

Visual Argument Structure Tool (VAST) Version 1.0

Daniel Leising, Oliver Grenke, and Marcos Cramer
Technische Universität Dresden

We present the first version of the Visual Argument Structure Tool (VAST), which may be used for jointly visualizing the semantic, conceptual, empirical and reasoning relationships that constitute arguments. Its primary purpose is to promote exactness and comprehensiveness in systematic thinking. The system distinguishes between concepts and the words (“names”) that may be used to refer to them. It also distinguishes various ways in which concepts may be related to one another (causation, conceptual implication, prediction, transformation, reasoning), and all of these from beliefs as to whether something IS the case and/or OUGHT to be the case. Using these elements, the system allows for formalizations of narrative argument components at any level of vagueness vs. precision that is deemed possible and/or necessary. This latter feature may make the system particularly useful for attaining greater theoretical specificity in the humanities, and for bridging the gap between the humanities and the “harder” sciences. However, VAST may also be used outside of science, to capture argument structures in e.g., legal analyses, media reports, belief systems, and debates.

Keywords: Modelling, Formalization, Narrative, Theory, Humanities, Science

Introduction

Argument structures are ubiquitous and important: We encounter them every day, in newspaper articles, in court rulings, in political and non-political debates on and off screen, and in our own more informal conversations with other people, and with ourselves. In an abstract sense, most arguments are about what is true, why and how things are related to one another, and about whether things are good or not. Too often, however, arguments seem to run in circles, end in stalemates, or just fizzle out and are given up upon, instead of being conclusively resolved. We argue (yes), that these things happen because people tend to lose sight of some of the claims that they or others have made before. Therefore, it is often advisable to aim for a comprehensive analysis of all relevant argument components. Another reason why so many arguments remain unproductive is that those who argue tend to overlook the actual complexity of their own and others’ claims, and the relative vagueness of many claims.

Interestingly, the same problems seem to plague much of the more “narrative” theorizing that is so common in the humanities, and in psychology. In fact, there has been no shortage of calls for better (e.g. more formalized) theorizing in psychology, precisely to counter these shortcomings (Devezer et al., 2021; Eronen & Bringmann, 2021; Fried, 2020; Glöckner & Betsch, 2011; Muthukrishna & Henrich, 2019; Robinaugh et al., 2021; Smaldino, 2017, 2019). Concrete advice on how

exactly such better theorizing may be achieved is largely missing, however (Borsboom et al., 2021).

In the present paper, we introduce a tool devised for dealing with those problems. The tool is called VAST (Visual Argument Structure Tool). The core idea is to visually display all relevant components of an argument structure at once, while at the same time aiming for exactness. A comprehensive display will make it harder to overlook or downplay related claims made earlier (e.g., because those previous claims do not align well with more recent ones). Visual displays may also be more intuitive and easier to digest for most users, especially when compared to the alternative of using algebraic expressions. After all, there is a reason why so many articles in scientific journals as well as in the news media are accompanied by figures illustrating their main points. Furthermore, visual displays tend to be more parsimonious: With formulae, the same variable name will have to be written again each time it is used as an input or output of some new equation. In contrast, a visual display may incorporate the variable only once, and establish all of the relevant relationships with other variables through arrows or lines pointing in various directions. This is how the matter is handled in VAST, and aligns well with the typical approaches in Structural Equation Modelling (SEM), Structural Causal Models (SCM) and Directed Acyclic Graphs (DAG) (Dablander, 2020; Pearl, 1995; Pearl & Mackenzie, 2018; Rohrer, 2018). The graphical structuring of arguments was also inspired to some extent by developments in formal ar-

Table 1. Components of the System

Element Name	Meaning	Symbolized by	Default Range
Concept	A feature that may apply to certain objects	Frame with abstract label	0; 1
Name	A natural-language label that may be used for a concept	Frame with label in quotation marks	0; 1
Higher-Order Concept	A specific combination of elements that may - in its entirety - apply to certain objects	Frame containing two or more elements	0; 1
Data	A set of actual observations	Frame with thick black edge on one side	0; 1
IS	How much X is the case	Pentagon containing IS in capitals	0; 1
OUGHT	How much X ought to be the case	Pentagon containing OUGHT in capitals	0; 1
Perspective	How much perspective-holder agrees with X	Oval connected to IS / OUGHT Containing name of perspective-holder	0; 1
Naming	How appropriate it is to call X by the name Y	Arrow accompanied by lowercase letter n	-1; 1
Conceptual Implication	How much thinking of something as being X also implies thinking of it as being Y	Arrow accompanied by lowercase letter i	-1; 1
Causation	How reliably X will trigger Y	Arrow accompanied by lowercase letter c	-1; 1
Transformation	How strongly X maps onto Y	Arrow accompanied by lowercase letter t	-1; 1
Prediction	How well Y may be predicted from X	Arrow accompanied by lowercase letter p	-1; 1
Reasoning	How much X is a reason to believe Y	Arrow accompanied by lowercase letter r	-1; 1

gumentation (see Baroni et al., 2018).

VAST overlaps significantly with all of these previous approaches, and also incorporates many elements of formal logic, in particular logical connectives (AND, OR and XOR) from classical propositional logic (Bünning & Lettmann, 1999) in the tradition of Boole (1854) and Frege (1879). Given that we allow truth-values between 0 and 1 (see below), VAST is also influenced by continuously-valued logics (Preparata & Yeh, 1972). However, VAST is comparatively broader and more integrative in that it explicitly accounts for various types of relationships between concepts (i.e., naming, conceptual implication, causation, prediction, transformation, and reasoning). The strength of all of these relationships may be expressed in terms of the same metric, as we will explain in more detail below. VAST also accounts for the possibility that concepts may be applied to different sets of objects, for claims as to whether something IS and/or OUGHT to be the case, and for different perspectives on these issues. We discuss some of the overlap and differences between VAST and previous tools with a similar scope further below.

The System

In this section, we introduce the different types of elements that, taken together, constitute our system in its entirety. Table 1 lists all of these elements alongside each other. To facilitate comprehension, we will use a variety of examples along the way to illustrate their potential uses.

Concepts

Concepts are the basic building blocks of cognition. Note that concepts are assumed to exist before language is being used (see below) — they may exert their influence irrespective of the words (“names”, see below) that are used to denote them. A concept assigns values to objects. Thus, concepts are very similar to mathematical functions in that they produce an output value for every input (i.e., object). In the most simple case, that output will be dichotomous, so the concept will yield a value of either 0 or 1 for each object. Here, 1 means that the object is an exemplar of the concept, and 0 means that the object is not an exemplar of the concept. For instance, a person may look at a number of objects and determine whether any of them are exemplars of the concept that one would refer to with the word “car”.

Note, however, that many concepts are not dichotomous but allow for continuous variation of output values. In VAST, this variation is usually normalized to a range between 0 and 1, to make comparisons between different concepts easier. This basically incorporates the so-called “prototype approach” to classification (Rosch, 1978), in which an object may be a more or less typical exemplar of a category / concept / class.

In VAST, concepts are usually displayed in the form of frames bearing abstract labels (e.g., X, Y, or XYZ). It is important to note that the display of a concept only symbolises the assumed existence and relevance of a cognitive process that would assign certain values to some objects and other values to other objects. It is agnostic in regard to the desired, assumed, or measured distribution of those values. The question of whether something is or should be the case is the realm of IS and OUGHT statements, which will be introduced later.

Empirical data constitute a special type of concept, which is symbolized by adding a thick black edge to one side of the respective frame. This is basically the same distinction that is made between “manifest” (measured) and “latent” (imagined) variables in Structural Equation Modelling (see Figure 12 and the accompanying text for an example).

Types of Relationships Between Concepts

The next major element of the system is the relationship that may exist between concepts. These are basically IF-THEN relations: IF one concept does apply (to some of the relevant objects), THEN another concept also applies, at least to some extent. There are different qualities of such relationships, however, which must be distinguished from another. In VAST, we explicitly account for naming, implication, causation, prediction, transformation, and reasoning relationships, which we consider to be among the most common and relevant ones. Specifying additional relationship types is also possible, if needed. All of this will now be explained in more detail.

Relationship Type 1: Naming (n).

People tend to use words to denote the different ways in which they think about objects. In VAST, these words are called “names”. To clearly distinguish concepts and their names from one another, VAST uses abstract labels (e.g., C, R) for the former and real language labels in quotation marks (e.g., “Cat”, “Rocket”) for the latter. The relationship between a concept and a word denoting that concept is called a “naming relationship”. It is symbolized by an arrow pointing from the concept to the name, accompanied by the lower letter n. As with all relationship types (see above), the arrow stands for

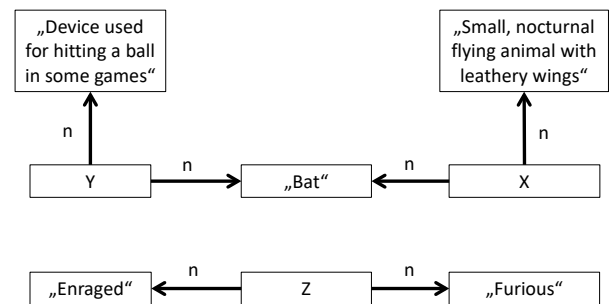


Figure 1

Selected naming relationships. The word “bat” is a homonym for concepts X and Y, whereas the words “furious” and “enraged” are synonyms for concept Z

an IF-THEN relation: IF an object is an exemplar of the respective concept, THEN one may call this object by the respective name.

Figure 1 displays some examples of naming relationships, including homonyms (the same name is used for different concepts) and synonyms (different names are used for the same concept). Note that both (a) the appropriateness and (b) the strengths of the relationships displayed in Figure 1 are treated as irrelevant for now.

