components will be processed faster than nondefault, non-coded, nonsalient interpretations which rely on contextual information for their derivation (e. g. *she's radiant*, said ironically of someone who is unhappy).

Contrary to the expectation hypothesis (e. g. Gibbs 2002; Ivanko/Pexman 2003), the Graded Salience Hypothesis predicts that it is not contextual expectation but degree of salience that will be the decisive factor during the early stages of sarcasm comprehension. Giora, Fein, Laadan et al. (2007) indeed found that the introduction of a sarcastic speaker into the discourse (see the 6th turn in example (5–6) below, in bold for convenience) was ineffective, resulting in longer reading times for ironic (5) than for literal (6) endings, the former involving their automatically retrievable default, salience-based interpretations.

- (5) Barak: *I finish work early today.*
 Sagit: *So, do you want to go to the movies?*
 Barak: *I don't really feel like seeing a movie.*
 Sagit: *So maybe we could go dancing?*
 Barak: *No, at the end of the night my feet will hurt and I'll be tired.*
 Sagit: **You're a really active guy ...**
 Barak: *Sorry, but I had a rough week.*
 Sagit: *So what are you going to do tonight?*
 Barak: *I think I'll stay home, read a magazine, and go to bed early.*
 Sagit: *Sounds like you are going to have a really interesting evening.*
 Barak: *So we'll talk sometime this week.*
- (6) Barak: *I was invited to a film and a lecture by Amos Gitai.*
 Sagit: *That's fun. He is my favorite director.*
 Barak: *I know, I thought we'll go together.*
 Sagit: *Great. When is it on?*
 Barak: *Tomorrow. We will have to be in Metulla* in the afternoon.*
 Sagit: **I see they found a place that is really close to the center.**
 Barak: *I want to leave early in the morning. Do you want to come?*
 Sagit: *I can't, I'm studying in the morning.*
 Barak: *Well, I'm going anyway.*
 Sagit: *Sounds like you are going to have a really interesting evening.*
 Barak: *So we'll talk sometime this week.*

[* Note that Metulla is far from the center]

Fein, Yeari, and Giora (2015) used the same stimuli used in Giora et al. (2007), only this time, sarcasm cues (e. g., mocking/winking) were strengthened even further, indicating the way in which the sarcastic speakers expressed themselves. Still, strengthening the cues did not affect the pattern of results. Default, salience-based targets were faster to read than equally highly expected, nondefault, nonsalient, sarcastically biased alternatives.

Giora, Fein, Kaufman et al. (2009) further measured reading times of statements (*This sure is an exciting life!*) following contexts featuring a frustrated, a realized, or no expectation. They found that contextual information, biasing targets toward their nondefault, nonsalient ironic interpretation, did not facilitate irony comprehension; while not differing from each other, both frustrated and realized expectations took significantly longer to read than default, literally/salience-based biased targets, which followed contexts featuring

no expectation. These results support the view that it is not strong contextual information (featuring an expectation for an ironic utterance) that plays an initial and crucial role in making sense of stimuli. Rather, it is degree of defaultness (here, default salience-based interpretation) that accounts for stimuli's processing speed or lack of it (i. e., when contextual fit requires rejecting the default, salience-based interpretations and activating nondefault nonsalient counterparts).

20.4 Default nonsalient interpretations

The speed superiority of default, nonsalient interpretations over nondefault, nonsalient counterparts

It is worth noting here that the Graded Salience Hypothesis considers all nonsalient interpretations ›non-default‹, and, as such, more difficult to process than default salience-based alternatives. But, in fact, there are a great number of constructions conveying default interpretation despite being nonsalient.

Given that the Graded Salience Hypothesis is now incorporated into the Defaultness Hypothesis (cf. Giora et al. 2015a), which accommodates it but also extends it, the focus on default yet nonsalient interpretations is in order here. In fact, the Defaultness Hypothesis is the only processing model that predicts that some non-coded nonsalient, i. e., novel constructions, will convey a non-literal interpretation *by default*. Indeed Giora, Livnat, Fein et al. (2013) found that negative constructions, controlled for novelty, such as *this is not a court of law* were rated as metaphorical (meaning ›stop arguing with each other‹) when presented outside of context; their nondefault nonpreferred alternative was rated as literal (meaning ›this is not where suspects appear before a judge‹). When embedded in strongly supportive contexts, nonsalient yet default negative constructions were read faster when biased toward the metaphorical than toward the equally strongly supported literal interpretation.

