

***H* as a Measure of Complexity of Social Information Processing**

Kenneth D. Locke

*Department of Psychology
University of Idaho*

*Many studies have used *H* (a measure of unpredictability derived from information theory) to quantify the complexity of descriptions of persons across multiple roles. Interpreting these studies is problematic, though, because *H* confounds unpredictability across roles (which is typically the construct of interest) and unpredictability within roles (which is simply a function of the proportion of traits endorsed). The need to control for unpredictability within roles was highlighted by 3 demonstration studies in which controlling for unpredictability within roles eliminated relations between well-being and *H*. I also show how, controlling for unpredictability due to the number of traits endorsed and number of roles described, *H* provides a unique measure of role dependence and independence. However, *H* does not measure the type of role overlaps that would predict “spillover effects” between roles; therefore, I recommend an alternative index of role similarity for future research on spillover effects.*

The *H* formula derived from information theory, when used to measure the complexity of person descriptions, has been frequently misused and misunderstood. For example, researchers have assumed *H* measures the distinctness of the roles in person descriptions, when in fact it does not. Conversely, researchers have not realized that the number of descriptors endorsed can strongly influence *H*, when in fact it can. This article aims to clarify what *H* does and does not measure, what problems may arise when using *H*, and how to avoid those problems in the future.

A Review of *H* in the Social Cognition Literature

During the 1950s and early 1960s, the increased interest in information processing approaches to understanding perception, cognition, and communication was accompanied by a search for means to quantify information. One quantitative measure of the information in a set of symbols is *H*. Although *H* was originally used in research on processing of nonsocial information (Attneave, 1959; Broadbent, 1958; Garner, 1962; Miller, 1953), over the past 2 decades *H* has become popular in research on social cognition.

Although some studies have used *H* as a measure of the complexity of descriptions of other people (Lin-

ville, 1982; Linville & Jones, 1980; Scott, 1969), most social cognition research has used *H* as a measure of the complexity of self-descriptions or “self-complexity.” Researchers have tested relations between self-complexity and a number of other variables, including psychiatric problems such as narcissism (Rhodewalt, Madrian, & Cheney, 1998; Rhodewalt & Morf, 1995, 1998), personality traits such as self-consciousness or sociotropy and autonomy (Davies, 1996; Solomon & Haaga, 1993), and cognitive capacities such as attentional resources (Conway & White-Dy-sart, 1999). However, the main focus of self-complexity research has been the ability of *H* to predict measures of physical and mental well-being.

Interest in the link between *H* and well-being can be traced to two articles by Linville (1985, 1987). The articles reported that people lower in self-complexity experienced greater changes in affect and self-appraisal following success or failure, greater mood lability over a 2-week period, and, following high levels of stressful events, were more prone to depression, physical symptoms, and the flu and other illnesses. Linville also offered a sensible model to explain these striking results. According to the model, *H* measures *self-complexity*, defined as “having more self-aspects and maintaining greater distinctions among self-aspects” (Linville, 1987, p. 664). Both the number and the distinctiveness of self-aspects help people to be “less affected by the ups and downs of life.” With respect to number, when something good or bad happens to one part of the self (e.g., me-as-husband), it will influence a smaller portion of the self if there are many self-aspects than if there are only a few. With respect to distinction, the “spillover hypothesis” states, “The more related two

I am grateful to Dennis Baughman and Tom Sneed for their assistance in conducting the studies, and to Len Horowitz, David Dunning, and my reviewers for their helpful advice.

Requests for reprints should be sent to Kenneth Locke, Department of Psychology, University of Idaho, Moscow, ID 83844-3043. E-mail klocke@uidaho.edu

self-aspects are, the more likely thoughts and feelings about one are to spill over to color thoughts and feelings about the other” (Linville, 1987, p. 664). Conversely, the more distinct the self-aspects (e.g., describing me-as-husband in ways distinct from me-as-father), the less the spillover. Given the important implications of both the results and the model described in these articles, it is not surprising that the Social Science Citation Index showed that 195 articles had cited Linville (1985), and 293 had cited Linville (1987) by the end of the year 2000.

However, the subsequent research was not always supportive. Rafaeli-Mor and Steinberg (2002) conducted a meta-analysis of 70 studies from 46 different articles published between 1985 and 2000, which tested relations between self-complexity and well-being (i.e., some measure of mood, affect, self-esteem, or depression). To be included in the analysis, the self-descriptions had to be generated using Linville’s (1985) procedure in which participants ascribe a fixed (experimenter-provided) set of traits to a variable (self-generated) set of self-aspects. The key conclusion of the meta-analysis was that there was little evidence that H buffers the impact of negative events. Moreover, there was great heterogeneity in the study-level effect sizes. Different studies yielded effect sizes varying from strongly positive to strongly negative.

One reason for the heterogeneity may be that H does not measure what it has been purported to measure. Specifically, this article will show that H does not measure role distinctness, is at best an indirect and inefficient measure of role numerousness, and is greatly impacted by the percentage of traits endorsed. So for years H has been used to test hypotheses about the relation between role numerousness and distinctness and important variables such as affective reactivity, personality disorders, and physical and mental well-being. However, to the extent that H does not measure role numerousness and distinctness, those hypotheses in fact remain untested. To clarify what H does and does not measure, I will now examine the mathematics of H in detail.

Deriving the H Formula

Information is the reduction of uncertainty. Imagine you are uncertain which one of S statements is true of a person, and you can reduce your uncertainty by asking yes–no questions. Each question can convey a maximum of one bit of information, which occurs when the answer eliminates half of the possibilities. If there are initially S possibilities, and each successive question eliminates half of the remaining possibilities, you must ask $\log_2 S$ questions to determine which statement is true. Because each answer conveys one bit of informa-

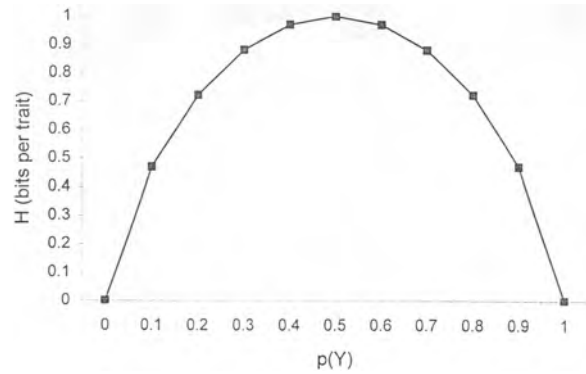


Figure 1. H as a function of the probability of trait endorsement.

tion, the total information conveyed, H , is $\log_2 S$ bits. If $p_i = 1/S$, where p_i is the probability that any particular statement i is the true statement, then $S = 1/p_i$. Therefore, $H = \log_2 S = \log_2(1/p_i) = -\log_2 p_i$.

Consider a simple example. We ask Jack if he is assertive. There are two possible answers: Yes (Y) or No (N). Therefore, $S = 2$, and $p(Y) = p(N) = .5$. Whether we compute H as $\log_2 S = \log_2(2)$ or as $-\log_2 p_i = -\log_2(.5)$, the answer is the same: 1 bit. Jack’s answer eliminates half the possibilities. It conveys one bit of information.

A question can yield one bit of information only if we are completely uncertain about the answer. For example, imagine we know that Jack says “yes” 75% of the time when asked if a trait describes him. The expected informational value of Jack’s answer is the information if Jack says “yes” and the information if Jack says “no,” weighted by the likelihood of each of those events; that is, $H = -\{(p[Y])\log_2(p[Y]) + (p[N])\log_2(p[N])\}$. In this case, $p(Y) = .75$ and $p(N) = .25$, so $H = -[.75]\log_2[.75] + [.25]\log_2[.25] = .81$. Note that Jack’s answer conveys .19 bits less information when $p(Y) = .75$ than when $p(Y) = .50$ because there is less uncertainty about what Jack will say. In other words, the response of a person who is predictable conveys less information than the response of a person who is unpredictable. Indeed, if Jack always answers yes, then $H = \log_2(1/1) = 0$. That is, if a response is completely predictable its occurrence conveys no information at all.