Distinguishing between concepts and their names is often necessary, because idiosyncratic word usage accounts for all sorts of problems (e.g., misunderstandings) in everyday arguments. The same issue is relevant for psychology, which continues to suffer from — often unacknowledged — jingle-fallacies (use of homonymic theoretical concepts) and jangle-fallacies (use of synonymic theoretical concepts) (Block, 1995).

Relationship Type 2: Conceptual Implication (i).

Conceptual implication is about the extent to which classifying objects as exemplars of one concept implies also categorizing the same objects as exemplars of another concept. Figure 2 displays an example. Here, when an object is considered to be a “sun”, the same object is also somewhat likely to be considered “hot” and “bright”. Conceptual implications are symbolized by arrows accompanied by the lowercase letter i.

Figure 2 contains three different ways of displaying basically the same information. In the display on the left-hand side, all concepts and the relevant relationships between them are displayed as a part of one coherent whole. This is the default mode that we suggest for use with most VAST analyses, as it maximises parsimony while retaining all of the relevant information. In the middle display, the conceptual implications (type i)

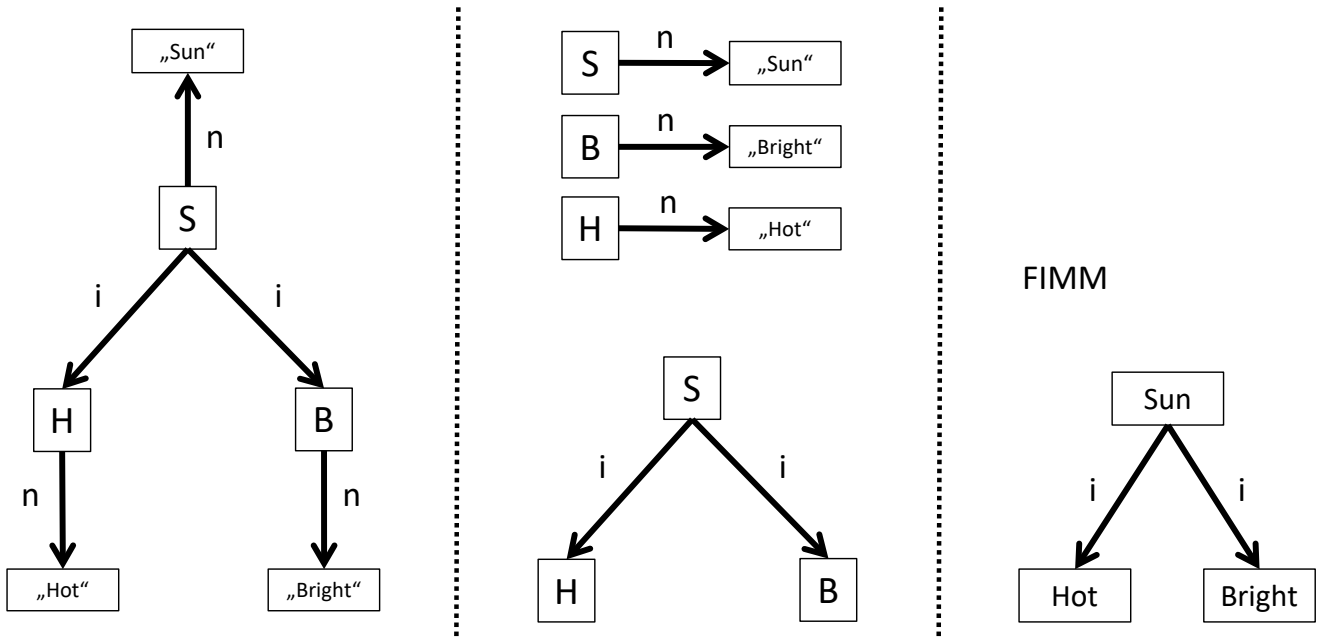


Figure 2

Conceptual implications among three concepts. Left: full (default) mode in which concepts and names are connected within the same structure. Middle: full (default) mode with concepts separated from their names. Right: Finger-is-Moon-Mode (FIMM) in which concept labels reference the concepts' names. The dashed lines symbolise that these are alternative ways of displaying the same set of concepts, names, and relationships, not three parts of the same VAST display

among the three concepts, and the naming relationships (n) have been separated from one another. This way of displaying things may sometimes be helpful to avoid clutter. However, this approach comes at the price of somewhat lower parsimony because every concept now has to be displayed twice. In the display on the right-hand side, we use concept labels that directly reference the concepts' names. This is what we call the “Finger-is-Moon-Mode” (FIMM), as it abolishes the explicit distinction between signifier and referent. This constitutes another possible way of reducing clutter, but comes at the significant risk of overlooking the importance of semantics, especially (partial) homonymity, synonymity, and antonymity. For example, another display using FIMM could show that the concept Star has the same conceptual implications (Hot, Bright) as the concept Sun. Here, the use of FIMM might obscure the fact that this is the case simply because “Sun” and “Star” are two different words for the exact same type of thing (S). To highlight the risk of semantic ambiguities like this one, we recommend explicating when the FIMM is being used, by adding the respective acronym in one corner of the display (see Figure 2). Also, many concept names are

too long to be used as concept labels. In these cases, we recommend the approach exemplified in the middle of Figure 2.

As a next step, we will introduce four more types of relationships between concepts that frequently feature in argument structures. Figure 3 displays the ways in which they are distinguished from one another (in terms of lowercase letters accompanying the respective arrows), along with a very simple example for each type. Note that, for simplicity, this figure uses FIMM, as signalled by the acronym in the upper right-hand corner.

Relationship Type 3: Causation (c).

Many important articles and books have been written about causation (Eronen & Bringmann, 2021; Pearl, 1995; Pearl & Mackenzie, 2018; Rohrer, 2018). In VAST, we use a concept of causation that is also reflected in how most experimentalists tend to think about their research designs. This concept involves temporal order as a necessary ingredient: Causes always precede effects, but never the other way round. Also, causation would become evident if we were able to manipulate the suspected cause variable and then observe subse-

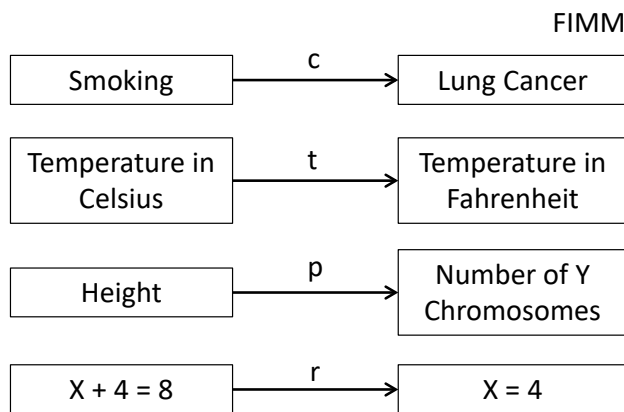


Figure 3

Four more types of relationships between concepts (c = causation, t = transformation, p = prediction, r = reasoning)

quent changes in another variable. Note that all of this concerns the ways in which we (and most people, presumably) *think* about causation, irrespective of whether such a suspected causal link may ever be proven or disproved in terms of data. Note also that most causal relationships between concepts could be — but do not have to be — decomposed into a number of intervening steps. For example, the causal relationship in Figure 3 reflects a relatively proximal link between cause (Smoking) and effect (Lung Cancer). It could be amended by inserting Tar Accumulation as a mediator that is caused by Smoking and that causes Lung Cancer.

Relationship Type 4: Transformation (t).

This relationship type is used to account for situations in which the applicability of one concept may be deduced from another concept by mere computation. The respective example in Figure 3 reflects a case in which one variable (Temperature in Celsius) is basically rescaled into another variable, by multiplying the former's values with a factor (1.8) and then adding a constant (32). The specific values for the factor and the constant are not displayed, but could be displayed. The transformation type of relationship may also be used to account for scoring procedures, such as the specific ways in which an operational measure of socio-economic status is derived from a number of indicators (e.g., highest degree attained, annual income).

Relationship Type 5: Prediction (p).

This type of relationship is about *knowing* something about the values of Y when we know something about

the values of X . Note that this is possible without knowing anything about the mechanism underlying the association. For example, type p relationships may ignore the direction of causal effects, as in the respective example in Figure 3: Here, a person's Height predicts the Number of Y Chromosomes that same person has, although certainly the former is not the cause for the latter. Often, such predictive relationships may be found and described first (e.g., a certain set of symptoms appearing together in patients), and only later be replaced by more specific explanations (e.g., in terms of a virus causing all of those symptoms).

Relationship Type 6: Reasoning (r).

This relationship type is about the conclusions that people draw from certain premises, on purely intellectual grounds. It reflects the idea that if some concept applies (e.g., $X + 4 = 8$, see Figure 3), one may infer that some other concept (e.g., $X = 4$, see Figure 3) also applies. Note that this is not limited to conclusions that would generally be regarded “logical”, but to just about any conclusion that someone thinks they may draw. In fact, VAST may be used to first explicate one person's line of reasoning and then refute that reasoning based on some other reasoning. For example, Peter may think that the results of some empirical study clearly suggest that Y is the case, whereas Trudy may think otherwise. Such discrepancies may then be explained using VAST, by analysing why exactly Peter and Trudy come to such different conclusions (e.g., because one of them trusts the authors of the study, whereas the other does not). So-called “logical conclusions” simply constitute a special case in which certain lines of reasoning are viewed as (in-)defensible by a group of people (e.g., scientists) who endorse some set of reasoning rules. That endorsement then serves as the premise for drawing conclusions as to whether “ X is reason to believe Y ” or not — which is another type r relationship.

Additional Relationship Types.