Similarly, Giora/Drucker et al. (2015b) studied novel negative constructions (*Intelligence is not her forte*; *Intelligence is not his prominent feature*) shown to convey a default albeit nonsalient, sarcastic interpretation (meaning ›s/he is stupid‹) when outside of context. When embedded in equally strong contexts, biasing them either toward the sarcastic or toward the literal interpretation, they were read faster in the sarcastically than in the literally biasing settings. The literally biasing contexts did not block default, yet nonsalient

sarcastic interpretations, rendering the literally biased targets slower to process.

Furthermore, Giora et al. (2015a) showed that nonsalient, novel negative constructions (*He's not the most organized student*), conveying a sarcastic interpretation by default (established as such when outside of context), were even read faster than their nondefault nonsalient *affirmative* sarcastic counterparts (*He's the most organized student*), involving their default literal interpretation in the process (the latter established as such when outside of context). Note, however, that both targets were embedded in equally strong contexts, biasing them toward the sarcastic interpretation. Still, processing affirmative counterparts in a sarcastic context does not block their automatically retrievable default, salience-based interpretation, which had to be rejected (albeit not discarded) in favor of a nondefault sarcastic alternative that meets contextual fit.

These findings are unprecedented. They demonstrate the superiority of defaultness, superseding degree of negation, degree of novelty/nonsalience, and degree of strength of contextual support. It's for the first time that novel, nonsalient, negative utterances are processed faster than their nonsalient affirmative counterparts despite both being equally strongly sarcastic. (For similar evidence from eye-tracking, collected from English items; cf. Filik/Howman/Ralph-Nearman et al. 2018; for the speed superiority in the cerebral hemispheres of default nonsalient negative sarcasm over nonsalient yet nondefault affirmative sarcasm; cf. Giora/Cholev/Fein et al. 2018).

The role of default, salience-based interpretations in affecting pleasure: The case of default and nondefault sarcasm

Recall that the Defaultness Hypothesis predicts that optimal innovation will result in greater pleurability as it involves the deautomatization of either default salient meanings or default salience-based interpretations. Giora et al. (2017a) tested the prediction with regard to interpretations, while using Giora et al.'s (2015a) sarcastic items, embedded either in linguistic or pictorial contexts. They expected that the default negative version of their constructions will not be pleasing, whereas the nondefault affirmative counterparts will. The rationale behind this relates to the fact that only the latter is a candidate for optimal Innovation, since it involves a default salience-based literal interpretation in the process. Given the automatic activation of default interpretations, this activation al-



Fig. 20.1 Pictorial stimuli presented in Giora et al. (2017a)

lows both, the nondefault sarcastic interpretation and the default salience-based interpretations to interact and be weighed against each other. Results indeed showed that when embedded in sarcastically biasing contexts, the affirmative items were rated as pleasurable; their negative counterparts were not perceived as pleasing. This was true of both linguistic as well as pictorial contexts (for an example of the latter, see fig. 20.1). It is deautomatized defaultness, then, that affects pleasure.

20.5 Conclusions

In this chapter we discussed the theoretical and empirical aspects of salience and defaultness and the way in which they interact. The debate in psycholinguistics and cognitive science, questioning whether contextual information affects initial processing, has continued for decades. Salience and defaultness provide key concepts in making sense of the data. The Graded Salience Hypothesis (cf. Giora 1997; 2003) and the Defaultness Hypothesis (cf. Giora et al. 2015a) attest to what is on our mind. They are able to predict when the speakers' accumulated, usage-based knowledge will kick in, context notwithstanding. These theories account for the different levels of processing (of meanings and interpretations) by assuming a modular view of the mind that expects encapsulation of certain cognitive processes (i. e., the lexicon) yet does not take a strict serial

view (à la Fodor 1983). Crucially, even though top-down mechanisms are always at work, they cannot put a default response down. The studies described here uncover the underlying mechanisms of understanding. They show that, to appreciate and make sense of the new, we first contemplate the default, whether a coded, lexicalized meaning or a novel, constructed interpretation. It is defaultness, then, that reigns.