Figure 1 graphs the relation between $p(Y)$ and H . H is maximized when $p(Y) = p(N)$, because that is when we are most uncertain about what will happen. As $p(Y)$ and $p(N)$ diverge, H declines. The function is symmetrical because $p(N) = 1 - p(Y)$. H approaches its minimum (zero) as the probability of any particular outcome approaches its maximum (unity).

When considering more response categories than simply Y or N, the same logic applies: Sum the expected value of the information conveyed by each response category weighted by the likelihood of those responses. Thus, the general equation for H is:

$$H = -\sum_{i=1}^C p_i \log_2 p_i, \quad (1)$$

where C = the number of response categories.¹

Note that if the response categories are simply Y or N, then the p_i are $p(Y)$ and $p(N)$, and Equation 1 is equivalent to the formula given in the preceding paragraph and displayed in Figure 1. But regardless of the number of response categories, Equation 1 shows that H increases as the p_i become more equal and decreases as the p_i diverge.

H has been used as an index of category diversity in a number of fields. For example, H is widely used in ecology as a measure of ecological diversity (Magurran, 1988). In this context, the categories are species and the p_i is the proportion of individuals found in the i th species. The more equal the populations of different species, the greater the H . In an interesting extension to the ecology of human communities, Kreiner et al. (2001) used H to measure the diversity of non-profit community-based organizational activities. The community activities were categorized according to their main focus (such as health, education, conservation, and so on) and the p_i was the proportion of activities in each category. An example from the social cognition literature is Brewer & Lui (1984), in which participants sorted photos of members of a particular social group into subgroups of three or more persons. The categories were the subgroups and the p_i was the proportion of photos in each subgroup.

In the preceding examples, H was used correctly as a measure of complexity within a particular situation: the complexity of species within an ecosystem, activities within a community, or perceived subgroups within a social group. However, almost all social cognitive research has been concerned with complexity of trait descriptors across different situations or roles. When H is applied to this type of Trait $\times$ Situation or Trait $\times$ Role matrix, problems can arise.

H Applied to a Trait $\times$ Role Matrix

To understand the problems, we first need to understand how H is computed on a Trait $\times$ Role matrix. When we asked Jack if he was assertive, there were two response categories: Y and N. If instead we were to ask Jack if he is assertive with his mother and if he is asser-

tive with his father, there are four response categories: yes with both mom and dad (YY), yes with mom but not dad (YN), no with mom but yes with dad (NY), and no with both mom and dad (NN). If we were to ask about assertiveness in three different roles, there would be eight categories: YYY, YYN, YNY, YNN, NYY, NYN, NNY, and NNN. The general rule is: If we ask about R different roles, the number of response categories, C , is equal to 2^R . Thus, if we ask about four roles, $C = 2^4 = 16$.

If we ask Jack to use eight traits to describe how he is with his mother and with his father, we can display his self-description as a Trait $\times$ Role matrix of the sort shown in Table 1. Note that the number of traits in the matrix does not affect H . Regardless of the number of traits, H (the expected information content of each trait) remains $-\sum p_i \log p_i$. In Example 1, $p(YY) = p(NN) = p(YN) = p(NY) = .25$. Therefore, $H = -\sum p_i \log p_i = -\{(p[YY])\log_2(p[YY]) + (p[YN])\log_2(p[YN]) + (p[NY])\log_2(p[NY]) + (p[NN])\log_2(p[NN])\} = -\{4[.25]\log_2[.25]\} = 2$. In Example 2, $p(YY) = p(NN) = .5$, and $p(YN) = p(NY) = .0$, so $H = 1$. The more equiprobable the response categories, the higher the H . The four categories were equally probable in Example 1 but not in Example 2, so H was higher in Example 1 than in Example 2. Now that we understand how to compute H on a matrix, let us consider four problems that can arise when H is used in this way and how future research can avoid these problems.

Problem 1: H Confounds Uncertainty Within and Between Roles

Equation 1 shows that H increases as the proportion of traits in each category, p_i , become more equal. The p_i becomes more equal not only as between-role uncertainty increases but also as within-role uncertainty increases. That is, as we become less certain about whether or not a person will endorse a trait within each role (i.e., as $p[Y]$ and $p[N]$ converge), we become less certain about whether that person will show a particular pattern of endorsing that trait across roles.

To appreciate why, consider again the example of Jack describing himself with his mother and father. If

Table 1. Trait $\times$ Situation Matrices Illustrating Basic Properties of H

Situation	Trait							
	T1	T2	T3	T4	T5	T6	T7	T8
Example 1 ($H = 2$)								
R1: "with mom"	Y	Y	Y	Y	N	N	N	N
R2: "with dad"	Y	Y	N	N	Y	Y	N	N
Example 2 ($H = 1$)								
R1: "with mom"	Y	Y	Y	Y	N	N	N	N
R2: "with dad"	Y	Y	Y	Y	N	N	N	N

Note: H = measure of unpredictability; Y = yes; N = no.

¹In the social cognition literature, the H formula is often written as $\log_2 n - (\sum n_i \log_2 n_i) / n$, where n is the total number of traits used, and n_i is the number of traits that appear in a given category. This equation is equivalent to the one I employ, because $\log_2 n - (\sum n_i \log_2 n_i) / n = \log_2 n - (\sum p_i n \log_2 p_i n) / n = \log_2 n - \sum p_i \log_2 p_i n = \log_2 n - \sum p_i \log_2 p_i - \sum p_i \log_2 n = -\sum p_i \log_2 p_i$.

the two role descriptions are independent, probability theory states that the probability of a pattern of responses across roles is the product of the probabilities of each component response. Therefore, $p(YY) = p(Y) \times p(Y)$, $p(YN) = p(Y) \times p(N)$, $p(NY) = p(N) \times p(Y)$, and $p(NN) = p(N) \times p(N)$. If $p(Y) = .5$, then $p(YY) = p(YN) = p(NY) = p(NN) = .25$, and $H = -\sum p_i \log p_i = -\{4[.25]\log_2[.25]\} = 2$. Maximum uncertainty within roles yields maximum uncertainty across roles. As the probabilities of Y and N diverge within roles, the probabilities of the various combinations of Y and N across roles also diverge. For example, if $p(Y) = .8$, then $p(YY) = .64$, $p(YN) = p(NY) = .16$, and $p(NN) = .04$. In this case, $H = -[.64]\log_2[.64] - (.16)\log_2(.16) - (.16)\log_2(.16) - (.04)\log_2(.04) = 1.4$. Thus, simply knowing that Jack has an “acquiescence bias” reduces H from 2.0 to 1.4.

Problem 2: The Proportions of Positive and Negative Traits Endorsed Influence Positive and Negative Complexity

Problem 2 is just a specific instance of Problem 1 but one worth highlighting. Some researchers advise computing an H for positive traits (H_{pos}) and a separate H for negative traits (H_{neg}) because the relations between H and H_{neg} and H_{pos} appear at best inconsistent (Morgan & Janoff-Bulman, 1994; Rafaeli-Mor, Gotlib, & Revelle, 1999; Woolfolk, Novalany, Gara, Allen, & Polino, 1995). The examples in Table 2 show why H_{neg} and H_{pos} sometimes diverge from H . In all three examples, T1–T4 are positive traits, and T5–T8 are negative traits, $H = 2$, and within each role $p(Y) = .5$. In Example 1, $H = H_{pos} = H_{neg}$. The reason is that within every response category, the proportion of positive traits endorsed (p_{pos}) equal the proportion of negative traits endorsed (p_{neg}). However, if p_{pos} and p_{neg} differ within categories, then H_{neg} and H_{pos} may differ from H . In Example 2, $H = 2$ and $H_{neg} = H_{pos} = 1$. The discrepancy exists because the response categories are equally

probable when examining all traits but not when examining positive and negative traits separately. Instead, two response categories {YY, NN} contain only positive traits, and the other two {YN, NY} contain only negative traits. In Examples 1 and 2, $p(Y) = p_{pos} = p_{neg}$. H_{pos} and H_{neg} become more likely to diverge from H to the degree that p_{pos} and p_{neg} diverge from $p(Y)$. In Example 3, for instance, $p(Y) = .5$, but $p_{pos} = .75$ and $p_{neg} = .25$. In this case, $H = 2$, but H_{pos} and H_{neg} cannot exceed 1.5.²

H_{neg} and H_{pos} have been reported to predict different outcomes. For example, Morgan and Janoff-Bulman (1994) found posttrauma adjustment was related to H_{pos} but not to H_{neg} , and Woolfolk et al. (1995) found depression was related to H_{neg} but not to H_{pos} . The problem is that just as $p(Y)$ can influence H , so too can p_{pos} influence H_{pos} and p_{neg} influence H_{neg} . Therefore, it is not clear if adjustment and depression are related to variations in positive versus negative complexity or simply variations in the numbers of positive versus negative traits endorsed. For example, Woolfolk et al. may have found a link between depression and H_{neg} because depression affects p_{neg} , which, in turn, affects H_{neg} .