In the present paper, we only address those types of relationships between concepts that we think feature prominently in many arguments — everyday ones as well as scientific ones. Needless to say, the selection is and has to be somewhat subjective. It is relatively easy to come up with examples of other relationship types that may be useful to employ under certain circumstances (e.g., metamorphosis (m): when X turns into Y over time; association (a): when thinking of X makes it likely to also think of Y ; element of (e): when X is among the ingredients that, together, constitute Y etc.). We assume that the principles laid out in the following (e.g., regarding relationship strength and the

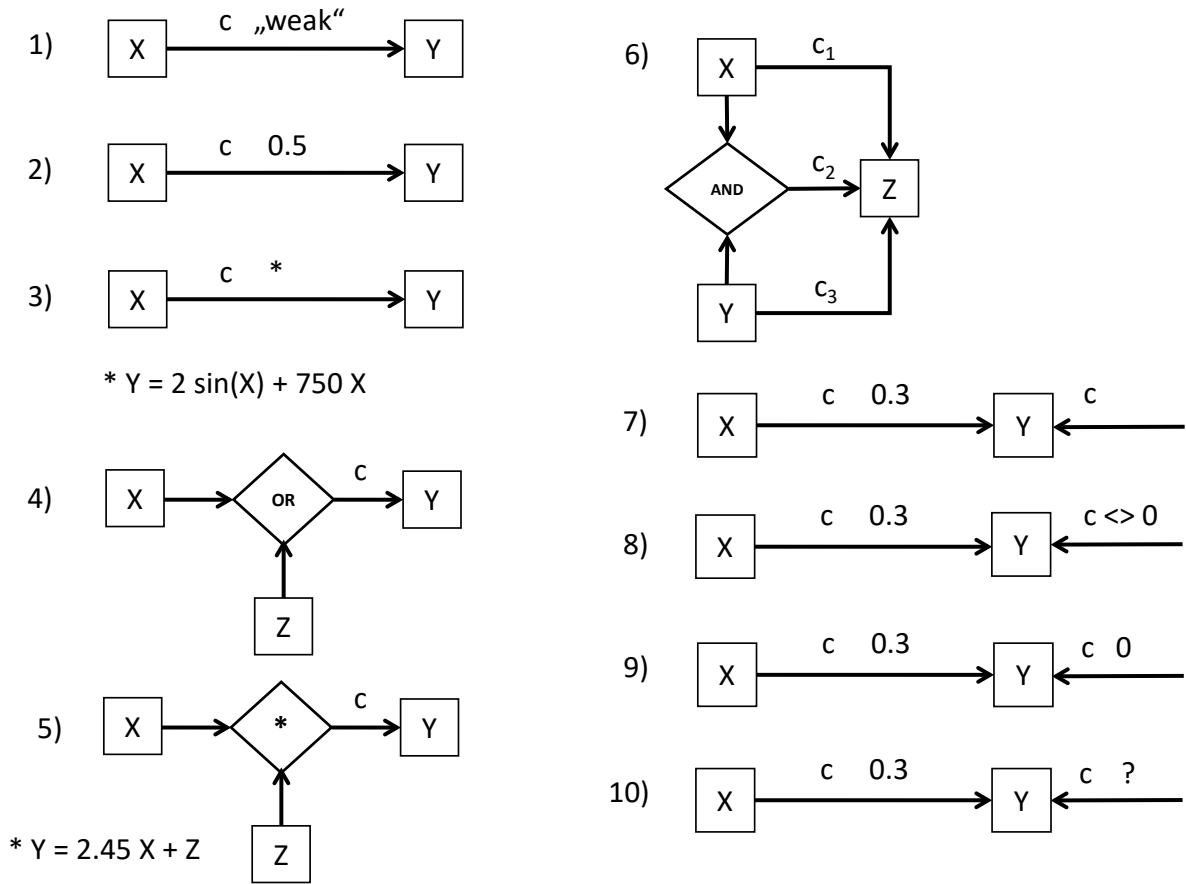


Figure 4

Relationship strengths and relationship patterns

construction of higher-order concepts) will still apply in these instances. In cases where a relationship between concepts is assumed to exist but the exact nature of that relationship is (yet) unknown, we recommend using the letter *u*.

Relationship Strength

In VAST, the default interpretation of an arrow that points from one concept (e.g., *X*) to another (e.g., *Y*) is that this relationship is considered relevant and positive (i.e., the more *X* the more *Y*). Thus, if an arrow is absent between *X* and *Y*, this means that the relationship is zero and/or that it is regarded unimportant for the present analysis. So far, we did not use any further specifications of relationship strength, and this approach may be perfectly sufficient in many cases. Sometimes, however, such specifications will be deemed useful or even necessary. VAST allows for the use of verbal labels such as “weak”, “strong”, “negative” etc. for this purpose. This

approach will often be appropriate when trying to visualise the structure of an existing argument that has been made using the natural language. It may also be the most useful approach when a numerical specification seems not possible (yet). Case 1 in Figure 4 displays an example. Note that we rather arbitrarily added the letter *c* (causation) to all the arrows in Figure 4. This may easily be replaced with any other relationship type, as everything we say here about relationship strength applies equally to all types.

If a simple numerical quantification of relationship strength is wanted, we propose using normalized coefficients ranging from -1 to 1 . “Normalized” means that these coefficients ignore the particular scales of the concepts that they connect, but rather quantify relationship strength in terms of proportions of these features’ ranges. Note that this is only possible if such a range may reasonably be assumed to exist. Based on our own experience, this seems to be the case with most psycho-

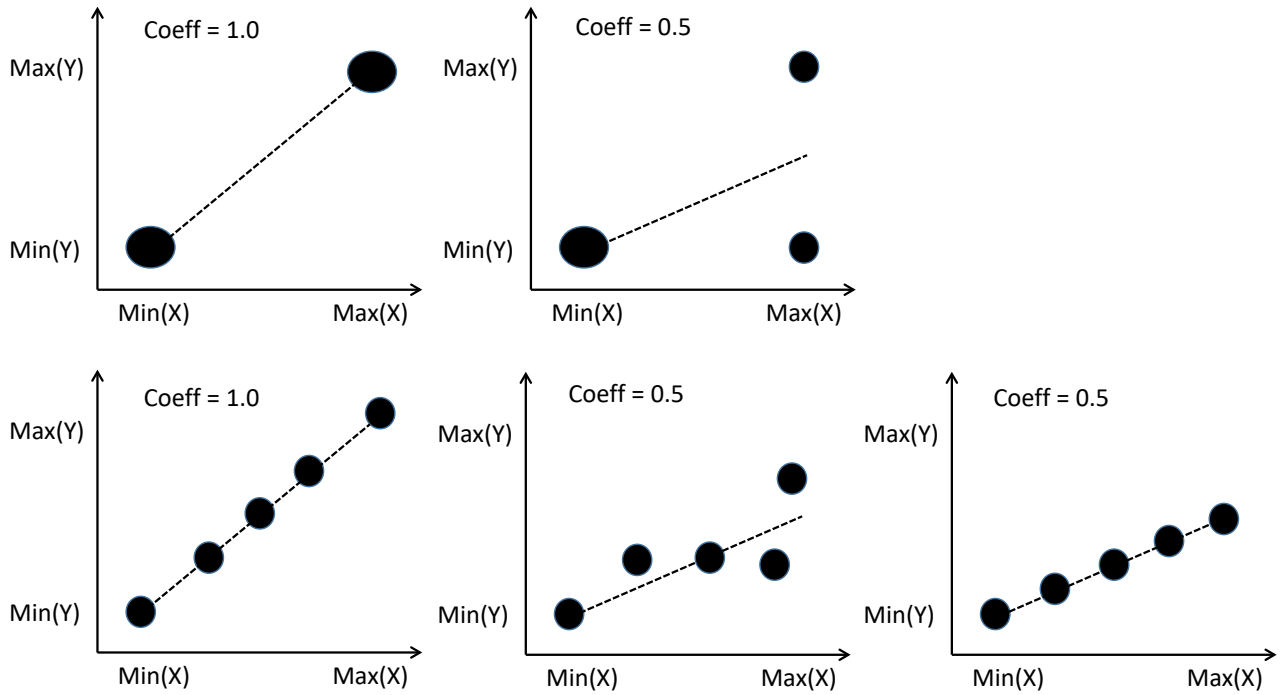


Figure 5

Relationship strength quantified in terms of default coefficients ranging from -1 to 1. Upper two examples: relationships between two dichotomous concepts. Lower three examples: relationships between two ordinal or metric concepts

logical concepts, however.

The default coefficient of relationship strength that we propose reflects the increase in Y (expressed as a percentage of the available range of that feature) that is associated with a “perfect” or “complete” increase in X (i.e., an increase that covers the whole available range of X, from the smallest conceivable value to the highest conceivable value). Case 2 in Figure 4 presents an example in which the coefficient is 0.5.

This type of coefficient may be applied to relationships between dichotomous concepts as well as to relationships between continuous concepts. In the former (dichotomous) case, it may be interpreted as the percentage of cases in which a complete change in X (from 0 to 1) is accompanied by a change in Y (from 0 to 1). In the latter (continuous) case, the coefficient reflects the average increase on the continuous Y scale that accompanies the largest possible increase (from 0 to 1) on the continuous X scale (again both expressed in terms of percentages of the respective ranges). Figure 5 contains a number of examples showcasing this broad applicability. Negative coefficients are to be interpreted accordingly: the more X is the case, the less Y is the case.

This default coefficient of relationship strength is generic enough to be applied to all types of relationships between concepts (e.g., Type p: “wearing glasses” makes it 70 percent likely for a person to also be “smart”; Type r: It is 90 percent reasonable to assume someone “is in love with you” when that person “giggles a lot while talking to you”; Type c: being “obese” makes it 50 percent likely for someone to develop “Diabetes Type II” as a consequence).

For non-dichotomous concepts (see the lower three panels of Figure 5), the default coefficient of relationship strength is largely agnostic regarding the distributions of concepts’ values: For example, the relationships displayed in the middle panel and in the panel to the right have the same strength coefficient (0.5) but in the latter case the relationship is deterministic whereas in the former case it is noisier. This difference may also be accounted for in VAST, as we will discuss in the next section (“Noise”).