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

Literatur

- Becker, Israela (2016): The Negation Operator is not a Suppressor of the Concept in its Scope. In fact, quite the Opposite. Unpublished MA thesis. Tel-Aviv University.
- Becker, Israela/Giora, Rachel (in press): The Defaultness Hypothesis: A quantitative corpus-based study of non/default sarcasm and literalness production. In: *Journal of pragmatics*.
- Clark, Herbert H./Clark, Eve V. (1977): *Psychology and Language: An Introduction to Psycholinguistics*. New York.
- Colich, Natalie L./Wang, Audrey-Ting/Rudie, Jeffrey D./Hernandez, Leanna M./Bookheimer, Susan Y./Dapretto, Mirella (2012): Atypical neural processing of ironic and sincere remarks in children and adolescents with autism spectrum disorders. In: *Metaphor and Symbol* 27(1), 70–92.
- Colston, Herbert L./Gibbs, Raymond W. (2002): Are irony and metaphor understood differently? In: *Metaphor and Symbol* 17, 57–60.
- Fein, Ofer/Yeari, Menahem/Giora, Rachel (2015): On the priority of salience-based interpretations: The case of irony. In: *Intercultural Pragmatics* 12(1), 1–32.
- Filik, Ruth/Howman, Hannah/Ralph-Nearman, Christina/Giora, Rachel (2018): The role of defaultness in sarcasm interpretation: Evidence from eye-tracking during reading. In: *Metaphor and Symbol* 33(3), 148–162.
- Frazier, Lyn/Rayner, Keith (1990): Taking on semantic commitments: Processing multiple meanings vs. multiple senses. In: *Journal of Memory and Language* 29, 181–200.
- Frisson, Steven/Pickering, Martin J. (2001): Obtaining a figurative interpretation of a word: Support for underspecification. In: *Metaphor and Symbol* 16(3&4), 149–171.
- Fodor, Jerry (1983): *The Modularity of Mind*. Cambridge, Mass.
- Gernsbacher, Morton Ann (1990): *Language Comprehension as Structure Building*. Lawrence Erlbaum Associates, Hillsdale, NJ.
- Gibbs, Raymond W. (1980): Spilling the beans on understanding and memory for idioms in conversation. In: *Memory/Cognition* 8, 149–156.
- Gibbs, Raymond W. (1994): *The Poetics of Mind*. Cambridge.
- Gibbs, Raymond W. (2002): A new look at literal meaning in