Overview of Demonstration Studies

To empirically demonstrate the preceding problems, three studies were conducted.³ In each study, participants used trait checklists to describe themselves in four different roles. The trait checklists contained equal numbers of positive and negative traits. In Study 1, participants were allowed to endorse any number of traits. In Study 2, participants were required to endorse a specific number of traits. In Study 3, participants were required to endorse a specific number of positive traits and the same number of negative traits.

The participants also completed measures of depression and esteem to show how relations between H and other variables can sometimes be misleading. Depression and esteem were chosen as the “other variables” simply because they were common in self-complexity research and were likely to influence p_{pos} and p_{neg} . By influencing p_{pos} and p_{neg} , a person’s depression and esteem scores could affect uncertainty about

Table 2. Trait $\times$ Situation Matrices Illustrating the Relationships of H_{neg} and H_{pos} to Overall H

Situation	Trait							
	T1	T2	T3	T4	T5	T6	T7	T8
Example 1								
R1: “with mom”	Y	N	Y	N	Y	N	Y	N
R2: “with dad”	Y	N	N	Y	Y	N	N	Y
Example 2								
R1: “with mom”	Y	Y	N	N	Y	N	Y	N
R2: “with dad”	Y	Y	N	N	N	Y	N	Y
Example 3								
R1: “with mom”	Y	Y	Y	N	N	N	Y	N
R2: “with dad”	Y	Y	N	Y	N	N	N	Y

Note: H = measure of unpredictability; Y = yes; N = no.

²In Example 3, because there are only four positive traits and $p_{pos} = .75$, H_{pos} is maximized when $p(YY) = .5$, $p(YN) = p(NY) = .25$, and $p(NN) = 0$. In this case, $H = -\sum p_i \log p_i = -[.5]\log_2[.5] - [.25]\log_2[.25] - [.25]\log_2[.25] - [0]\log_2[0] = 1.5$. The same would be true if $p(Y) = .25$, as is the case for the negative traits in this example.

³I had originally hoped to demonstrate these points by reanalyzing data from previous studies. I attempted to obtain the data from 12 articles published between 1991 and 1998 that had employed the H statistic. In all cases I was unable to conduct the reanalysis because the researchers were either unable (due to lost data or corrupted data files) or unwilling to share the data. I undertook these demonstration studies only after becoming discouraged about being able to use previously collected data.

whether or not that person would endorse positive or negative traits within roles, and thus overall uncertainty (H_{pos} and H_{neg}), even if those scores did not affect uncertainty about whether or not that person would make consistent responses across roles.

The influence of p_{pos} and p_{neg} on H can be controlled either experimentally or statistically. In Study 3, it was done experimentally by controlling p_{pos} and p_{neg} . In Studies 1 and 2, it was done statistically by controlling for H_{array} , the uncertainty that would exist if the responses were arranged in a one-dimensional array rather than a two-dimensional Trait $\times$ Role matrix. Mathematically, $H_{\text{array}} = -\{p[Y]\log_2(p[Y]) + p[N]\log_2(p[N])\}$. Conceptually, H_{array} is one's uncertainty about whether a trait will be endorsed in a particular context when one does not know if it was endorsed in other contexts. Because knowing how a trait was endorsed in other contexts cannot increase uncertainty, H cannot exceed H_{array} . Controlling for H_{array} , the residual H reflects uncertainty about how the endorsements are organized across roles. H_{array} can be computed on all traits, on just positive traits (i.e., $H_{\text{pos-array}}$), or on just negative traits (i.e., $H_{\text{neg-array}}$).

Study 1

In Study 1, $p(Y)$ was allowed to vary across participants. As $p(Y)$ varies from 0 to 1, H_{array} first rises and then falls (see Figure 1). To the extent that H_{array} influences H , $p(Y)$ will show a similar relation with H . Moreover, if depression or esteem influence p_{pos} (and thus $H_{\text{pos-array}}$) or p_{neg} (and thus $H_{\text{neg-array}}$), then depression or esteem might show a similar relation with H_{pos} or H_{neg} . If those relations are only due to effects on $H_{\text{pos-array}}$ or $H_{\text{neg-array}}$, however, then controlling for $H_{\text{pos-array}}$ or $H_{\text{neg-array}}$ should eliminate them.

Method

Participants. College students (85 women, 40 men, 1 unknown) participated for extra credit in undergraduate psychology courses.

Beck Depression Inventory-2 (BDI-2). The BDI-2 (Beck, Steer, & Brown, 1996) is a widely used 21-item self-report measure of symptoms associated with depression.

Rosenberg Self-Esteem Inventory (RSEI). The RSEI (Rosenberg, 1965) is a widely used 10-item self-report measure of overall feelings of value and worth.

Traits. The traits were selected from a pool of 400 adjectives for which there were published social

desirability norms from two independent samples (Hampson, Goldberg, & John, 1987; Norman, 1967). In both samples, students rated trait desirability on a 1 (*extremely undesirable*) to 9 (*extremely desirable*) scale. I eliminated traits whose ratings across the two samples differed by more than 2 scale points. Then I averaged the ratings across the two samples to obtain a more stable index of desirability. Then, so that the content and valence of the traits would be somewhat independent, I paired together traits that were contrasting in meaning but whose mean desirability ratings were within .5 units of each other.

Finally, I selected 18 traits pairs: 3 very positive (lively-relaxed, independent-sociable, adaptable-stable; desirability rating between 7 and 8), 3 medium positive (humble-bold, dignified-playful, frank-sensitive; desirability between 6 and 7), 3 mildly positive (soft-tough, outspoken-quiet, cautious-carefree; desirability between 5 and 6), 3 mildly negative (docile-dominant, conventional-rebellious, shy-dramatic; desirability between 4 and 5), 3 medium negative (impatient-indecisive, submissive-argumentative, meek-demanding; desirability between 3 and 4), and 3 very negative (irritable-apatetic, distrustful-gullible, vain-insecure; desirability between 2 and 3). Participants in the 75% positive (POS) condition received checklists containing nine positive pairs and three negative pairs. (The three negative pairs varied across participants with the constraint that one pair be mildly negative, one pair be medium negative, and one pair be very negative.) The checklists in the 75% negative (NEG) condition were constructed in the same way.

Procedure. In small classroom settings, the participants completed the BDI-2 and RSEI and described the following four roles: (a) "you when you are engaging in school or work-related activities," (b) "you when you are with peers of the same sex," (c) "you when you are engaging in recreational activities," and (d) "you when you are with peers of the opposite sex." The participants described each role by circling traits in an alphabetically ordered list of 24 traits (containing either 18 positive and 6 negative traits or 18 negative and 6 positive traits). The roles were presented in two different orders. The affect measures were also presented in two different orders: Either the BDI-2 came prior to and the RSEI came after the self-description, or the RSEI came prior to and the BDI-2 came after the self-description.

Participants were randomly assigned to one of the eight conditions of a 2 (order of roles) $\times$ 2 (order of affect measures) $\times$ 2 (POS vs. NEG) design. There were no significant effects of order or gender, so these variables are not discussed further. The number of participants in the POS and NEG conditions were, respectively, 65 and 61.