VAST’s default coefficient of relationship strength is based on percentages of the ranges of the concepts that the relationship connects. Sometimes, however, there may be good reasons to deviate from the default (e.g., when earthquake magnitude on the unbounded Richter

scale is part of an argument). In such cases, using other measures of relationship strength (e.g., an exact function translating X into Y) is possible. As the exact function connecting certain concepts will often be too long to be written above an arrow in its entirety, we suggest placing it somewhere else in the display and referencing it using an asterisk (e.g., Case 3 in Figure 4).

Diamonds should be used when several concepts are jointly related to another concept. This includes logical connectives such as AND, OR and XOR (exclusive OR), as shown in Case 4 in Figure 4. When a more specific formula is needed to derive a joint output from several inputs (e.g., a scoring procedure), the diamond and asterisk elements may be combined, as shown in Case 5. A diamond with AND inside it may also be used to symbolise the interaction effect that two concepts (X and Y) have on a third concept (Z). Case 6 in Figure 4 displays this possibility along with the two main effects of X and Y on Z, so this is basically a VAST-type depiction of two-factor ANOVA. We also use Case 6 to showcase the possibility of indexing coefficients (c_1 , c_2 , c_3). Doing so is often useful to facilitate discussions among analysts.

Noise

Sometimes, we may not only wish to display the relationship between specific concepts, but acknowledge that there are additional unspecified influences (“noise”) on these concepts, as well. To symbolise these influences, we recommend using “noise arrows” similar to the ones that are used in Structural Equation Modelling (SEM). A noise arrow always points toward a given concept (i.e., the one that is affected by the noise) but does not originate in a specific concept. Cases 7, 8, 9 and 10 in Figure 4 provide examples. Note that, other than in Structural Equation Modelling, noise arrows in VAST do not stand for the residuals that remain between observed Y values and the Y values that one predicts from X. Rather, they stand for other influences apart from X that may move the values of Y toward its maximum (default, positive coefficient), or toward its minimum (negative coefficient). Case 7 displays the default situation in which noise may lead to an increase in concept Y. If noise may move the values of a concept in either direction, this may be specified using “< > 0”, as shown in Case 8. If it is important to specify that there is no noise, this may be expressed using a coefficient of zero, as shown in Case 9. This expresses the idea that the only factor whose values may make a difference with regard to the values of concept Y is concept X. Finally, when an influence may exist and be relevant for the analysis, but one is not really sure yet, we recommend using a question mark as “coefficient” (see Case 10 in Figure 4).

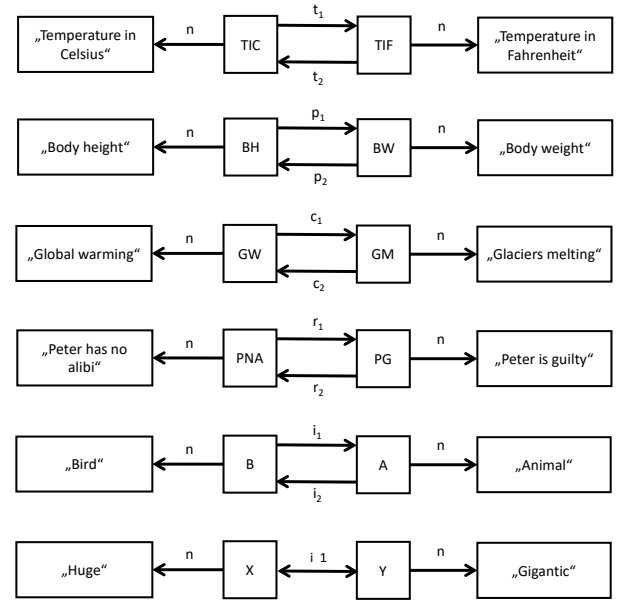


Figure 6

Relationship arrows connecting concepts both ways. Use of two arrows implies that the exact ways in which the concepts are related depend on direction. Use of a bidirectional arrow (as in the last case) implies that direction does not matter. The $i\ 1$ coefficient in the latter case reflects the assumption that concepts X and Y are identical

Relationship Direction

Arrows in VAST stand for IF-THEN relationships between concepts. We will now briefly address the direction into which arrows may point, and how this differs between relationship types. Generally speaking, if there is an arrow pointing from X to Y, there may also be an arrow pointing from Y to X. In most cases, the shape of the respective relationship will differ depending on its direction. If both directions are of interest to the current analysis, we recommend signalling this difference by using two separate arrows, one for each direction. Figure 6 presents a few examples.

A few specifics need to be briefly discussed in this regard: First, naming relationships are special in that arrows may only point from a concept to its name but not the other way round. Second, when the strengths of conceptual implications (type i) between X and Y differ depending on direction, this means that one concept (the one that is the target of the arrow with the higher coefficient) is broader and more inclusive than the other. This is highly relevant to all kinds of concep-

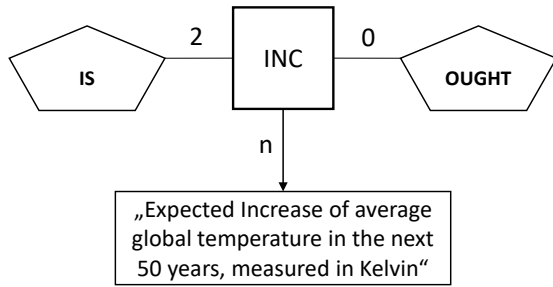


Figure 7

A simple example with IS- and OUGHT-statements regarding the same concept

tual hierarchies (taxonomies). Third, VAST does allow for causal (type c) arrows pointing from X to Y and back. This is relevant for displaying all kinds of positive and negative feedback loops. Note that this diverges from recommendations in the literature on Directed Acyclic Graphs (Rohrer, 2018). Fourth, VAST also allows for displays of circular reasoning (type r), because the purpose of VAST is not to prescribe rules as to how one should think, but to make visible the ways in which someone thinks. This includes the possibility of displaying beliefs that others may find unconvincing or even irrational. Fifth, we recommend using a bidirectional arrow with an “i 1” coefficient for expressing the idea that X and Y are identical (see Figure 6). Likewise, a bidirectional arrow with a coefficient of “i -1” would imply that one concept is the exact opposite of the other (e.g., between concept H named “huge” and concept T named “tiny”).

The IS and OUGHT elements

In many arguments, the extent to which something is considered to be the case and the extent to which something should be the case play important roles. To capture these extents, VAST uses two special elements, called IS and OUGHT. Both denote specific values on a given concept. They are important in a variety of ways: First, disagreements often arise because people start from different premises regarding the extent to which something IS the case (e.g., whether vaccines are safe) or the extent to which something OUGHT to be the case (e.g., whether one should trust the government). Second, discrepancies between IS and OUGHT-values on the same concept often explain why people decide to act in certain ways — often they do so in order to move the IS-value closer to the OUGHT-value.

In VAST, IS and OUGHT are symbolized by pentagons which include the respective term (IS or OUGHT) in

capitals. These are connected to one or more concepts using simple lines rather than arrows, in order to distinguish them from relationships between concepts. The specific IS- and OUGHT-values are written next to the respective lines. Figure 7 shows a very simple example.

IS and OUGHT are not concepts themselves but rather denote specific locations within the range of values that a given concept may have. If IS/OUGHT is used, specifying the metric of measurement may sometimes be helpful (e.g., Figure 7). If it is not possible or useful to specify this metric, we recommend expressing IS and OUGHT values in terms of fractions of the normalized range (between 0 and 1) of the respective concept. If no specific value is given, we recommend using “applies more likely than not” (> 0.5) as the default interpretation.

Note that IS may be interpreted as a measure of the respective concept’s central tendency (e.g., the arithmetic mean). It is possible, however, to provide whole ranges of IS- or OUGHT-values, if a single value is deemed insufficient.

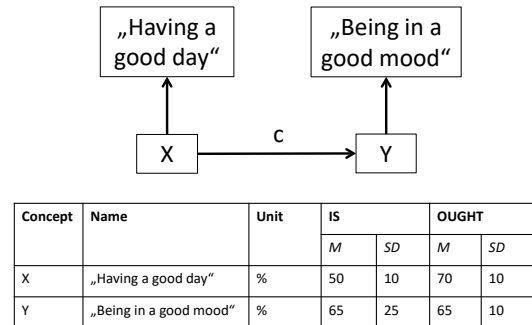


Figure 8

An example of how assumed and desired concept features may be displayed separately in a table, to avoid clutter

Going further, it may sometimes be helpful to specify the assumed and/or desired distribution characteristics of a concept even more. In such cases, we recommend providing the respective information in a separate Table on the side, to avoid clutter. An example is shown in Figure 8. The example tells us that, if a person (= object) “has a good day”, that person will be more likely to also “be in a good mood”, and that this effect is a causal one. The table also tells us that (a) it would be good if the average percentage of people having a good day would increase (from IS: 50 to OUGHT: 70), and that (b) the variation (SD) in the percentage of people being in a good mood would go down (from IS: 25 to OUGHT:

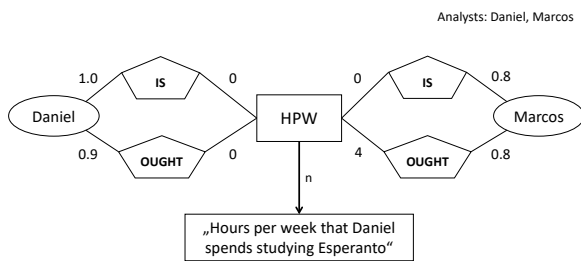


Figure 9

Different IS and OUGHT values for different perspective-holders

10). The latter goal may be rooted e.g., in the assumption that it is important to avoid extreme unhappiness in people.

The Perspective Element

Almost by definition, arguments tend to involve disagreements between viewpoints. To account for this, VAST incorporates a so-called “perspective” element that reflects how strongly a given entity agrees with something. This “entity” is usually a person, but it may also be a group of people or something more abstract like a corporation.

