

“Analyzing ‘Spooky Action at a Distance’ Concerning Brand Logos”

AUTHORS	Robert Schorn  https://orcid.org/0000-0002-9580-0558 Gottfried Tappeiner Janette Walde
ARTICLE INFO	Robert Schorn, Gottfried Tappeiner and Janette Walde (2006). Analyzing ‘Spooky Action at a Distance’ Concerning Brand Logos. <i>Innovative Marketing</i> , 2(1)
RELEASED ON	Friday, 10 March 2006
JOURNAL	"Innovative Marketing "
FOUNDER	LLC “Consulting Publishing Company “Business Perspectives”



NUMBER OF REFERENCES

0



NUMBER OF FIGURES

0



NUMBER OF TABLES

0

© The author(s) 2026. This publication is an open access article.

Analyzing 'Spooky Action at a Distance' Concerning Brand Logos

Robert Schorn, Gottfried Tappeiner, Janette Walde

Abstract

Attractive brand logos and packaging elements are gaining more and more importance as decisive competitive advantages in a world flooded by stimuli. Based on the assumption that there exists a kind of collective knowledge beyond individual experience, the authors found that in respect to logos humans are more likely to respond to stimuli if many people in other parts of the world do or did know them, even though they personally are not consciously familiar with the logos. An improved favorability of 20% for original symbols versus comparable control symbols can be regarded as a solid competitive advantage. This benefit regarding brand logos was analyzed by means of latent class models. Additionally, the heterogeneity in the participant's characteristics as well as the heterogeneity in the analyzed symbols were incorporated by means of random and fixed effects models. Furthermore, this effect was shown to be neither culture-specific nor linked to age, gender, level of extraversion, and education of the participants.

Key words: limited dependent variables model, random and fixed effects, brand logos, collective knowledge, consumer behavior.

1. Introduction

As visual representations of products, companies or organizations, brand logos are an important factor. Logos easily memorized and recognized are a significant competitive advantage (Aaker, 1996; Henderson & Cote, 1998; Keller, 2002). Psychologists concerned with the field of learning investigate how brands are perceived, memorized and retrieved (Kuehn, 1962; Ratchford, 2001). The term "learning" here refers to all processes that alter an organism in such a way that it will be able to react differently – even if that is only a little faster – in a comparable situation (Lefrancois, 1972; Mazur, 2002). Marketing experts use the findings in the field of psychology of learning by trying to activate such learning processes within consumers so that they memorize or recognize products – or rather, their distinctive features – more easily (Williams, 1990). It is widely adhered to that learnt information is stored in the form of memory tracks in the brain of each individual. Whereas earlier approaches of brain research proposed a tight local storage of knowledge, more recent ones propose storage of knowledge in the form of propositional networks within the brain (Anderson, 2000).

Besides these classical approaches, there is a row of unconventional theories, some of which postulate a kind of supra-individual knowledge. Of these, C.G. Jung's concept of the collective unconscious is most known (Jung, 1981); one of the more recent theories in this field comes from Rupert Sheldrake. He postulates that everything that was learned by a variety of people is easier to learn for people who have not been consciously familiar with it than something comparably difficult to learn but previously unknown to humanity. According to Sheldrake, this is due to "morphic fields" that all humans have access to. These fields contain some sort of cumulative memory; they become more efficient by repeated access and they do not deteriorate with spatial and/or temporal distance. According to Sheldrake's theories, the brain functions rather as a transmission or receiving device than as a place for storing information (Sheldrake, 1981; 1988). If this proved true, it would have enormous effects on designing brand logos, packaging, product and design elements, etc., it could help to improve the economic effectiveness of these elements. Regarding brand logos this would mean that logos which were familiar to many people in earlier times but have fallen into oblivion, or logos which are popular in parts of the world different from where they are to be employed should have strong morphic fields in terms of Sheldrake's hypothe-

sis, and thus should be easier to recognize, process and memorize than similar symbols that have not been known anywhere or at any times.

The idea of collective knowledge that is accessible for all humans – even if by the indirect way of the unconscious – may seem unfathomable for humans used to rational ways of thinking. Motivated by this unconventional idea, we conceptualized an experiment that was generally suitable for either falsifying or corroborating Sheldrake's hypothesis. This article presents a controlled experiment with which hypotheses based on Sheldrake's ideas can be statistically tested. The hypotheses examined are:

H1: People are more likely to respond to symbols that were or have been widely known but that they are not consciously familiar with, as compared to similar symbols (control symbols) that have been artificially created for the test.