Results

Effect of number of traits endorsed. There is an inverted *U* relation between $p(Y)$ and H_{array} . To the extent that H_{array} influences H , there may be a similar relation between $p(Y)$ and H . In this study, H_{array} did in fact explain most of the variance in H , $r(126) = .80, p < .001$. Therefore, if the $p(Y)$ for most participants is less than .5, then one would expect a positive correlation between $p(Y)$ and H . If most $p(Y)$ values are greater than .5, one would expect a negative correlation between $p(Y)$ and H . If $p(Y)$ averages about .5, one would expect no linear relation. As Figure 2 shows, $p(Y)$ was less than .5 for all but 2 participants, so $p(Y)$ had a positive effect on H , $r(126) = .68, p < .001$. As Figure 3 (top) shows, p_{neg} was less than .5 for all participants, so the p_{neg} also had a positive effect on H_{neg} , $r(126) = .82, p < .001$. In contrast, Figure 3 (bottom) shows that p_{pos} values were spread over the middle portion of the distribution, and consequently there was no linear association between p_{pos} and H_{pos} , $r(126) = .01, ns$.

Effects of depression and esteem. In the NEG condition, H_{neg} was positively related to BDI-2 scores, $r(59) = .35, p = .005$, and negatively related to RSEI scores, $r(59) = -.30, p = .02$. When such results occurred in previous studies they were interpreted as evidence that a more complex negative self-image predicts higher levels of depression and lower levels of esteem (e.g., Woolfolk et al., 1995). But we now recognize that depression and esteem can affect H simply by affecting the number of traits endorsed. Indeed, once we controlled for $H_{neg-array}$ (which is solely a function of p_{neg}), the partial r between H_{neg} and BDI-2 scores was $r(58) = .05, ns$, and the partial r between H_{neg} and RSEI scores was $r(58) = -.09, ns$. Thus, the effects of BDI-2 and RSEI scores on H_{neg} were due to their effects on p_{neg} —that is, the overall probability of endorsing negative traits—rather than their effects on the complexity of negative trait endorsements across roles. BDI-2 and RSEI scores did not predict H_{neg} in the POS condition (perhaps because p_{neg}

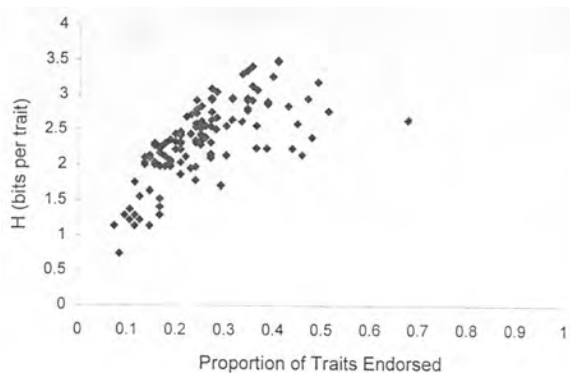


Figure 2. H as a function of the proportion of traits endorsed in Study 1.

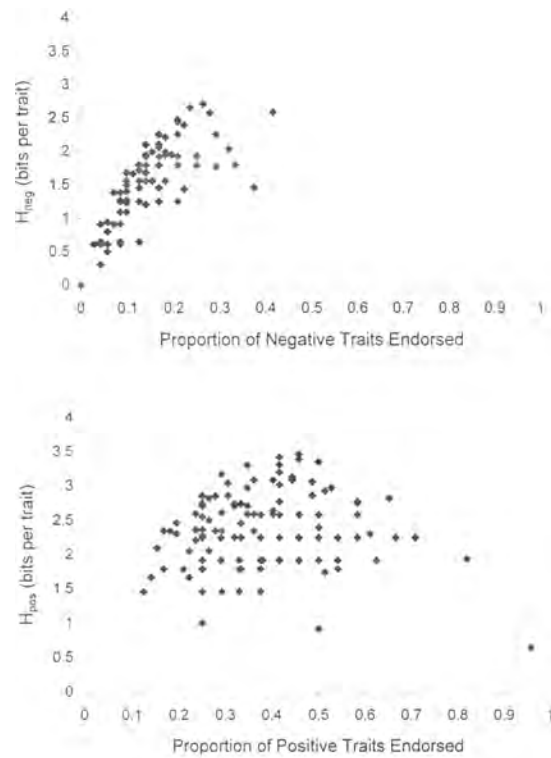


Figure 3. H_{pos} as a function of the proportion of positive traits endorsed and H_{neg} as a function of the proportion of negative traits endorsed in Study 1.

varied less in that condition) nor H_{pos} in either condition (perhaps because p_{pos} did not predict H_{pos}).

Study 2

Study 1 clearly demonstrated that H confounds different sources of uncertainty and that sometimes the predominant source is within-role uncertainty. The purpose of Study 2 was to demonstrate one way to separate uncertainty within roles from uncertainty across roles—namely, requiring participants to endorse a specific number of traits. Controlling the total number of traits endorsed should solve Problem 1. However, because participants can still endorse different numbers of positive and negative traits, it should not solve Problem 2—the confounding of uncertainty within and across roles when considering positive and negative traits separately.

Method

Participants. College students (65 women, 37 men, 10 unknown) participated for extra credit in undergraduate psychology courses.

Traits. Using the same trait pool and procedure as in Study 1, I selected 5 desirable pairs (desirability

greater than 6) and 5 undesirable pairs (desirability less than 4). The desirable traits were: bold, dignified, frank, humble, independent, lively, playful, relaxed, sensitive, and sociable. The undesirable traits were: apathetic, distrustful, gullible, impatient, indecisive, insecure, irritable, nosy, withdrawn, and vain.

Procedure. The procedure was identical to that of Study 1 except that the participants were asked to describe each self-aspect by circling a specific number of traits (namely, 4, 6, 8, 10, 12, 14, or 16 traits) from a list of 20 traits. Participants were randomly assigned to one of the conditions of a 2 (order of roles) $\times$ 2 (order of affect measures) $\times$ 7 (number of traits endorsed) factorial design. There were no significant effects of order or gender, so these variables are not discussed further. Participants who circled one too many or too few traits when describing a particular role were not excluded from the analyses, but those ($n = 2$) who circled more than one too many or too few traits were excluded. The number of participants in the 20%, 30%, 40%, 50%, 60%, 70%, and 80% endorsed conditions were, respectively, 17, 14, 17, 15, 16, 13, and 18.

Results

Effect of number of traits endorsed. Manipulating the number of traits endorsed ensured that the distribution of $p(Y)$ in the sample was approximately symmetrical around the midpoint of .5. When the $p(Y)$ distribution is symmetrical around .5, there should be no linear relation between $p(Y)$ and H , and there was none, $r(108) = .01$, *ns*. Thus, controlling $p(Y)$ eliminates Problem 1; however, it does not control p_{pos} and p_{neg} , so it does not eliminate Problem 2. As Table 3 shows, in every condition the participants endorsed more positive than negative traits, so the distributions of p_{pos} and p_{neg} were not symmetrical around .5. The mean p_{pos} was .65 (range = .15–1.00), and the mean p_{neg} was .35 (range = .00–.80). Because p_{pos} tended to exceed .5, the relation between p_{pos} and H_{pos} was negative, $r(108) = -.67$. Because p_{neg} tended to be less than

.5, the relation between p_{neg} and H_{neg} was positive, $r(108) = .74$.

Table 3 shows the mean H , H_{pos} , and H_{neg} for each condition. As $p(Y)$ varied from .2 to .8, H showed the expected inverted U curve, but H_{pos} showed only the decreasing half of the curve, and H_{neg} showed only the increasing half. For participants who endorsed less than 50% of the traits ($n = 48$), p_{pos} was close to .5 and p_{neg} was not. Consequently, H_{pos} was greater than H_{neg} , and $p(Y)$ had a linear relation with H_{neg} , $r(46) = .90$, but not with H_{pos} , $r(46) = .09$. Conversely, for participants who endorsed more than 50% of the traits ($n = 47$), p_{neg} was close to .5 and p_{pos} was not. Consequently, H_{neg} was greater than H_{pos} , and $p(Y)$ had a linear relation with H_{pos} , $r(45) = -.92$, but not with H_{neg} , $r(45) = .05$. When participants endorsed exactly 50% of the traits, p_{neg} and p_{pos} were equally close to .5 (i.e., .36 and .64), so H_{neg} and H_{pos} were almost the same.

Effects of depression and esteem.