The perspective element may only be used to condition IS and OUGHT statements. When used to condition an IS-statement, it reflects the extent to which a perspective-holder agrees that the given level of the concept applies. When used to condition an OUGHT-statement, it reflects the extent to which the perspective-holder agrees that this is the most desirable level of the concept. Levels of agreement may be quantified using values between 0 (“does not agree at all”) and 1 (“agrees completely”). If a perspective-holder’s level of agreement is left unspecified, we recommend using “tends to agree” (> 0.5) as the default interpretation.

If a (part of a) VAST display does not explicate who the perspective-holder is, then IS and OUGHT statements reflect the view of the analyst who created the display (see next section).

Figure 9 showcases the use of the perspective element in VAST: The name of the perspective-holder is displayed inside an oval, and the strength of the perspective-holder’s belief is again expressed using coefficients ranging from 0 to 1. Note that a value of 0 would only imply a complete lack of agreement with X, not the belief that the opposite of X is true. This is necessary because many of the concepts that we use in everyday life do not have clearly defined opposites. If

it seems necessary to not only visualise a perspective-holder’s lack of agreement with X but also what else (Y) they believe in, that alternative view will have to be specified, as well.

The specific example given in Figure 9 conveys a wealth of information at one glance: Daniel and Marcos both think that Daniel does not spend any time studying Esperanto. However, Daniel is perfectly certain about that (1.0) whereas Marcos — who cannot really know for sure — is a little less certain (0.8). Also, Daniel is almost certain (0.9) that he should not take up any Esperanto-learning, whereas Marcos is also quite sure that Daniel should spend 4 hours per week learning Esperanto. Making differences such as these visible may go a long way in explaining the different behavioural choices that people make.

The perspective element of VAST incorporates the important issue of subjective certainty, which plays a key role in scientific theorizing. In fact, if one plotted all of the possible IS-values (X-axis) against a perspective-holder’s subjective certainties (Y-axis) and rescaled the latter such that their sum is 1, one would basically obtain a density distribution very much akin to a Bayes prior. Outside of scientific theorizing, however, inspecting entire distributions of possible IS-values is relatively rare. Thus, we recommend displaying the IS-value with the highest subjective certainty as a default. If necessary, alternative IS-values and their respective certainties may be displayed in addition. The perspective feature may also be used to express someone’s “hunches” (e.g., the suspicion that there may be another yet unrecognized factor involved in accounting for some effect). This also includes suspicions as to what someone may or may not have meant by saying something (i.e., implications).

The Analyst Element

Each VAST display has to be created by someone. Notably, the persons creating such displays are responsible for arranging the various concepts and their relationships with one another in the most accurate or helpful ways possible, but not for the actual content of the respective argument. In fact, it is possible to display the structure of an argument with great precision while at the same time disagreeing wholeheartedly with most or all of the points that are being made. This is why we prefer to call these persons “analysts” rather than “authors”. We recommend naming the analyst who created a display in a header, as also shown in the upper left corner of Figure 9. The persons named as analysts are responsible for the display in its entirety (again: irrespective of how much they agree with the display’s actual content). To make this point clear, we

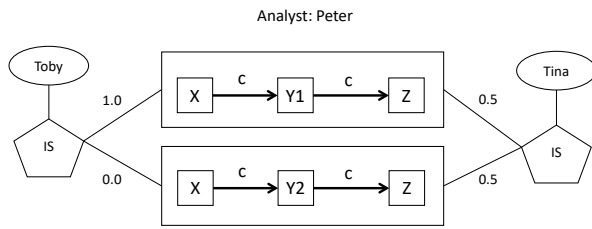


Figure 10

Higher-order concepts. Here, IS and OUGHT statements for different perspective-holders refer to two higher-order concepts, each of which contains a number of causal relationships between lower-order concepts

decided to include the same persons (Daniel and Marcos) in Figure 9, in two different roles: as analysts, and as perspective-holders. Altogether, the Figure tells us that Daniel and Marcos (analysts) agree that Daniel and Marcos (perspective-holders) have different views on how much time Daniel should spend learning Esperanto (the specifics of their viewpoints were discussed above and will not be repeated here). In other words, they “agree to disagree”.

Higher Order Concepts

After having introduced all of the different ways in which concepts may be related to one another, as well as the IS, OUGHT, perspective and analyst elements, it is now time to introduce a final element of great importance: In VAST, any combination of elements may itself become a “higher-order concept” (HOC) and thus be related to other (higher-order) concepts or be the subject of IS or OUGHT statements. In this, all of the rules explained so far do apply as well.

Figure 10 displays a very simple example. Here, two persons (Toby and Tina) are portrayed (by Peter) as disagreeing in regard to the question of whether the causal effect of X on Z is mediated by Y1, or by Y2. Toby is certain that the former is the case, whereas Tina is undecided. Similar displays may be used to account for all sorts of differences between viewpoints, such as naming conventions (i.e., what it IS that a given label refers to, or what that label OUGHT to be used for). Such issues are as greatly relevant to contemporary psychology as they were decades ago (Block, 1995).

It is important to note that a higher-order concept binds all of its components together in an inseparable fashion: The higher-order concept applies (to some object) if and only if all of its lower-order components apply.

By using higher-order concepts, VAST users may

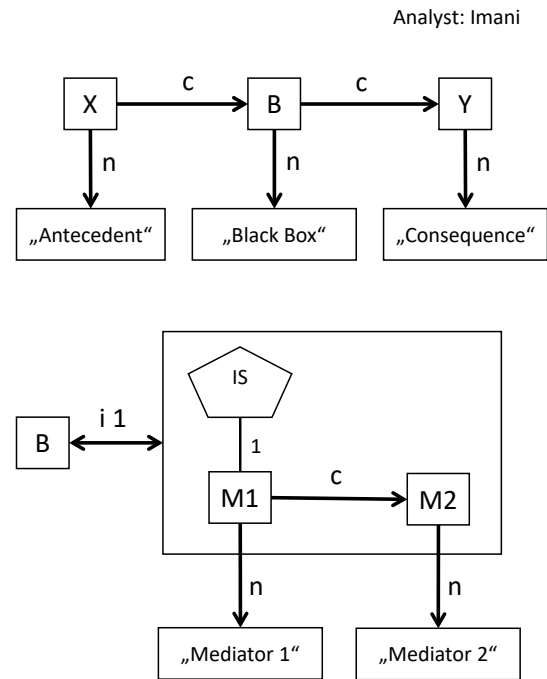


Figure 11

Using higher-order concepts to “zoom in” on a particular part of a VAST display. The lower part of the display details what concept B is about

“zoom in” on certain parts of a display if they wish to elaborate on its details, or “zoom out” when they decide to rather ignore some of the details for some time. Figure 11 provides an example: The lower part of the figure “zooms in” on the meaning of one of the concepts (B) that features in the causal chain displayed in the upper part of the figure. Specifically, we learn that B stands for M1 being the case (IS) and M1 eliciting M2. Furthermore, all of this elicits Y.

Higher-order concepts may be used to display a hierarchy of concepts that apply to different sets of objects. So far, we abstained from explicating the sets of objects that the concepts in a VAST analysis apply to. Instead, we tacitly assumed that those objects were the same across all concepts. Sometimes, however, specifying the sets of objects to which different concepts apply is necessary. We recommend using Greek letters for this purpose.

Figure 12 displays a hypothetical example. Here, intelligence test scores and school grades were obtained from students (τ). Remember that the thick black edges of the respective concept frames symbolise the fact that these concepts were actually measured. We assume

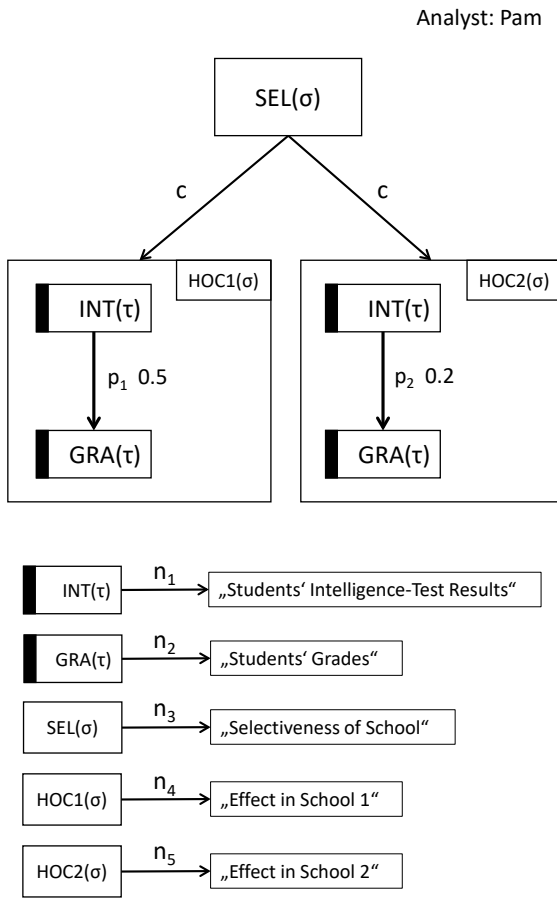


Figure 12

Use of concepts (*INT*, *GRA*, *SEL*) and higher-order concepts (*HOC1*, *HOC2*) that are applied to different sets of objects (τ vs. σ)

that the students' intelligence test scores do predict their grades to some extent. Note that we use p instead of c relationships in this display, because the two types of data are empirically related (probably because they both reflect the students' actual cognitive abilities), but the students' test results are not the cause of their grades.

The figure also tells us that this predictive relationship is assumed to be different for students in School 1 as compared to students in School 2. We learn this from comparing coefficients p_1 and p_2 , and from how the two higher-order concepts (*HOC1* and *HOC2*) are named. At this point, a new set of objects (σ) has to be introduced to distinguish the two schools from one another. We also learn that something else (*SEL*) named "selectivity" is assumed to vary across the same objects

(σ), and that this variation causally explains the difference between p_1 and p_2 .