'People being more likely to respond' means that facilitatory or repetition effects (priming effects) such as in the case of experiments considering implicit memory can be found when presenting symbols that were or have been known by many people. Such priming effects cause subconsciously available material (in this case symbols) to be recognized faster, processed better and memorized more easily, as well as to be perceived as more familiar and appealing.

Our second hypothesis is based on Sheldrake's postulate which says that collective knowledge does not decrease with spatial distance.

H2: Regarding symbols that were or have been very popular, people who are not consciously familiar with them are not more likely to respond if they originate in the same region as the symbols tested, compared to people who originate in other regions where these symbols have been less known or not known at all.

This means that the researched phenomenon is equally strong everywhere in the world, i.e. not dependent on location and therefore not stronger in or near the region of origin and popularity of a symbol as compared to other parts of the world.

To consider whether people differ in responding to popular symbols as to artificially created symbols, we control for characteristics of the participants such as age, gender, levels of education, or extraversion and test the impact of these variables.

H3: Age, gender, levels of education, and extraversion do not have any influence on the fact that people are more likely to respond to symbols that were or have been very popular but are not consciously known to the participants compared to artificially created control symbols.

This means that the researched phenomenon is independent of the characteristics of a person.

2. Method

In order not to create methodological artifacts, our test method had to fulfil two requirements:

- The method must be suitable for detecting even minute differences in the perception of symbols.
- The method must have been tested with other, comparable experiments, and it must be scientifically acknowledged.

In keeping with these requirements, methods for assessing implicit memory provided valuable clues and parallels.

2.1. *Measuring Implicit Knowledge*

Implicit memory tests are indirect procedures that enable the detection of memory without creating consciousness of the memory involved (Dienes & Berry, 1997; Richardson-Klavehn

& Bjork, 1988). According to Richardson-Klavehn and Bjork (1988), implicit memory tests are classified into the following four categories:

- tests of conceptual, factual, lexical, and perceptual knowledge,
- tests of procedural knowledge (i.e. skilled performance, problem solving),
- measures of evaluative response, and
- other measures of behavioral change, including neurophysiological response and conditioning measures.

Methods for measuring evaluative response (c) are based upon the assumption that recurrent contact with a certain stimulus alters the affective attitude, as well as the cognitive judgment concerning the stimulus. One effect that is based upon this assumption is the mere exposure effect, which is well known as a factor influencing consumer behavior (Janiszewski, 1990; 1993; Janiszewski & Meyvis, 2001). The mere exposure effect was documented for the first time by Zajonc (1968) and has been researched in over two hundred experiments since (for an overview see Bornstein (1989)). It consists in the fact that familiar people or objects are more willingly accepted and preferred than less familiar people or objects.

Kunst-Wilson and Zajonc (1980) attempted to find out whether the mere exposure effect can also be proved if the period of exposition of the test participants to the test stimuli is lowered to a level where the recognition performance decreases to a random level. In a comparison of two irregular octagons (black color on white background), of which one was presented for a very short time beforehand and the other was not, the respondents had to decide (1) which one they liked better and (2) which one had been shown previously. The responses to the question concerning which of the two octagons had been shown previously approached a random distribution (48% guessed correctly), whereas for the question concerning which one of the two octagons was liked better (affective response) 60% of the responses coincided with the previously presented octagon. The authors concluded from their experiments that people develop certain preferences for objects without consciously recognizing them, and that affective discrimination could be possible without participation of the cognitive system (Kunst-Wilson & Zajonc, 1980).

This effect that had been found by Kunst-Wilson and Zajonc was replicated and amended by different experimenters (e.g., Bonnano & Stillings, 1986; Seamon et al., 1983; Seamon et al., 1984). Thus it was found that affective preferences cannot only be assessed through appraisals concerning like or dislike, but also through appraisals regarding familiarity (Bonnano & Stillings, 1986; Mandler, 1980). Mandler et al. (1987) arrived at even more extensive conclusions. The results of their experiment, which only differed from the experiment of Kunst-Wilson and Zajonc (1980) in the kind of questions asked, indicate that appraisals concerning the degree of brightness or darkness reach the same values as assessments concerning preferences, and thus they concluded "that any relevant dimensions can be related to the activation of the stimulus representation" (Mandler et al., 1987, p. 648).