BDI-2 scores were positively related to H_{neg} , $r(108) = .25$, $p < .01$, but not to H_{pos} , $r(108) = .06$, *ns*. RSEI scores were inversely related to H_{pos} , $r(108) = -.23$, $p < .05$, but not to H_{neg} , $r(108) = -.11$, *ns*. Whereas previous studies might have concluded that depression and esteem predicted the complexity of thinking about positive and negative features of the self, we now know that they may simply be predicting the proportion of positive and negative features endorsed. Indeed, controlling for $H_{neg-array}$, the partial r between H_{neg} and BDI-2 scores was $r(107) = .12$, *ns*. Controlling for $H_{pos-array}$, the partial r between H_{pos} and RSEI scores was $r(107) = -.09$, *ns*. Thus, as in Study 1, depression and esteem predicted the proportions—and thus the predictability—of positive and negative trait endorsements but not the complexity of patterns of endorsement across roles.

Study 3

Whereas the participants in Studies 1 and 2 could endorse different numbers of positive and negative traits, Study 3 participants were asked to endorse the

Table 3. Numbers of Traits Endorsed and H as a Function of Experimental Condition and Trait Valence in Study 2

Condition (% endorsed)	Traits Endorsed			H		
	Total	Positive	Negative	All Traits	Positive	Negative
20	4	3.0	1.0	2.2	2.3	1.2
30	6	4.8	1.3	2.5	2.6	1.3
40	8	6.3	1.7	2.7	2.4	1.5
50	10	6.4	3.6	2.9	2.2	2.4
60	12	7.4	4.5	3.0	2.0	2.4
70	14	8.1	5.9	2.7	1.7	2.6
80	16	9.3	6.7	2.1	0.9	2.3

Note: H = measure of unpredictability; $N = 110$. The proportion of positive traits endorsed = the number of positive traits endorsed divided by 10, and the proportion of negative traits endorsed = the number of negative traits endorsed divided by 10.

same specified number of positive and negative traits. So, whereas Studies 1 and 2 controlled for the effects of variations in p_{pos} and p_{neg} statistically, Study 3 did so experimentally.

Method

Participants. College students (72 women, 32 men, 17 unknown) participated for extra credit in undergraduate psychology courses.

Procedure. The procedure was identical to that of Study 2 except that instead of presenting the traits as a single list, the desirable traits and undesirable traits were presented in two separate lists. The participants were asked to describe each self-aspect by circling a specific number of traits (i.e., 2, 3, 4, 5, 6, 7, or 8) from the list of 10 positive traits and then the same number of traits from the list of 10 negative traits. Participants were randomly assigned to one of the conditions of a 2 (order of roles) $\times$ 2 (order of affect measures) $\times$ 7 (proportion of traits endorsed) factorial design. There were no significant effects of order or gender, so these variables are not discussed further. Participants who circled one too many or too few traits when describing a particular aspect were not excluded from the analyses, but participants ($n = 1$) who circled more than one too many or too few traits were excluded. The number of participants in the 20%, 30%, 40%, 50%, 60%, 70%, and 80% endorsed conditions were, respectively, 16, 19, 20, 18, 15, 15, and 17.

Results

The experimental design created roughly symmetrical distributions of $p(Y)$ —for positive traits, negative traits, and all traits combined. Thus, the relation between $p(Y)$ and H approximated an inverted U —for positive traits, negative traits, and all traits. Thus, there were no significant linear relations between $p(Y)$ and H —for positive traits, negative traits, or all traits, all $p_s > .05$. Nor were there any significant linear relations between BDI-2 or RSEI and H_{neg} or H_{pos} , all $p_s > .05$. Thus, once we controlled p_{pos} and p_{neg} , and concomitantly $H_{\text{pos-array}}$ and $H_{\text{neg-array}}$, there was no longer any evidence that BDI-2 and RSEI scores predicted variations in the complexity with which participants used those traits to describe themselves across roles.

Discussion of Findings

H_{array} is your uncertainty about whether a particular response was Y versus N when all you know is overall how many responses were Ys versus Ns. In all three studies, H_{array} explained a significant proportion of the variance in H . Consequently, the relations between

$p(Y)$ and H approximated the inverted U relation between $p(Y)$ and H_{array} shown in Figure 1. Accordingly, when $p(Y)$ tended to be less than .5 (as with negative traits in Study 1), the relation between $p(Y)$ and H was positive. When $p(Y)$ tended to be greater than .5 (as with positive traits in Study 2), the relation between $p(Y)$ and H was negative. When $p(Y)$ tended to be distributed evenly around .5 (as with all traits in Study 3), there was no linear relation between $p(Y)$ and H .

The important implication is that any variable may influence H simply because it influences $p(Y)$. For example, these results revealed significant relations between measures of well-being and H_{neg} or H_{pos} . However, statistically controlling for the impact of p_{neg} and p_{pos} on H_{neg} and H_{pos} eliminated these effects in Studies 1 and 2, as did experimentally controlling p_{neg} and p_{pos} in Study 3. Thus, the well-being measures predicted the positivity, not the complexity, of the self. But these studies were not concerned with the particular question of whether H is related to well-being. Rather, their purpose was to show why it is critical to consider the influence of $p(Y)$ in any study using H .

If the $p(Y)$ distribution varies across studies, then how variables that influence $p(Y)$ influence H also will vary across studies. Knowing this, the heterogeneity in the findings relating H and measures of well-being (Rafaeli-Mor & Steinberg, 2002) is not surprising. Unfortunately, previous studies neither controlled for nor even reported $p(Y)$. Therefore, we cannot know if the results of previous research were due to differences in the complexity of patterns of trait endorsement across roles or simply differences in the numbers of traits endorsed within each role. To address Problems 1 and 2, future research at least should report $p(Y)$ and, better yet, test the extent to which any relations involving H are due to uncertainty within roles, between roles, or both.⁴

Problem 3: H Confounds Role Numerosity and Role Independence

Once we control for the influence of uncertainty within roles, however, we face another problem. The residual H is a function of two conceptually and em-

⁴The appropriate measure of uncertainty within roles will vary across studies. In my demonstration studies, in which all participants were forced to apply the same set of traits to the same number of roles, the simplest possible control variable (H_{array}) was adequate. More complex studies may require more complex measures. For example, to control for variations in the number of roles (R) across subjects, one could use $R * H_{\text{array}}$. To also control for variations in $p(Y)$ between roles within a single description, one could compute H_{array} for each role separately and then sum them. Moreover, one could take into account the ways in which the number of traits under consideration (T) constrains the maximum possible H , which is normally $\log_2 C = \log_2(2^R) = R$. The most important such case (in terms of practical consequences) is when $T < C$. In this case, the number of trait categories that can actually occur is T rather than C , so the maximum H is $\log_2 T$.

pirically distinct sources of uncertainty between roles: role numerosity and role independence. Increasing R only permits, and does not necessitate, a greater H . An additional role increases H only to the extent that it increases uncertainty about how a trait will be used. Specifically, the increment in H due to a particular role is:

$$H_{inc} = -\sum_{i=1}^C p_i ((pY_i \log 2pY_i) + (pN_i \log 2pN_i)), \quad (2)$$

where C is the number of trait categories prior to adding the role, and the p_i are the probabilities of each of those categories. H_{inc} shares several properties with H_{array} : It can range from 0 to 1 bits, approaching its maximum as the $p(Y_i)$ and $p(N_i)$ converge toward .5, and declining as the $p(Y_i)$ and $p(N_i)$ diverge. The reason is that Equation 2 actually says: (a) divide the role into C parts, (b) compute H_{array} on each part separately, and (c) compute a weighted average of the H_{arrays} . Indeed, if $p(Y_i)$ is the same for all C , then the H_{inc} function will be identical to Figure 1.