Discussion

In this paper, we presented the first version of VAST, the Visual Argument Structure Tool. It provides a set of clear rules by which the ways in which people think and speak about things may be visually organized, in order to help understand those ways better. VAST may be used for constructing new arguments, as well as for analysing, completing, revising and/or (partly) refuting existing ones. The system captures some of the key types of elements that many everyday arguments and scientific theories share by means of a relatively small set of graphical symbols. In our view, its appeal lies in its intuitiveness and relative economy, in its capacity to account for any degree of specification or relative fuzziness, and in its applicability to basically any content domain. In the next section, we give a brief overview of the system's possible uses. Depending on the material at hand, one or several of the following goals may play a more prominent role in a VAST analysis: (a) afford comprehensiveness, (b) explicate premises in terms of what IS the case and what OUGHT to be the case, (c) clarify the views that different (groups of) people have, including areas of (dis-)agreement, (d) identify areas of under-specification, inconsistency or outright contradiction, (e) deduce defensible conclusions.

Potential Uses of the System

Theory Specification.

There is no shortage of complaints that the mainly narrative theories which are so common in the humanities and in psychology are of limited value because of their relative fuzziness and under-specification. At the same time, proposals as to how this situation may be improved are largely lacking, and the existing ones are often relatively unspecific themselves. We think that VAST may offer a solution to this problem, for two reasons: First, a VAST analysis may help pinpoint those parts of a narrative theory that may and should be better specified, and then aid in the specification process. We consider this the preferable approach compared to the alternative of rejecting narrative theories altogether. By specifying a theory better, its "empirical content" (empirischer Gehalt; Glöckner and Betsch, 2011; Popper, 2002, page 96) and thus ultimately its utility will be improved (e.g., it will become easier to refute).

Second, VAST allows for accommodating any level of fuzziness that seems acceptable or unavoidable at present. This very much aligns with the idea of scientific theory-development as an incremental process of

gradually increasing specificity. Furthermore, by being able to incorporate natural-language components of an argument, a VAST display may help bridge the chasm that exists between the humanities and the “harder” sciences, with psychology dangling somewhere in between. In the Appendix we provide a somewhat more complex example, showcasing an attempt to clarify the meaning of a short theory paragraph from a research paper.

Nomenclature Issues.

Psychology in particular has long suffered — and continues to suffer — from significant jingle- and jangle-problems (Block, 1995): To this day, psychologists often use different words to denote the same thing, or the same words to denote different things. Both practices are at odds with scientific ideals of efficiency and parsimony. VAST may be used to help improve on the present situation quite a bit by making visible (at a glance), (a) which terms are used, (b) by whom, (c) to denote what (including the relationships among the concepts that the terms refer to). The according displays will almost certainly involve naming and implication relationships as well as the perspective element. For example: Is the thing that is called “narcissism” by author A the same as the thing that is called “narcissism” by author B (e.g., in terms of its assumed or shown relationships with other concepts)? To what extent are “arrogance”, “dominance”, and “self-enhancement” just different words for the same thing, and is that the same thing that is also called “narcissism” by some scientists? And so on.

Facilitating Scientific Discourse.

Through its perspective element, VAST analyses may very well be used to account for the inherently social nature of all scientific discourse (Oreskes, 2020), as they enable an explicit and comprehensive showcasing of points of convergence or disagreement between the views of different scholars studying the same subject. We assume that scientific debates may become significantly more efficient when making sources of disagreement visible and then working through them, one after another, possibly in an iterative fashion. This way, a VAST analysis may actually serve as a kind of road-map to help guide the scientific process (e.g., in systematic attempts at forming consensus).

Peer Review.

VAST may also be used as a tool in peer review. This seems particularly promising when reviewing a paper from a research field that the reviewer is not that familiar with. In such cases, it may be helpful to first

organise the available information in the paper as to (a) how many relevant concepts there are (b) how they are assumed to relate to one another, (c) how they are named, (d) how they were measured, (e) how well these measurements reflect the assumed relationships between the concepts of interest, and so on.

Use as a Tool for Gathering Research Data.

The use of VAST may also be helpful when people’s belief systems (e.g., so-called conspiracy theories) are the research domain of interest. For example, how aware are people of the logical (in-)consistencies within their own worldviews? What happens if you make them aware of the existing inconsistencies (e.g., do they add new components post hoc that mitigate them)? Which components of people’s belief systems are particularly hard to change (e.g., the ones that are of key importance to several intertwined belief systems)? And: do people find it easier to map arguments they agree with, as compared to arguments with which they do not agree?

Explicating the Paths from Premises to Conclusions (and Back).

VAST may be used to derive defensible conclusions from a given set of premises, or to elucidate the ways in which a given perceiver seems to draw conclusions from such premises. Likewise, VAST may be used to infer the premises upon which some existing set of conclusions was built. Often, perspectives may be of particular importance in such analyses. This is because believing different things to be true (IS) or desirable (OUGHT) goes a long way in explaining wildly different conclusions (e.g., in terms of how one should act).

Finding Common Ground.

The potential use of VAST for working toward consensual viewpoints is not limited to scientists (see above). We hope that VAST may just as well be used to enable more traceable and rational conversations among proponents of viewpoints that may seem irreconcilable at first (e.g., regarding abortion, second amendment rights, vaccination etc.). The extent to which this hope is warranted will have to become the subject of future research, however.

Teaching Critical Thinking and the Art of Argumentation.

VAST may be used as a teaching tool, helping teachers explain to students the various ways in which concepts may be related to one another, and the important roles that IS and OUGHT statements as well as different perspectives play in many arguments. Ideally, these

things would be taught by way of analysing existing arguments together, or by jointly developing new ones (cf. Cullen et al., 2018).

Comparison with Related Tools

VAST's intended domain of use overlaps very significantly with those of many other systems, most prominently Directed Acyclic Graphs (DAG; Pearl, 1995; Pearl and Mackenzie, 2018; Rohrer, 2018) and Structural Equation Modelling (SEM). Major points of convergence with these systems are obviously the display of concepts or variables ("nodes"), and the display of relationships between them ("edges"). In SEM, numerical coefficients are often used to express the strengths of relationships between variables. A similar route is taken when creating so-called "research maps" summarizing the theoretical assertions and the evidence speaking for or against them, for a given research field (Matiasz et al., 2018). From SEM, VAST has further borrowed the use of "noise arrows" to symbolise additional, unspecified influences. Diverging from SEM conventions, however, VAST uses a different default meaning for the absence of arrows between concepts: Whereas in SEM this usually signals an unrelatedness of variables, in VAST it means "unrelated or not related in ways relevant to the current argument". This is a somewhat more liberal approach, and more in line with how everyday arguments are structured, according to our experience: When people do not talk about relationships between concepts, this usually means that they see no reason for doing so, but not necessarily that they assume the respective coefficient to be zero.

The coverage of VAST exceeds those of DAGs and SEMs by a wide margin: Like DAGs and SEMs, VAST does cover the (measured or unmeasured) features of objects and the relationships (causation and association) among those features. Unlike those other two frameworks, however, VAST also covers some more "psychological" relationship types such as naming, conceptual implication and reasoning. For this reason, the default coefficient of relationship strength in VAST expresses the covariation of two concepts in terms of percentages of ranges (e.g., how much more will I consider an object to "be a car" if it "has tires?"). VAST also goes beyond the aforementioned systems in that it enables an explicit accounting for assumptions as to how much something IS and OUGHT to be the case, and for differences between people in regard to such assumptions. All of this is unquestionably of key importance for many everyday arguments.

The coverage and methodology of VAST overlaps considerably with tools developed in philosophy, such as MindMup (<https://maps.simoncullen.org/>), Reason!

Able (Van Gelder, 2002) and others (for an overview of argument visualisation approaches, see Okada et al., 2014). These other tools usually enable users to zoom in on any parts of a verbal argument and deduce the logical relationships among them. This concerns reasons for drawing certain conclusions as well as objections to doing so. In VAST, these are captured using positive or negative type *r* relationships. Many tools account for premises that may lead to certain conclusions either by themselves, or in combination. In VAST, these would be distinguished from one another in terms of separate vs. combined (AND/OR) arrows pointing toward a concept. Several tools also afford the possibility of making whole strains of argument (e.g., "X is a reason to believe Y") the subject of further reasoning (e.g., "Z is a reason not to believe that X is a reason to believe Y"). In VAST, this is captured using higher-order concepts. Variants of IS-statements and quantifications of reasoning strength are also found in some existing tools (e.g., Reason!Able). However, a major difference between these tools and VAST is that the former deal exclusively with relationships of the reasoning (*r*) type, whereas the latter also accounts for many other possible types of relationships between concepts, while still capturing their strengths with the same (default) metric.

Limitations and Outlook

At this early stage in the development of VAST, it is difficult to predict how eagerly it will be picked up and eventually be used by others. We have spent significant amounts of time over the course of approximately three years developing, testing, revising and refining the system through numerous iterations, trying to make it work with analyses of diverse sets of examples both from within science and outside of it. We are convinced that the current version does work reasonably well, but we certainly expect additional improvements in the future. To facilitate these, we encourage our readers to give it a try, to put the system to use on whatever arguments they find interesting, and to let us know about the experiences they make. News media articles, statements by political agents (e.g., parties or office-holders), court sentences, advertisements, and of course science texts are all fair game. Based on our own experiences, we predict that, like us, most readers will find this type of analysis intellectually challenging. We hope that the substantial effort that tends to be associated with specifying argument structures this way will not deter people from trying.