Such experiments show that subjective judgments or item selections concerning emotive impressions (prefer), personal appraisals or intuitive appraisals guarantee a hit ratio far beyond the random level, even if such a discrimination does not take place on the conscious level, i.e. without retracing it consciously and explaining it rationally. The fact that questions regarding recognition led to results that did not differ from random expectations is explained with the argument that with such questions the critical-rational attitude of the respondent, as well as the fear of doing something wrong or to fail may lead to tension and stress – even if these are not consciously perceived – and thus block the utilization of resources that are embedded in the intuitive "sensing" of subconsciously stored knowledge (Moreland & Zajonc, 1977; Perrig et al., 1993; Reber, 1996).

Similarly to our experiment, the tests described tried to show that stimuli that lie below the perception threshold can cause learning effects. The main difference is that in all the experiments cited above the participants were exposed to stimuli beforehand, even though they were very weak and lay below the psychological perception threshold, while our test was aimed at detecting effects of "morphic fields" without any prior exposure of the participants to a test stimulus.

2.2. Experimental Implementation

Our experiment was aimed at investigating whether a systematic bias concerning the appraisal of symbols can also be found if the participants have not, as in the experiments mentioned above, been exposed to the symbols earlier, but if the symbols tested were or have been very familiar to a large group of people distinct from that of the participants. The method we applied was the same as in the experiment conducted by Mandler et al. (1987), with the sole difference that our respondents were not exposed to the stimuli beforehand. Instead, we used symbols that were unfamiliar to our respondents but were or have been known very well in earlier times or in other regions.

2.3. Material

We used two classes of test stimuli. The first one consisted of political, religious, and economic symbols such as flags, trademarks, emblems, faces of people, coinages, etc. that were very popular earlier but have fallen into oblivion nowadays, or that are familiar to many people in countries our respondents did not originate in, such as the Chinese Coca-Cola symbol, Indian trademarks, or Far Eastern religious symbols. For each of these 20 symbols, a corresponding control symbol was created; these differed marginally from the original symbols, and it was made sure that they were not or have not been known anywhere (Figure 1).

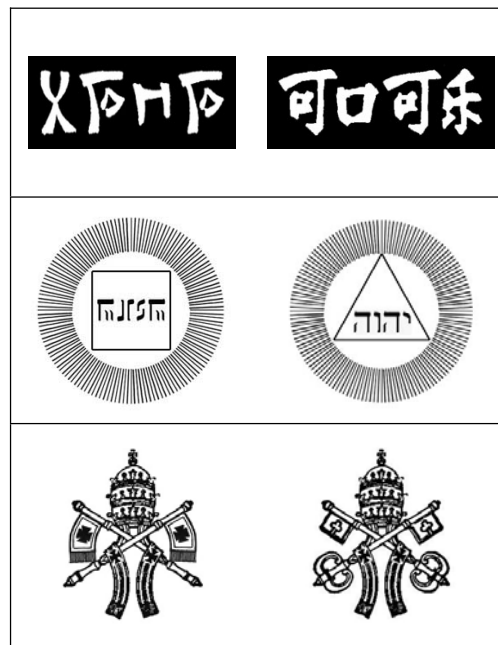


Fig. 1. Original symbols with corresponding control symbols

When creating the control symbols, we conducted seven pre-tests with more than 200 participants. The pairs of symbols were discussed with the participants and the participants were asked to indicate whether they perceived one of the symbols of each pair as less credible or real than the other one. We employed only those symbols in our experiment that were not classified as less credible than the original symbols in the pre-tests.

The second category of test stimuli consisted of 20 very frequently encountered Russian words in Cyrillic characters. Many generations of people in the countries of the former Soviet Union have been familiar with these words (in 1991, the former Soviet Union had 288 million inhabitants). An expert created a control word for each of these words through permutations of the letters of the original word, so that the newly created word had no semantic meaning (Figure 2). While in

respect to the symbols pertaining to category 1 critics may object that the original symbols are equipped with specific perceptive-intrinsic characteristics and were originally created due to these reasons and thus are bound to prevail in the experiment versus the control symbols, this objection cannot be made for the category “Russian words”. The argument that certain words are used more frequently than others simply because of their attractiveness in a written form, and that therefore they would have to prevail in the experiment versus inexistent words, is not valid, because writing was invented much later than language. Owing to the fact that the control words contained the same letters as the original words, it cannot be argued either that certain words were chosen more frequently only because they consisted of “prettier” letters.