Consider the examples in Table 4. In each example, $p(Y) = .5$, $H_{array} = 1$, and $R = 2$. What varies is H_{inc} . In Examples 1 and 5, $H_{inc} = 0$ because, within the categories defined by R1, the $p(Y_i)$ for R2 are zero or one. R1 predicts R2 perfectly, so R2 does not increase H at all. In Examples 2 and 4, $H_{inc} = .811$ because the $p(Y_i)$ are .75 or .25. In Example 3, $H_{inc} = 1$ because the $p(Y_i)$ are .5; R1 does not predict R2 at all.⁵

Thus, in studies in which R has been allowed to vary, any particular H may reflect a smaller number of relatively independent roles or a larger number of nonindependent roles. For example, Jane and Joe may both have an $H = 4$ and a $p(Y) = .5$, but Jane's self-concept may consist of 4 completely independent roles, whereas Joe's may consist of 12 highly similar roles. These very different organizations may have very different psychological implications. Therefore, my recommendation is that R at least be reported and ideally be entered as a separate predictor. Based on their analysis of empirical data, Rafaeli-Mor et al. (1999) also concluded that using "measures that separately and independently reflect the two underlying components of [role] quantity and overlap" (p. 351) may be more informative than using H alone.

Table 4. *Trait $\times$ Situation Matrices Illustrating the Relationship of H to Role Independence*

Situation	Trait							
	T1	T2	T3	T4	T5	T6	T7	T8
Example 1 ($H = 1.00$)								
R1: "with mom"	Y	Y	Y	Y	N	N	N	N
R2: "with dad"	Y	Y	Y	Y	N	N	N	N
Example 2 ($H = 1.81$)								
R1: "with mom"	Y	Y	Y	Y	N	N	N	N
R2: "with dad"	Y	Y	Y	N	Y	N	N	N
Example 3 ($H = 2.00$)								
R1: "with mom"	Y	Y	Y	Y	N	N	N	N
R2: "with dad"	Y	Y	N	N	Y	Y	N	N
Example 4 ($H = 1.81$)								
R1: "with mom"	Y	Y	Y	Y	N	N	N	N
R2: "with dad"	Y	N	N	N	Y	Y	Y	N
Example 5 ($H = 1.00$)								
R1: "with mom"	Y	Y	Y	Y	N	N	N	N
R2: "with dad"	N	N	N	N	Y	Y	Y	Y

Note: H = measure of predictability; Y = yes; N = no.

Problem 4: H is Not an Appropriate Measure of Spillover

Controlling for R and H_{array} , the residual H is a measure of role independence. Yet, H is instead often described as a measure of role distinctness. High H self-concepts have been repeatedly defined as having "greater distinctions among self-aspects" (Linville, 1987, p. 663), "different attributes in different roles or situations" (Dixon & Baumeister, 1991, p. 364), or "subelves that differ considerably from one another in terms of their defining attributes" (Morgan & Janoff-Bulman, 1994, p. 64). These definitions suggest that greater H should predict lesser spillover—spillover being when "feelings and inferences associated with the originally activated self-aspect spill over and color feelings and inferences regarding associated self-aspects" (Linville, 1987, p. 664). However, H does not measure differences between roles, and therefore is unlikely to predict spillover.

To appreciate why, imagine that Table 4 shows how Jack conceptualizes his relationship with his parents, and imagine Jack has had an upsetting falling out with his mom. In Example 1, Jack is saying: "How I am around mom is the same as how I am around dad." In Example 3, Jack is saying: "How I am with mom is in some ways the same as when I am with dad, but some things are different, too." In Example 5, Jack is saying: "How I am around mom is completely different from how I am with dad." Intuitively, one would expect Jack's upsetting thoughts and feelings about "me-with-mom" to color his thoughts and feelings about "me-with-dad" the most in Example 1 and the least in Example 5. If so, one would want a measure of

⁵The H_{inc} computations for the examples in Table 4 are shown here. In Example 1, $H_{inc} = -\{.5([1]\log_2[1] + [0]\log_2[0]) + .5([0]\log_2[0] + [1]\log_2[1])\} = 0$. In Examples 2 and 4, $H_{inc} = -\{.5([.75]\log_2[.75] + [.25]\log_2[.25]) + .5([.75]\log_2[.75] + [.25]\log_2[.25])\} = .811$. In Example 3, $H_{inc} = -\{.5([.5]\log_2[.5] + [.5]\log_2[.5]) + .5([.5]\log_2[.5] + [.5]\log_2[.5])\} = 1$.

spillover that is greatest in Example 1, smallest in Example 5, and decreases monotonically across the intervening examples. Instead, H is greatest in Example 3 (when the roles are completely independent), and smallest in Examples 1 and 5 (when the roles are completely identical or completely opposite). Thus, self-descriptions that should maximize spillover (such as Example 1) and those that should minimize spillover (such as Example 2) yield the same H value. The implication is clear: Because H does not measure spillover, the conclusions of dozens of published studies that have used H to measure spillover are invalid.

How Should Role Similarity be Measured?

On the basis of the data showing a lack of covariation between H and an index of feature overlap, which they called OL, Rafaeli-Mor et al. (1999) also concluded that H was a poor predictor of spillover. They further suggested that OL might be an appropriate alternative in future research on spillover. Although I concur with their conclusions in general, I question whether OL in particular will always be the best alternative.

OL is the conditional probability that if a trait is endorsed in one role it is also endorsed in another role, averaged across all pairs of roles. OL is thus a proportion that can range from 0 to 1. Note that OL defines similarity only in terms of shared Y responses, so it is more accurate to call it OL_Y . The problem with only considering shared Ys is that (all else equal) simply increasing $p(Y)$ will increase OL_Y . Consider the examples in Table 5. In each example the roles are independent ($r = 0$), so the probability of endorsing a trait in one role is independent of whether that trait was endorsed in another role. In other words, the conditional $p(Y)$ equals the unconditional $p(Y)$; and so the mean conditional $p(Y)$ —that is, OL_Y —equals the mean unconditional $p(Y)$. Specifically, $OL_Y = p(Y) = .25$ in Example 1, $.5$ in Example 2, and $.75$ in Example 3, even though in each case the roles are not correlated.

That OL_Y ignores shared Ns is only a problem if shared Ns influence perceptions of role overlap or similarity. Do they? Although there is consensus that similarity is a function of objects' common and distinctive features (Tversky, 1977), there is no consensus as to what features should enter into that function and how those features should be weighted. More to the point, there is no simple formula for computing the relative impact of shared Ys versus shared Ns.

However, one important predictor is the diagnosticity principle (Tversky, 1977). The diagnosticity principle states that features that reduce more uncertainty—that is, convey more information—about important classifications should have more weight in similarity judgments. This principle predicts that whether shared Ys have more or less weight than shared Ns depends in part on whether shared Ys are more or less diagnostic than shared Ns. What determines diagnosticity? According to information theory, the more diagnostic feature is the less likely feature. So, if Ys are relatively common, then shared Ys should be less diagnostic—and thus receive less weight—than shared Ns.

For example, if respondents are asked to use any words in their lexicon to describe different roles, then, of course, the few descriptors that might be shared by two roles are much more diagnostic than the thousands that are not applied to either role, and a measure such as OL_Y would be appropriate. In contrast, if respondents are asked (as some were in Study 3) to endorse 8 of 10 positive traits, then they are likely to focus more on which 2 traits they lack than on which 8 traits they have. Consequently, whether the positive features they lack in one role are also lacking in another role may be a potent determinant of perceived role similarity.

A second potential moderator of the impact of shared Ys and shared Ns is the tendency for judgments to be more influenced by events than nonevents (e.g., Brendl, Higgins, & Lemm, 1995; Fazio, Sherman, & Herr, 1982). Thus, saying "yes" or saying "no" typically has more weight than not saying "yes" or not saying "no." So, if people are asked

Table 5. *Trait $\times$ Situation Matrices Illustrating Properties of OL*

Traits															
T1	T2	T3	T4	T5	T6	T7	T8	T9	T10	T11	T12	T13	T14	T15	T16
Example 1															
Y	Y	Y	Y	N	N	N	N	N	N	N	N	N	N	N	N
Y	N	N	N	Y	N	N	N	N	Y	N	N	Y	N	N	N
Example 2															
Y	Y	Y	Y	Y	Y	Y	Y	N	N	N	N	N	N	N	N
Y	Y	N	N	Y	Y	N	N	Y	Y	N	N	Y	Y	N	N
Example 3															
N	N	N	N	Y	Y	Y	Y	Y	Y	Y	Y	Y	Y	Y	Y
N	Y	Y	Y	N	Y	Y	Y	N	Y	Y	Y	N	Y	Y	Y

Note: OL = index of features overlap; Y = yes; N = no.

to mark the features that apply to them and not mark those that do not, then (all else equal) Ys should have more weight than Ns. Conversely, if people are asked to mark the features that do not apply and not mark those that do, then Ns should have more weight than Ys. If people are asked to mark whether or not a feature applies—for example, if given a true–false or agree–disagree response format—then Ys and Ns should have equal weight.