- understanding what is said and implicated. In: *Journal of Pragmatics* 34, 457–486.
- Giora, Rachel (1997): Understanding figurative and literal language: The graded salience hypothesis. In: *Cognitive Linguistics* 8/3, 183–206.
- Giora, Rachel (2003): *On our Mind: Salience, Context, and Figurative Language*. New York.
- Giora, Rachel (2006): Anything negatives can do affirmatives can do just as well, except for some metaphors. In: *Journal of Pragmatics* 38, 981–1014.
- Giora, Rachel/Cholev, Adi/Fein, Ofer/Peleg, Orna (2018): On the speed superiority of defaultness – a divided visual field study. In: *Metaphor and Symbol* 33(3), 163–174.
- Giora, Rachel/Drucker, Ari/Fein, Ofer/Mendelson, Itamar (2015b): Default sarcastic interpretations: On the priority of nonsalient interpretations. In: *Discourse Processes* 52(3), 173–200.
- Giora, Rachel/Fein, Ofer (1999a): On understanding familiar and less-familiar figurative language. In: *Journal of Pragmatics* 31, 1601–1618.
- Giora, Rachel/Fein, Ofer (1999b): Irony: context and salience. In: *Metaphor and Symbol* 14, 241–257.
- Giora, Rachel/Fein, Ofer/Kaufman, Ronie/Eisenberg, Dana/Erez, Shani (2009): Does an »ironic situation« favor an ironic interpretation? In: Geert Brône/Jeroen Vandaele (ed.): *Cognitive Poetics. Goals, Gains and Gaps (Applications of Cognitive Linguistics series)*. Berlin/New York, 383–399.
- Giora, Rachel/Fein, Ofer/Kotler, Nurit/Shuval, Noa (2015c): Know hope: Metaphor, optimal innovation, and pleasure. In: Geert Brône/Kurt Feytaerts/Tony Veale (ed.): *Cognitive Linguistics Meet Humor Research. Current Trends and New Developments*. Berlin/New York, 129–146.
- Giora, Rachel/Fein, Ofer/Kronrod, Ann/Elnatan, Idit/Shuval, Noa/Zur, Adi (2004): Weapons of mass distraction: Optimal innovation and pleasure ratings. In: *Metaphor and Symbol* 19, 115–141.
- Giora, Rachel/Fein, Ofer/Laadan, Dafna/Wolfson, Joe/Zeitun, Michal/Kidron, Ran/Kaufman, Ronie/Shaham, Ronit (2007): Expecting irony: Context vs. salience-based effects. In: *Metaphor and Symbol* 22, 119–146.
- Giora, Rachel/Gazal, Oshrat/Goldstein, Idit/Fein, Ofer/Stringaris, K. Argyris (2012): Salience and context: Interpretation of metaphorical and literal language by young adults diagnosed with Asperger's syndrome. In: *Metaphor and Symbol* 27, 22–54.
- Giora, Rachel/Givoni, Shir/Fein, Ofer (2015a): Defaultness reigns: The case of sarcasm. In: *Metaphor and Symbol* 30/4, 290–313.
- Giora, Rachel/Givoni, Shir/Heruti, Vered/Fein, Ofer (2017a): The role of defaultness in affecting pleasure: The optimal innovation hypothesis revisited. In: *Metaphor and Symbol* 32/1, 1–18.
- Giora, Rachel/Livnat, Elad/Fein, Ofer/Barnea, Anat/Zelman, Rakefet/Berger, Iddo (2013): Negation generates nonliteral interpretations by default. In: *Metaphor and Symbol* 28, 89–115.
- Giora, Rachel/Meytes, Dalia/Tamir, Ariella/Givoni, Shir/Heruti, Vered/Fein, Ofer (2017b): Defaultness shines while affirmation pales. In: Angeliki Athanasiadou/Herbert Colston (ed.): *Irony in Language, Use and Communication*. FTL Series. Amsterdam, 219–236.
- Givoni, Shir/Giora, Rachel/Bergerbest, Dafna (2013): How speakers alert addressees to multiple meanings. In: *Journal of Pragmatics* 48(1), 29–40.
- Gold, Rinat/Faust, Miriam (2012): Metaphor comprehension in persons with Asperger's syndrome: Systemized versus non-systemized semantic processing. In: *Metaphor and Symbol* 27(1), 55–69.
- Grice, Paul H. (1975): Logic and conversation. In: Peter Cole/Jerry Morgan (ed.): *Speech acts*. In: *Syntax and semantics* 3. New York, 41–58.
- Hasson, Uri/Gluckberg, Sam (2006): Does understanding negation entail affirmation? An examination of negated metaphors. In: *Journal of Pragmatics* 38, 1015–1032.
- Ivanko, Stacey L./Pexman, Penny M. (2003): Context incongruity and irony processing. In: *Discourse Processes* 35, 241–279.
- Kaup, Barbara (2001): Negation and its impact on the accessibility of text information. In: *Memory/Cognition* 29, 960–967.
- Libben, Maya R./Titone, Debra (2008): The multidetermined nature of idiom processing. In: *Memory/Cognition* 36/6, 1103–1121.
- MacDonald, Maryellen C./Just, Marcel A. (1989): Changes in activation levels with negation. In: *Journal of Experimental Psychology: Learning, Memory, and Cognition* 15, 633–642.
- Mayo, Ruth/Schul, Yaacov/Burnstein, Eugene (2004): »I am not guilty« vs »I am innocent«: Successful negation may depend on the schema used for its encoding. In: *Journal of Experimental Social Psychology* 40, 433–449.
- Peleg, Orna/Eviatar, Zohar (2008): Hemispheric sensitivities to lexical and contextual constraints: Evidence from ambiguity resolution. In: *Brain and Language* 105(2), 71–82.
- Peleg, Orna/Giora, Rachel/Fein, Ofer (2001): Salience and context effects: Two are better than one. In: *Metaphor and Symbol* 16, 173–192.
- Swift, Jonathan (1726): *Gulliver's Travels*. Ireland.
- Van de Voort, Marlies E. C./Vonk, Wietske (1995): You don't die immediately when you kick an empty bucket: A processing view on semantic and syntactic characteristics of idioms. In: Martin Everaert/Erik-Jan van der Linden/Andre Schenk et al. (ed.): *Idioms: Structural and Psychological Perspectives*. Hillsdale.
- Williams, John N. (1992): Processing polysemous words in context: Evidence from interrelated meanings. In: *Journal of Psycholinguistic Research* 21, 193–218.

Shir Givoni / Rachel Giora