ОЦНЛСЕ	СОЛНЦЕ
ачплъе	печаль

Fig. 2. Original Russian words with corresponding control words

For analyzing the selection of the Russian words, a range of control variables were included (positive or negative semantic content, adjectives/nouns/verbs, number of letters of the word and the corresponding control word (3-11 letters), positions of the original and control words in the experiment (left/right), as well as pronounceable and unpronounceable control words).

2.4. Procedure and Design

The symbols of category 1 were presented in pairs of two, next to each other (authentic symbol and corresponding control symbol), for the duration of 10 seconds each. In order to avoid biases caused by a possible tendency of the respondents to prefer one side or the other, in half of the cases the authentic symbols were presented on the left-hand side, in the other half of the cases they were shown on the right-hand side of the control symbols. In keeping with the type of questions posed by Mandler et al. (1987), the participants were asked to select the symbol they perceived as containing more spirit; in the pre-tests we conducted with 80 participants for finding out the most suitable question, this one had proved best. A question can be regarded as suitable if the respondents are indifferent towards a stimulus in the case of a conscious decisive process, such as for example the question concerning brightness and darkness of octagons of the same color (black) in the experiment of Mandler et al. (1987). Only if the respondents cannot take a conscious decision, the possibility for decisions influenced by unconscious processes opens up. Only if both octagons have exactly the same color (as in the experiments of Kunst-Wilson and Zajonc (1980) or Mandler et al., (1987)), and thus no conscious decision can be taken, the voice of the unconscious can be heard.

Additionally, the participants were asked to indicate the meaning of a symbol in case they were familiar with it. This served as a control question in order to recognize those cases where a symbol was consciously known to the respondent, and to subsequently exclude this pair of symbols for the given participant. Besides those symbols to which the participants could attribute the correct meanings, also the symbols to which the participants attributed a meaning that coincided with a similar – but not identical – symbol were excluded.

The procedure followed for the Russian words in Cyrillic writing was the same as that for the symbols, only that in this case the respondents were asked to select the word which had more visual appeal to them. In respect to the kind of the problem this question could have been researched as well with symbols. The reason why this was not done is because it could be argued by critics that the real symbols were selected more frequently than the control symbol not due to col-

lective knowledge, but rather that these symbols have a particular effect on people and were created for this reason and prevail due to their general favourability versus other symbols.

After testing the stimuli, which were presented according to the two categories, the participants were asked to answer 24 questions concerning extraversion in the Eysenck Personality Inventory (EPI, Form B) (see Eysenck 1964), as well as to indicate their age, gender, level of education, and the country they resided in. Furthermore, the date and time when the experiment had been conducted were registered.

The information about the country of origin served for creating the variable "distance". This variable indicates whether the interviewee and the symbol originate from the same region or not. All participants and all symbols were assigned to one of the following geographical categories: Europe (without Great Britain), Great Britain (this was introduced as a separate category because certain symbols were not familiar in Great Britain as opposed to continental Europe), America, Asia, and Australia. By means of this variable it was to be verified whether participants who originated in the same region as a certain symbol selected this symbol (in comparison with its control symbol) significantly more often than participants from other regions.

The data obtained in the course of the experiment indicating the date and time of the participation of each respondent served as the basis for creating the variable "rank". This was a control variable created for checking whether the order in which the respondents participated in the experiment had any influence on their answers.

It was calculated how often each participant had chosen an authentic stimulus as the one with more spirit in category 1 or as the one having more visual appeal in category 2. This "morphic ability" of each participant for each of the two categories of stimuli entered our analyses as a variable.

The experiment was conducted in an offline and an online version. In the case of the offline version, the test stimuli were projected onto a wall with a video or overhead projector, and the participants were asked to indicate their judgments on response sheets. Two hundred and seventy-two participants were tested in the offline version. In order to include participants from as many different countries as possible, also an online version of the experiment was created, so that it could be conducted at any computer with access to the internet. Four hundred and eight persons participated in the experiment via the online version. Thus the sample comprised a total of 680 persons from Europe, Asia, Australia, and North and South America.

3. Statistical Modeling

The present analyses shall answer two central questions in accordance with the hypotheses mentioned above:

Are the stimuli that were very popular in the past or are still well known today but not consciously known to the participants selected significantly more often than the control stimuli? In other words, we would like to investigate whether there is evidence that allows for the conclusion that there is indeed some sort of collective knowledge that people have access to or that influences their decision-making.