So, if in the examples in Table 5 respondents actively asserted “I am ...” or “I am not ...” for each trait, then (all else equal) shared Ys and shared Ns should have equal weight. To make this concrete, imagine that trait T1 is “unassertive” in Example 1 and “assertive” in Example 3. The shared Y responses in Example 1 (saying “I am unassertive with mom and unassertive with dad too”) probably have the same psychological meaning and weight as the shared N responses in Example 3 (saying “I am not assertive with mom and not assertive with dad either”). Across all traits, Examples 1 and 3 are identical, except that the Ys and Ns are inverted. If the Ys and Ns have equal weight, this should not affect role similarity. Yet, whereas $OL_Y = .25$ (suggesting low overlap) in Example 1, $OL_Y = .75$ (suggesting high overlap) in Example 3. The implication is that ignoring shared Ns sometimes can lead OL_Y to ignore psychologically meaningful sources of role overlap.

Therefore, I recommend that researchers consider using the following, expanded measure of overlap: $S = w_Y OL_Y + w_N OL_N$, where OL_N is the mean probability that if a trait is marked N in one role it also is marked N in another role (i.e., the analog of OL_Y for N responses), and $w_Y + w_N = 1$. For example, the weights could be $w_Y = 2/3$ and $w_N = 1/3$ (if shared Ys are deemed twice as important as shared Ns), or they could be $w_Y = 1/3$ and $w_N = 2/3$ (if shared Ns are deemed twice as important as shared Ys). If shared Ns are deemed irrelevant ($w_Y = 1$, $w_N = 0$), then the formula would reduce to OL_Y . If there is no basis for weighting Ys more or less than Ns (as is true when Ys and Ns occur with similar frequencies, and both require active responses), then they would be given equal weight ($w_Y = w_N = 1/2$). In this case, S will give equivalent results as computing the mean r across all role pairs, which has already been used as a measure of role differentiation in several studies (Block, 1961; Donahue, Robins, Roberts, & John, 1993; Locke, 2002).

What Does H Measure That Other Indexes Do Not?

In sum, for many purposes computing H will be unnecessary. The critical information in a matrix could instead be summarized in terms of the number of rows,

the rates of endorsement (perhaps translated into bits in the form of H_{array}), and some measure of row similarity. However, for some purposes, H may still be necessary, such as when row independence is conceptualized as a property of the matrix, rather than a property of pairs of rows. Consider the self-descriptions of Jack and Jill in Table 6. For both Jack and Jill, all between role r s are zero, implying that we cannot predict the traits that will appear in one role from the traits that will appear in other roles. For Jill this is true, but for Jack it is not. Jack shows with his wife only those traits that he shows with both parents or does not show with either parent. Thus, knowing how Jack is with his mother and father, one can predict how Jack is with his wife perfectly. This difference between Jack and Jill is missed by all measures of pairwise association (such as correlation coefficients) and all procedures that operate on measures of pairwise association (such as factor analytic or multidimensional scaling techniques), but it is captured by H . For Jill, $H = 3$, whereas for Jack, $H = 2$. We are one bit more certain about how Jack will use a trait to describe himself.

But what can we conclude from the fact that Jill’s H is higher? A high H may indicate that Jill is carefully evaluating each Trait $\times$ Situation combination. Or instead it may indicate that Jill is responding randomly. A high H can result from either effortful or lazy responding. A low H can result from either effortful or lazy responding, too. A low H may indicate that a person is simply mindlessly endorsing the same traits in every situation. Or instead it may indicate that the person is trying to weave the sundry strands of a description into a meaningful pattern. For example, if you know that Jack and Jill are describing themselves to a marriage counselor and that T1–T4 are positive and T5–T8 are negative traits, you will realize that Jack is implying that the positive behaviors he exhibits with Jill (T1 and T2) are how he is with everyone, whereas the negative behaviors he exhibits with Jill (T7 and T8) are unique to his interactions with her. Jack’s H is lower because he is linking pieces of his self-description to tell a coherent (and self-serving) story.

Table 6. Trait $\times$ Situation Matrices Illustrating the Difference Between H and Measures of Pairwise Similarity

Situation	Trait							
	T1	T2	T3	T4	T5	T6	T7	T8
“Jill”								
R1: “with mom”	Y	Y	Y	Y	N	N	N	N
R2: “with dad”	Y	Y	N	N	Y	Y	N	N
R3: “with spouse”	Y	N	Y	N	Y	N	Y	N
“Jack”								
R1: “with mom”	Y	Y	Y	Y	N	N	N	N
R2: “with dad”	Y	Y	N	N	Y	Y	N	N
R3: “with spouse”	Y	Y	N	N	N	N	Y	Y

Note: H = measure of unpredictability; Y = yes; N = no.

Complexity Versus Differentiation

So, a higher H does not always reflect the amount of care and effort involved in a self-description, but does it at least reflect complexity? Although there is no consensus about how to define complexity, most definitions refer to a combination of differentiation and integration (Suedfeld, Tetlock, & Streufert, 1992). *Differentiation* refers to distinguishing elements within a stimulus domain. *Integration* refers to linking or organizing the differentiated stimuli. A complex self-description, therefore, would distinguish among a number of different traits and situations (differentiation) and link or organize those traits and situations in meaningful ways (integration). Although differentiation does tend to raise H , integration does not. Indeed, integration may typically lower H , as it did when Jack linked his roles as son versus husband.

However, just as H cannot distinguish thoughtful differentiation from thoughtless randomness, neither can H distinguish thoughtful integration from thoughtless simplification. Neither H nor any other mathematical formula can measure integration, because ultimately the input for a judgment of integration is not numeric but semantic. One cannot extract meaning from a matrix of Y and N without knowing what each Y and N means. Therefore, researchers have assessed integration by having human judges apply detailed coding manuals, such as those developed by Baker-Brown et al. (1992) or Woike (1989), to narrative protocols.

Conclusions and Future Directions

In conclusion, when H is computed on a matrix, it measures the unpredictability or independence of one part of a matrix from any other part. The problem with interpreting H is that it confounds three sources of unpredictability: unpredictability due to the number of rows in the matrix, unpredictability due to the independence among the rows, and unpredictability due to the rates of endorsement within roles. Therefore, when research has found effects of H , there is no way to know whether the effects were due to the number of rows, the independence of rows, the overall rates of endorsement, or an interaction of those variables. This is true of any research involving H , regardless of the particular hypotheses, measures, or analyses involved.

Given that these sources of unpredictability are conceptually and empirically distinct, I recommend distinguishing the influence of each source. To study the influence of rates of endorsement, one could test the effects of H_{array} or some variant thereof (see Footnote 2). To study the influence of the number of roles or situations considered, one could test the effects of R . To study the influence of role independence, one could test the effects of H , controlling for R and H_{array} . Alter-

natively, one could control the number of roles and the rate of endorsement experimentally, as long as adding such constraints would not directly or indirectly prevent the expression of important individual differences. In addition, to study the impact of independence of particular roles from other roles (or of a particular set of roles from another set of roles), one could directly compute and test the effects of H_{inc} . To study the impact of role overlap, as opposed to role independence, one could test the effects of a measure of association (such as OL, S , or r).

By clarifying some of the conceptual and methodological issues surrounding the use of H , this article hopefully will facilitate progress on the substantive issues. For example, some (typically personality psychologists) have claimed that integrated, consistent selves are healthy, contributing to positive affect, adjustment, and role satisfaction (e.g., Block, 1961; Donahue et al., 1993). Others (typically social psychologists) have claimed differentiated, multifaceted selves are healthy, buffering the impact of negative life events (e.g., Linville, 1985, 1987), and offering a rich repertoire of behavioral responses (e.g., Sande, Goethals, & Radloff, 1988). Perhaps one reason the existing literature on self-complexity is inconclusive is that there are no consistent effects at the level of H (i.e., overall uncertainty). However, clearer patterns may emerge once we distinguish among the effects of the number of traits endorsed, the number of roles articulated, the interdependencies among those roles, and the interactions of these variables.