VAST analyses may be real eye-openers in regard to the, well, vast level of complexity that does permeate many arguments but tends to be overlooked when sticking to purely narrative ways of formulating them. In

this regard, a reviewer (Julia Rohrer) alerted us to a potential conflict of interest, especially for scientists who consider using the system: Given the current incentive structure in academia, there may be good (i.e. rational) reasons to *avoid* greater levels of theory specification. In fact, low levels of specification may be selected for because using them is likely to involve a lower risk of being proven wrong. Ill-specified theories will also be less likely to incur strong negative reactions from reviewers and may be easier to “sell” to the public. All of this is true and may in fact explain part of why weak theorizing is so persistent in psychology to this day. We do consider it unlikely that the present paper will incite a mass movement of theory specification enthusiasts. However, based on our own experience with the tool, we do believe that for those psychologists who already are genuinely interested in improving on the specificity of their (and others’) theorizing, using VAST is definitely worth a try. Considerably greater clarity is usually achieved.

At present, VAST merely consists of a set of conceptual distinctions and related rules for how they should be visualized. As this is the core of the system, complete and satisfactory VAST analyses are already possible using any standard graphics tool, or even just paper and pencil. However, our ultimate goal is to implement the system as a free web resource that will be capable of (a) developing VAST displays by asking users the right questions and (b) checking any given display for consistency and completeness.

Appendix: A More Complex Example

The task at hand is to clarify the meaning of the following theory paragraph from the paper by Theves et al. (2020), using VAST. Such a clarification attempt may — and usually does — lead to an identification of areas of underspecification, ambiguity or even contradiction. Note, however, that this particular paragraph was picked for no other reason than being a relatively typical example of narrative theorizing in psychology (and because it deals with the subject of “concepts”, which play a key role in VAST). We do not consider this a particularly problematic case, but rather just use it to showcase a typical application of VAST. Thousands of other paragraphs may have been used just as well. The paragraph from the (Theves et al., 2020) paper (page 7318) goes like this:

“Concepts are organizing structures that define how contents are related to each other and can be used to transfer meaning to novel input (Smith & Medin, 1981; Kemp, 2010). Their formation thus inherently depends on generalization over, and integration of experiences. Thus, a role of the hippocampus in generalization seemed considerable due to its roles in binding elements into spatial

and episodic context (Davachi et al., 2003; Komorowski et al., 2013; Davachi, 2006; Ranganath, 2010) as well as integration of information over episodes (Schlichting et al., 2015; Davis et al., 2012; Collin et al., 2015; Milivojevitich et al., 2015 [...])”

Figure 13 shows a VAST display created by Daniel (analyst) to reflect the paragraph above. More specifically, it shows what this analyst thinks the authors (perspective) of the paragraph are saying (IS), which is everything inside the largest frame. Note that, for simplicity, a minor portion at the end of the narrative paragraph was excluded, as indicated by the use of brackets at the end (see above).

Daniel used the default mode of VAST in which all concepts and their names are separated from one another. He identified 8 relevant concepts to be accounted for. Five of these (C1 to C5) reflect key sentences from the paragraph, as shown in the lower part of the figure where the respective naming relationships are listed. Note that some minor inferences and modifications were necessary to enable each concept name to speak for itself. The other three concepts (L1 to L3) represent the references to the literature that also feature in the paragraph. These are empirical in nature, as indicated by the use of thick black edges on the respective frames. They are also assumed to be given, as indicated by the use of IS pentagons.

In this first VAST display, Daniel focuses on what he thinks is the reasoning structure in the paragraph. His use of VAST’s default mode along with his decision to set the naming relationships aside allows us to fully concentrate on this reasoning structure, without being distracted by the actual content of specific concepts. There is no need for indexing the naming relationships in this context, because they will not be individually discussed.

Daniel identified seven relevant relationships of the reasoning type (r_1 to r_7): Because L1 is given (IS), we may believe C1 to be the case, via r_1 . This is how Daniel interprets the first citation. If C1 is the case, we may also believe C2 to be the case, via r_2 . This is how Daniel interprets the first “thus”. If all of this is the case, we may also believe C3 to be the case, via r_3 . This is Daniel’s interpretation of the second “thus”. However, this (r_3) is only true if C4 and C5 are also the case, via r_4 and r_5 . This latter conditioning is how Daniel interprets the “due to” in the paragraph. Furthermore, L2 is given, which is a reason (via r_6 ; reflecting the second citation) to believe C4 to be the case. Finally, L3 is also given, which is a reason (via r_7 ; reflecting the third citation) to believe C5 to be the case.

Note that the analyst had to make numerous auxiliary assumptions about things that were not entirely clear in the narrative paragraph itself: For example, the analyst

Analyst: Daniel

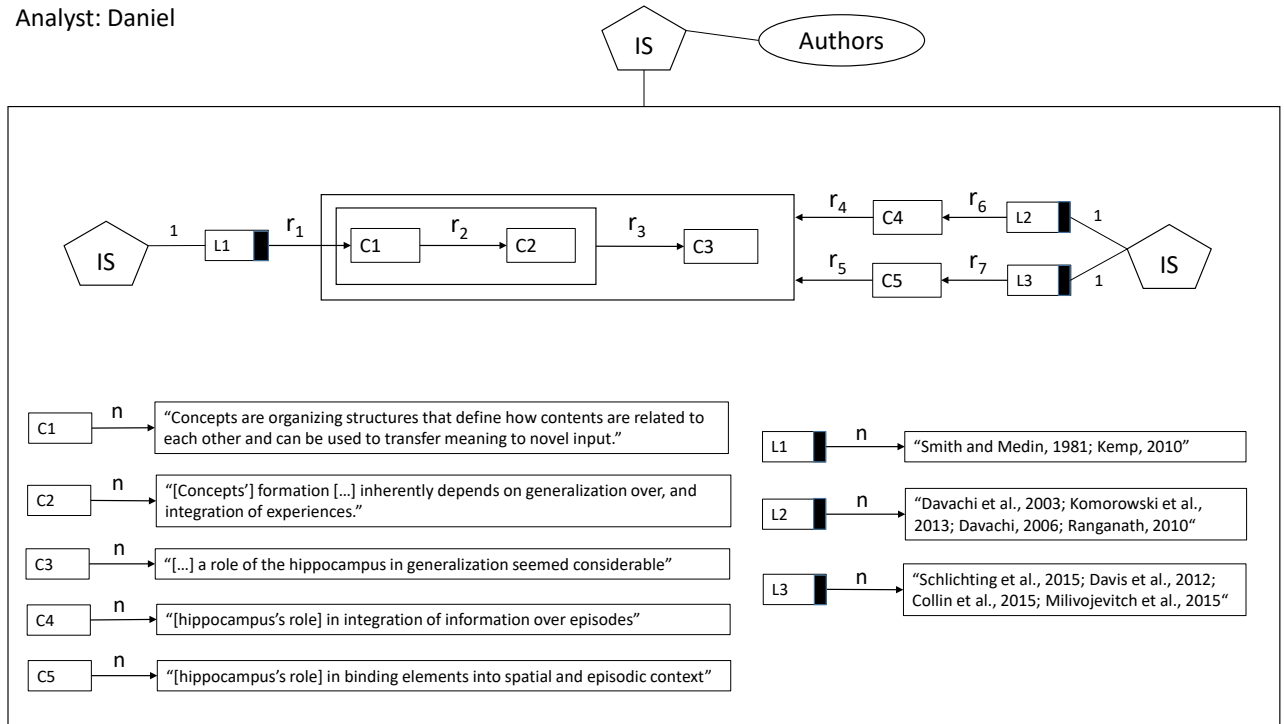


Figure 13

Reasoning structure in a theory paragraph by Theves et al. (2020), according to Daniel (analyst)

treated all of the papers supposedly supporting a given proposition as a unified whole (i.e., a single concept). Alternatively, he could have used a separate concept for each paper, which would have meant that each paper by itself supports the respective proposition. Also, Daniel treated r_4 and r_5 as two entirely separate paths by which r_3 is supported. Alternatively, he could have interpreted the “as well as” in the respective sentence to mean that only the *combination* of C4 and C5 is a reason to believe C3, which he would then have had to express using an AND diamond.

Needless to say, each of these additional assumptions, and any other element in the display, may be challenged any time — but only if they are made explicit, which is the whole point of VAST analyses. By making specification gaps and ambiguities explicit this way, a VAST analysis may eventually help drive theory development.

As a next step, Daniel decides to zoom-in a bit more on two of the key concepts in the paragraph. The first of these is C1, which he interprets to be mostly about conceptual implications. This is how he interprets the word “are” in this concept’s name. Figure 14 shows the outcome of this zooming-in:

Daniel believes the content of C1 and the entire con-

tent of the largest frame underneath it to be interchangeable. To him, these two concepts have the exact same meaning, as indicated by his using a bidirectional implication arrow with coefficient 1 (signalling identity). What Daniel does here is specify the meaning of a higher-order concept (C1) by breaking it down into four components (C6 to C9) and some relevant relationships (i_2 to i_4) between them: Objects that are considered exemplars of C6 (named “concepts”) are also considered exemplars of C7, C8 and C9. Note that these latter implications use unidirectional arrows only. Note further that all of this is expressly Daniel’s opinion, and not necessarily shared by the authors of the paragraph. This is because only Daniel is named as the analyst in the byline, but the original authors of the paragraph do not appear as perspective-holders anywhere in the figure.

The second concept that the analyst chooses to “zoom-in” on is C2. Figure 15 showcases the result. As in the previous analysis, Daniel assumes that the content of C2 and the content of the largest frame underneath it are mutually interchangeable (i_5). Again, he explicates the meaning of a concept (C2) by breaking it down into a few subordinate concepts (C10, C11, and

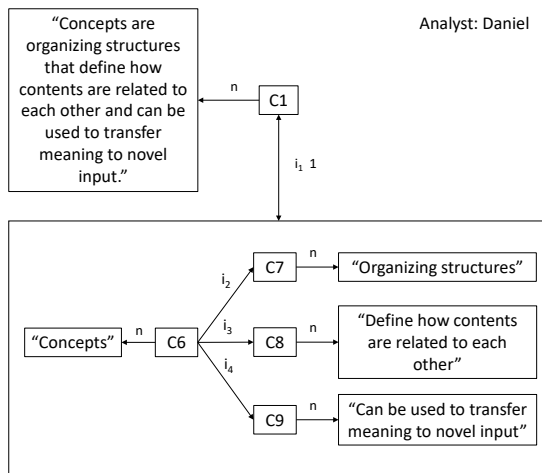


Figure 14

“Zooming-in” in on concept C1 (ignoring citations)

C12) and some relationships among them (this time of the *c* type, which is how Daniel interprets the word “depends” in the text).