In case there is evidence for the fact that people have access to collective knowledge, does this ability correlate with characteristics of the respondents (age, gender, education, extraversion, distance, and rank) or is it independent of such characteristics?

In order to be able to answer these questions and to get a valid level of significance (in contrast to a series of independent *t*-tests), we employed panel econometric models (Baltagi, 2001; Hsiao, 1999) accounting for heterogeneity between the stimuli and the participants. On the one hand, we used the information as to whether the participants had chosen the left (coded as "0") or the right (coded as "1") symbol/word of the presented pair as dependent variable. This served for testing whether the original stimuli were chosen significantly more often than the control stimuli. As independent variable we used the information whether the original stimulus was presented at the right-hand side (coded "1") or at the left-hand side (coded as "0"). If this variable significantly influenced the choice of the respondents, we could answer our first question positively.

On the other hand, we coded the choice of the authentic stimulus as “1” and of the control stimulus as “0” and used this coding scheme as dependent variable. Age, gender, rank, distance, extraversion, and education were included as independent variables. Morphic ability for the corresponding stimulus category (symbols or words), which was not included as dependent variable in the analysis, was used as additional input variable. By means of these independent variables we attempted to estimate each participant’s chance of choosing the authentic stimuli.

A suitable method for modeling limited dependent variables is logistic regression. Among other things a basic assumption for a valid logistic regression analysis is that the coefficients measuring the effects of the independent variables, as well as the level-parameters of all respondents, are equally large. Also, these parameters are independent from which stimulus is involved (homogeneity, poolability). The relationship between them is described by equation 1.

$$y_{ij} = \begin{cases} 1 & \text{if } y_{ij}^* > 0 \\ 0 & \text{if } y_{ij}^* \leq 0 \end{cases}, \text{ where } y_{ij}^* = \alpha + \mathbf{x}_{ij}' \boldsymbol{\beta} + \varepsilon_{ij}, \quad (1)$$

and where y_{ij} is the value of the dependent variable for individual i ($i = 1, \dots, 680$) and stimulus j ($j = 1, \dots, 20$), y_{ij}^* is the corresponding value of the latent variable, ε_{ij} is the unobservable noise parameter, $\mathbf{x}_{ij}$ is the column vector with the values of the independent variables from individual i and for stimulus j as its components, and $\alpha, \boldsymbol{\beta}$ are the parameter and parameter vector respectively we are interested in and want to estimate.

In case the assumptions of homogeneity and poolability are violated, analyses not controlling for heterogeneity run the risk of obtaining biased results. However, panel econometric models allow for the consideration of differences between the stimuli or the participants. In the case of the stimuli, this means that we should account for differences between the stimuli regarding their age, shape, etc. These differences are modeled by means of a level-parameter that may have different values depending on the stimulus or participant. If it is assumed that this coefficient is a fixed parameter for each stimulus (fixed effects model, cf. equation 2), no effects of variables that are constant for all participants (e.g. number of letters) can be estimated. Consequently in order to test the effect of e.g. gender, age, level of education we can only control the heterogeneity between stimuli in the model.

$$y_{ij}^* = \alpha_j + \mathbf{x}_{ij}' \boldsymbol{\beta} + \varepsilon_{ij}, \quad (2)$$

where α_j denotes the ‘unobservable’ stimulus-specific effect. In the case of the fixed effects model, the α_j are assumed to be fixed parameters. For the random effects model, the α_j can be assumed random: $\alpha_j \sim N(0, \sigma_\alpha^2)$.

When estimating the effect of control variables like the number of letters, the order of presentation of the words, the semantic content, etc. for the Russian words we accounted for the differences between the participants.

We also calculated a logistic regression for each stimuli pair (j). This is the model with the highest number of degrees of freedom:

$$y_i^{(j)*} = \alpha^{(j)} + \mathbf{x}_i^{(j)}' \boldsymbol{\beta}^{(j)} + \varepsilon_i^{(j)}. \quad (3)$$

Which of these models could be regarded as the correctly specified one was established by means of statistical tests (e.g. Hausmann specification test). Beforehand, however, both the quality and the validity of the models had to be checked.