This article targeted self-complexity research only because H has been applied most often to that topic. The problems and solutions presented here pertain to any data in which a set of attributes is applied to a set of entities, which is a data structure pervasive in social information processing research. Consider research on the self. Any time there are data on multiple attributes (such as different motives or different attitudes or different behaviors) across multiple entities (such as the self with different individuals, with different groups, or in different states of mind), the data can be organized as the type of matrix analyzed in this article. The same is true of data on how a person conceptualizes attributes of another individual across different situations. In a simple extension to social cognition about a group (instead of an individual), the entities could be group members (instead of roles); in this case, a high H implies that attributes of one member may not predict attributes of other members. In a further extension, the entities could be different groups and the attributes themselves could be different individuals; in this case, a high H means that membership information on one individual or group may not generalize to other individuals or groups.

Thus, mental representations of social information can often be modeled as Attribute $\times$ Entity matrices, and

H is just one of many structural properties that can be computed on such matrices (Scott, 1969). However, we must be cautious about reifying these properties—assuming that because we can compute an index of some structural property (such as H), it actually is a psychologically meaningful and unitary variable (such as “complexity”), and not bother to even report the simpler elements (such as the number of rows or the rates of endorsement within rows). Instead, we always should consider how a structural index might be influenced by several simpler elements, and how those elements may each be related to distinct underlying psychological processes that have distinct causes and effects. Otherwise, by only testing the structural properties of information processing, a potentially productive line of research may produce only confusion and frustration.

References

- Attneave, F. (1959). *Applications of information theory to psychology*. New York: Holt-Dryden.
- Baker-Brown, G., Ballard, E. J., Bluck, S., DeVries, B., Suedfeld, P., & Tetlock, P. E. (1992). The conceptual/integrative complexity scoring manual. In C. Smith (Ed.), *Handbook of thematic analysis* (pp. 401–418). Cambridge, England: Cambridge University Press.
- Beck, A. T., Steer, R. A., & Brown, G. K. (1996). *BDI-II Manual*. San Antonio, TX: Psychological Corporation.
- Block, J. (1961). Ego identity, role variability, and adjustment. *Journal of Consulting Psychology*, 25, 392–397.
- Brendl, C. M., Higgins, E. T., & Lemm, K. M. (1995). Sensitivity to varying gains and losses: The role of self-discrepancies and event framing. *Journal of Personality & Social Psychology*, 69, 1028–1051.
- Brewer, M. B., & Lui, L. (1984). Categorization of the elderly by the elderly: Effects of perceiver's category membership. *Personality and Social Psychology Bulletin*, 10, 585–595.
- Broadbent, D. E. (1958). *Perception and communication*. London: Pergamon.
- Conway, M., & White-Dysart, L. (1999). Individual differences in attentional resources and self-complexity. *Social Cognition*, 17, 312–331.
- Davies, M. F. (1996). Self-consciousness and the complexity of private and public aspects of identity. *Social Behavior and Personality*, 24, 113–118.
- Dixon, T. M., & Baumeister, R. F. (1991). Escaping the self: The moderating effect of self-complexity. *Personality and Social Psychology Bulletin*, 17, 363–368.
- Donahue, E. M., Robins, R. W., Roberts, B. W., & John, O. P. (1993). The divided self: Concurrent and longitudinal effects of psychological adjustment and social roles on self-concept differentiation. *Journal of Personality and Social Psychology*, 64, 834–846.
- Fazio, R. H., Sherman, S. J., & Herr, P. M. (1982). The feature-positive effect in the self-perception process: Does not doing matter as much as doing? *Journal of Personality and Social Psychology*, 42, 404–411.
- Garner, W. R. (1962). *Uncertainty and structure as psychological concepts*. New York: Wiley.
- Hampson, S. E., Goldberg, L. R., & John, O. P. (1987). Category-breadth and social-desirability values for 573 personality terms. *European Journal of Personality*, 1, 241–258.
- Kreiner, P., Soldz, S., Berger, M., Elliott, E., Reynes, J., Williams, C., et al. (2001). Social indicator-based measures of substance abuse consequences, risk, and protection at the town level. *Journal of Primary Prevention*, 21, 339–365.
- Linville, P. W. (1982). The complexity-extremity effect and age-based stereotyping. *Journal of Personality and Social Psychology*, 42, 193–211.
- Linville, P. W. (1985). Self-complexity and affective extremity: Don't put all your eggs in one cognitive basket. *Social Cognition*, 3, 94–120.
- Linville, P. W. (1987). Self-complexity as a cognitive buffer against stress-related illness and depression. *Journal of Personality and Social Psychology*, 52, 663–676.
- Linville, P. W., & Jones, E. E. (1980). Polarized appraisals of out group members. *Journal of Personality and Social Psychology*, 38, 689–703.
- Locke, K. D. (2002). Are descriptions of the self more complex than descriptions of others? *Personality and Social Psychology Bulletin*, 28, 1094–1105.
- Magurran, A. E. (1988). *Ecological diversity and its measurement*. Princeton, NJ: Princeton University Press.
- Miller, G. A. (1953). What is information measurement? *American Psychologist*, 8, 3–11.
- Morgan, H. J., & Janoff-Bulman, R. (1994). Positive and negative self-complexity: Patterns of adjustment following traumatic versus non-traumatic life experiences. *Journal of Social and Clinical Psychology*, 13(1), 63–85.
- Norman, W. T. (1967). *2800 personality trait descriptors: Normative operating characteristics for a university population*. Ann Arbor: University of Michigan, Department of Psychology.
- Rafaeli-Mor, E., Gotlib, I. H., & Revelle, W. (1999). The meaning and measurement of self-complexity. *Personality and Individual Differences*, 27, 341–356.
- Rafaeli-Mor, E., & Steinberg, J. (2002). Self-complexity and well-being: A research synthesis. *Personality and Social Psychology Review*, 6, 31–58.
- Rhodewalt, F., Madrian, J. C., & Cheney, S. (1998). Narcissism, self-knowledge organization, and emotional reactivity: The effect of daily experiences on self-esteem and affect. *Personality and Social Psychology Bulletin*, 24(1), 75–87.
- Rhodewalt, F., & Morf, C. C. (1995). Self and interpersonal correlates of the Narcissistic Personality Inventory: A review and new findings. *Journal of Research in Personality*, 29(1), 1–23.
- Rhodewalt, F., & Morf, C. C. (1998). On self-aggrandizement and anger: A temporal analysis of narcissism and affective reactions to success and failure. *Journal of Personality and Social Psychology*, 74, 672–685.
- Rosenberg, M. (1965). *Society and the adolescent self-image*. Princeton, NJ: Princeton University Press.
- Sande, G. N., Goethals, G. R., & Radloff, C. E. (1988). Perceiving one's own traits and others': The multifaceted self. *Journal of Personality and Social Psychology*, 54, 13–20.
- Scott, W. A. (1969). Structure of natural cognitions. *Journal of Personality and Social Psychology*, 12, 261–278.
- Solomon, A., & Haaga, D. (1993). Sociotropy, autonomy, and self-complexity. *Journal of Social Behavior and Personality*, 8, 743–748.
- Suedfeld, P., Tetlock, P. E., & Streufert, S. (1992). Conceptual/integrative complexity. In C. Smith (Ed.), *Handbook of thematic analysis* (pp. 393–400). Cambridge, England: Cambridge University Press.
- Tversky, A. (1977). Features of similarity. *Psychological Review*, 84, 327–352.
- Woike, B. A. (1989). *Categories of cognitive complexity: A scoring manual*. Unpublished manuscript. East Lansing: Michigan State University.
- Woolfolk, R. L., Novalany, J., Gara, M. A., Allen, L. A., & Polino, M. (1995). Self-complexity, self-evaluation, and depression: An examination of form and content within the self-schema. *Journal of Personality and Social Psychology*, 68, 1108–1120.

Copyright of Personality & Social Psychology Review is the property of Lawrence Erlbaum Associates and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.