Daniel’s use of an AND diamond signals that C12 may only be the case if C10 and C11 are both given, but not if only C10 or C11 are given. This is how he interprets the word “and” in the text. Furthermore, the *c 0* influence on C12 makes it clear that there are no other causal pathways by which C12 may come about. This is how Daniel interprets the word “inherently” in the text. However, the question mark above the third arrow pointing toward the AND diamond signals that Daniel is not sure whether another influence is needed to bring about C12. This is because the name of concept C2 only seems to say that C10 and C11 are *necessary* for C12 to happen, but not that they are *sufficient*. Here, the analysis points to a need for greater specification.

Finally, a brief word on the sets of objects that the concepts in the three figures pertain to. These are obviously not the same. In Figure 13, the objects are conceivable realities: the one that is described as given (in which certain rules apply and certain research papers exist), and possible alternative ones. In Figure 14, the objects are rather ill-defined. It may be useful to think of them as being various kinds of mental phenomena here (e.g., perceptions, memories etc.). In Figure 15, the objects may be thought of as variants of people’s developmental trajectories: Only in those in which C10 and C11 (and maybe something else, see question mark) take place, will we also see C12 happening. For simplicity, we did not explicitly account for the different sets of objects in Figures 13, 14, and 15.

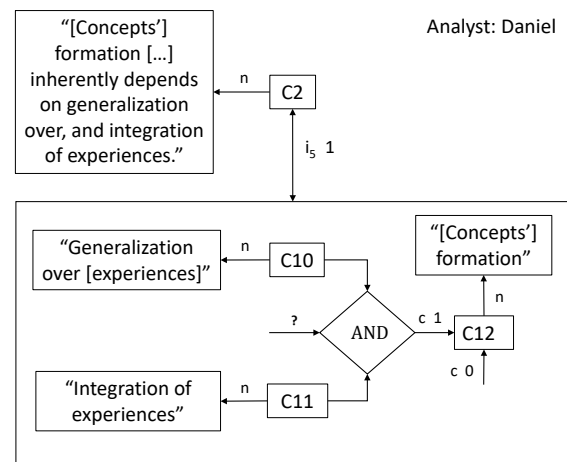


Figure 15

“Zooming-in” in on concept C2 (ignoring citations)

Author Contact

Corresponding Author: Daniel Leising (ORCID: 0000-00001-8503-5840), Technische Universität Dresden, daniel.leising@tu-dresden.de

Conflict of Interest and Funding

The authors state that they have no conflict of interest to declare. This project was not connected to any particular source of third-party funding.

Author Contributions

Daniel Leising: Conceptualization, Methodology, Visualization, Writing — Original Draft, Writing — Review & Editing; Oliver Grenke: Conceptualization, Methodology, Visualization, Writing — Review & Editing; Marcos Cramer: Conceptualization, Methodology, Visualization, Writing — Review & Editing.

Acknowledgments

The authors would like to thank Moritz Leistner for help with layouting this paper.

Open Science Practices

This article is theoretical and as such is not eligible for Open Science badges. The entire editorial process, including the open reviews, is published in the online supplement.

References

- Baroni, P., Gabbay, D., Giacomin, M., & van der Torre, L. (2018). *Handbook of formal argumentation*. College Publications.
- Block, J. (1995). A contrarian view of the five-factor approach to personality description. *Psychological Bulletin*, 117, 187–215. <https://doi.org/10.1037/0033-2909.117.2.187>
- Boole, G. (1854). *An investigation of the laws of thought: On which are founded the mathematical theories of logic and probabilities* (Vol. 2). Walton; Maberly.
- Borsboom, D., van der Maas, H. L. J., Dalege, J., Kievit, R. A., & Haig, B. D. (2021). Theory construction methodology: A practical framework for building theories in psychology. *Perspectives on Psychological Science*, 16(4), 756–766. <https://doi.org/10.1177/1745691620969647>
- Büning, H. K., & Lettmann, T. (1999). *Propositional logic: Deduction and algorithms* (Vol. 48). Cambridge University Press.
- Cullen, S., Fan, J., van der Brugge, E., & Elga, A. (2018). Improving analytical reasoning and argument understanding: A quasi-experimental field study of argument visualization. *npj Science of Learning*, 3, 21. <https://doi.org/10.1038/s41539-018-0038-5>
- Dablander, F. (2020). An introduction to causal inference. <https://doi.org/10.31234/osf.io/b3fkw>
- Devezer, B., Navarro, D. J., Vandekerckhove, J., & Buzbas, E. O. (2021). The case for formal methodology in scientific reform. *Royal Society Open Science*, 8(3), 200805. <https://doi.org/10.1098/rsos.200805>
- Eronen, M. I., & Bringmann, L. F. (2021). The theory crisis in psychology: How to move forward. *Perspectives on Psychological Science*, 16(4), 779–788. <https://doi.org/10.1177/1745691620970586>
- Frege, G. (1879). *Begriffsschrift: Eine der arithmetischen nachgebildete formelsprache des reinen denkens*. Verlag von Louis Nebert.
- Fried, E. I. (2020). Lack of theory building and testing impedes progress in the factor and network literature. *Psychological Inquiry*, 31(4), 271–288. <https://doi.org/10.1080/1047840X.2020.1853461>
- Glöckner, A., & Betsch, T. (2011). The empirical content of theories in judgement and decision making: Shortcomings and remedies. *Judgement and Decision Making*, 6(8), 771–721.
- Matiasz, N. J., Wood, J., Doshi, P., Speier, W., Beckemeyer, B., Wang, W., Hsu, W., & Silva, A. J. (2018). Researchmaps.org for integrating and planning research. *PLoS ONE*, 13(5), e0195271. <https://doi.org/10.1371/journal.pone.0195271>
- Muthukrishna, M., & Henrich, J. (2019). A problem in theory. *Nature Human Behaviour*, 3(3), 221–229. <https://doi.org/10.1038/s41562-018-0522-1>
- Okada, A., Buckingham Shum, S. J., & Sherborne, T. (2014). *Knowledge cartography: Software tools and mapping techniques*. Springer.
- Oreskes, N. (2020). *Why trust science?* Princeton University Press.
- Pearl, J. (1995). Causal diagrams for empirical research. *Biometrika*, 82(4), 669–688. <https://doi.org/10.1093/biomet/82.4.669>
- Pearl, J., & Mackenzie, D. (2018). *The book of why*. Basic Books.
- Popper, K. R. (2002). *The logic of scientific discovery*. Routledge Classics.
- Preparata, F. P., & Yeh, R. T. (1972). Continuously valued logic. *Journal of Computer and System Sciences*, 6(5), 397–418. [https://doi.org/10.1016/S0022-0000\(72\)80011-4](https://doi.org/10.1016/S0022-0000(72)80011-4)
- Robinaugh, D. J., Haslbeck, J. M. B., Ryan, O., Fried, E. I., & Waldorp, L. J. (2021). Invisible hands and fine calipers: A call to use formal theory as a toolkit for theory construction. *Perspectives on Psychological Science*, 16(4), 725–742. <https://doi.org/10.1177/1745691620974697>
- Rohrer, J. M. (2018). Thinking clearly about correlations and causation: Graphical causal models for observational data. *Advances in Methods and Practices in Psychological Science*, 1(1), 27–42. <https://doi.org/10.1177/2515245917745629>
- Rosch, E. (1978). Principles of categorization. In E. Rosch & B. B. Lloyd (Eds.), *Cognition and categorization* (pp. 27–48). Lawrence Erlbaum Associates.
- Smaldino, P. (2017). Models are stupid, and we need more of them. In R. R. Vallacher, S. J. Read, & A. Nowak (Eds.), *Computational social psychology*. Routledge. <https://doi.org/10.4324/9781315173726>
- Smaldino, P. (2019). Better methods can't make up for mediocre theory. *Nature*, 575, 9. <https://doi.org/10.1038/d41586-019-03350-5>
- Theves, S., Fernández, G., & Doeller, C. (2020). The hippocampus maps concept space, not feature space. *The Journal of Neuroscience*, 40(38), 7318–7325. <https://doi.org/10.1523/JNEUROSCI.0494-20.2020>

Van Gelder, T. (2002). Argument mapping with reason! able. *The American Philosophical Association Newsletter on Philosophy and Computers*, 2(1), 85–90.

Correction notice for Leising et al. (2023)

Date of correction: 2026-08-16

Correction to Visual Argument Structure Tool (VAST) Version 1.0

By Daniel Leising, Oliver Grenke and Marcos Cramer

Published in *Meta-Psychology*, 7, 2023.

<https://doi.org/10.15626/MP.2021.2911>

The first published version of this article contains (on page 9) a mistake in the text explaining Figure 8. The published sentence reads

“The table also tells us that (a) it would be good if the average percentage of people having a good day would increase (from IS: 50 to OUGHT: 70), and that (b) the variation (SD) in this percentage would go down (from IS: 25 to OUGHT: 10).”

The correct version would be

“The table also tells us that (a) it would be good if the average percentage of people having a good day would increase (from IS: 50 to OUGHT: 70), and that (b) the variation (SD) in the percentage of people being in a good mood would go down (from IS: 25 to OUGHT: 10).”

The authors would like to thank Felix Schönbrodt for catching this mistake.

This PDF contains the corrected version.