The quality of the models was measured via the hit ratio. For example by means of the models we were able to predict the probability for an individual to choose the authentic stimuli. If this predicted probability was greater than 50%, we classified the individual into class 1 (likely to choose the authentic stimuli), otherwise the individual was classified into class 0. In the end, we compared the predicted class with the decision each participant made. The hit ratio in this case is the

ratio of the number of correctly classified participants to all participants. The hit ratio obtained by our model had to be compared with the random chance of classifying a person correctly. This random chance is the ratio of the number of participants of the larger class to all participants. Concerning the quality of the model, this means that if the hit ratio of the model was significantly higher than the hit ratio by chance, we could conclude that the independent variables were able to partially explain why one participant was more likely to choose the real stimuli than another participant.

4. Results

4.1. Symbols

We analyzed how often the respondents preferred an authentic symbol to its corresponding control symbol, and whether this percentage was significantly higher than 50%. As demonstrated by the results in Table 1, in 15 out of 20 cases the original symbol was selected significantly more often. Overall, the original symbols were chosen in 61.7% of the cases. Only for three symbols the control symbols prevailed in the comparison.

Table 1

Choices of the respondents concerning the symbols

Symbol	+	-	+%	Sig. (1-tailed)
Doordarshan	593	80	88.1	.000
SS-bolts	376	119	76.0	.000
Om	487	156	75.7	.000
PX	403	137	74.6	.000
Sun Yat-sen	462	193	70.5	.000
Coca-Cola	469	200	70.1	.000
Trinity	460	211	68.6	.000
Confederate Battle Flag	379	186	67.1	.000
Ku Klux Klan	449	220	67.1	.000
Islam	411	231	64.0	.000
Fascies	409	258	61.3	.000
China flag	352	223	61.2	.000
Shiva Lingam	407	262	60.8	.000
Bangladesh flag	397	273	59.3	.000
Key of Petrus	379	265	58.9	.000
Tata	336	333	50.2	.469
Roma	335	339	49.7	.454
Chubais	271	395	40.7	.000
Alexander the Great	260	406	39.0	.000
Jerusalem cross	252	412	38.0	.000
Total	7887	4899	61.7	.000

The column “+” in Table 1 indicates in absolute figures how often the authentic symbol was preferred to its control symbol, whereas the column “-” indicates in absolute figures how often the control symbol prevailed. The column “+%” indicates the percentage of the cases in which the authentic symbol was preferred. For example, in the case of the symbol “Doordarshan” this means that the authentic symbol was selected 593 times, i.e. in 88.1% of all cases, whereas its corresponding control symbol was chosen 80 times only. The column “Sig. (1-tailed)” shows the level of significance of the binomial test with the alternative hypothesis at a percentage higher than 50%.

However, measuring the significance of the individual tests is not suitable for calculating the significance of the test as a whole. Hence, we employed a random effects model with the inde-

pendent variable “position of the correct symbol (left/right-hand side)” for testing whether the original stimuli were picked significantly more often than the control stimuli. The hit ratio of the random effects model was 61.7%. Compared to the hit ratio by chance (50%), this was a significant improvement of the model performance. As the log likelihood test for the necessity of considering heterogeneity was highly significant ($p < 0.000$), a stacked logistic regression model was not adequate. The odd of making the cross at the right-hand side was 2.214, the corresponding 95%-confidence interval was [2.048, 2.394]. Hence, the probability of choosing the authentic stimulus was more than twice as high as that of selecting the artificially created control stimulus.

4.2. Words

The results for the test with the Russian words were analyzed analogously to the results of the symbols-test: the column “+” in Table 2 shows how often the authentic word was preferred to its control word; the column “-” indicates how often the control word prevailed. The column “+ %” illustrates the percentage of the cases in which the authentic word was chosen. The column “Sig. (1-tailed)” indicates the probability of type I error of the binomial test (alternative hypothesis: $\pi > 50\%$). In 11 of the 20 cases, the original words were chosen significantly more often than their corresponding control words, for five words there was no significant difference ($p > .05$), and in four cases the control words prevailed. The last row of Table 2 demonstrates that the authentic words were selected in 56.4% of all cases.

Table 2

Choices of the respondents concerning the Russian words

Word	+	-	+%	Sig. (1-tailed)
Good	474	169	73.7	.000
War	448	208	68.3	.000
To kill	446	210	68.0	.000
Peace	428	205	67.6	.000
Beautiful	428	216	66.5	.000
Disease	436	220	66.5	.000
Water	406	234	63.4	.000
Sun	393	266	59.6	.000
Life	371	278	57.2	.000
Aggressive	366	283	56.4	.001
Mourning	366	295	55.4	.003
To live	342	318	51.8	.186
To help	327	326	50.1	.500
Pain	322	337	48.9	.293
Death	316	337	48.4	.217
Love	306	340	47.4	.097
To die	301	350	46.2	.030
Bad	299	349	46.1	.027
To give	290	361	44.5	.003
To hate	275	381	41.9	.000
Total	7340	5683	56.4	.000

In this case it is even more difficult to calculate the overall result of the test when basing ourselves on the data in Table 2. Therefore we used the random effects model for testing whether the participants chose the original stimuli more often than the control stimuli. This time the hit ratio of the model was 56.01%. Compared to the hit ratio by chance (50%), this modeling approach has some power to explain the variability of the dependent variable. The log likelihood test

for the necessity of considering heterogeneity again was highly significant ($p < 0.000$). The odd of choosing the right-hand side was 1.723, the corresponding 95%-confidence interval was [1.568, 1.8934]. The probability of choosing the original Russian word was more than 1.5 times as high as that of selecting the artificially created control word.

In our second analysis we wanted to evaluate whether the observed personal characteristics of the participants have a significant effect on the probability of preferring an original stimulus to its control stimulus. This includes the variables of age, gender, education, extraversion, distance, rank, as well as morphic ability for the corresponding stimulus category (symbols or words) that is not included as the dependent variable in the analysis. This time the dependent variable was the probability for each respondent to choose the original stimulus. If s/he indeed selected the authentic stimuli, the variable was coded as "1", otherwise as "0".

We modeled the limited dependent variables for both categories of stimuli separately. The analysis of the Russian words additionally included the following control variables: number of letters, position (left/right), type of word (noun, adjective, verb), semantic content (positive, negative), pronounceability of the control word, and order of presentation.

Another factor taken into consideration was the following: research concerning implicit learning demonstrates that people are able to learn implicit rule-based knowledge (tacit knowledge) without being able to explicitly express it or being consciously aware of it at all (Polanyi, 1966; Reber, 1996). Words pertaining to human language are subjected to a system of rules, and one could object that this system could be implicitly perceived by the participants in the context of the experiment involving Russian words; thus, it could have caused improved discrimination between the authentic words and the control words favoring the prior ones. The existence of such a connection was statistically tested.

The calculation of the hit ratio, however, showed that classification according to the simple logistic regression model and the fixed and random effects models did not produce any better results than classification according to the most frequently encountered class (prior probability).

Furthermore, for each pair of stimuli an independent model was estimated. In the case of the symbols, three out of 20 (cf. Table A-3), and for the Russian words two out of 20 models showed an enhanced hit ratio in comparison to the prior probability ($p < .05$) (cf. Table A-4). However, in case of the symbols this is a hardly significant number, and in case of the Russian words, it is an amount that falls into the acceptance region of the null hypothesis, thus indicating that those were only random improvements.

This second set of analyses allows for the conclusion that the investigated characteristics are not suitable for modeling the probability of choosing the authentic stimulus, and that they do not have an influence on this probability.

5. Discussion and Conclusions

The results of our experiments absolutely support hypothesis 1: both in the test with the symbols and that with the Russian words, the original stimuli were selected significantly more often than the artificially created control stimuli. The result concerning the symbols is even more unambiguous than that regarding the Russian words. This can probably be attributed to the fact that the two approaches implicitly analyzed different contents: while with the symbols it was a comparison between an authentic stimulus and a completely fictitious stimulus, in the case of the Russian words both stimuli (although the control stimulus did not bear any meaning) were made up of existing Cyrillic letters, i.e. the respondents assessed the word faces. It could be argued that it is easier to perceive the difference between two symbols than that between two word faces; this would be worth further in-depth investigation.

We would like to state that in our experiment we did not try to investigate the reasons for the effect we observed. Jung's collective unconscious or Sheldrake's morphic fields could be regarded as possible explicatory approaches. What can be observed, however, is that the results of the present investigation are compatible with Sheldrake's theories, inasmuch as they do not contradict the hypothesis that objects that were or have been familiar to a large number of people are stored in some sort of collective memory.

The results of the present experiment are also compatible with Sheldrake's theory that claims that the effect of morphic fields does not diminish with spatial and/or temporal distance, insofar as the statistical influences of the researched control variables and independent variables were not significant. The fact that the variable "distance" did not have any effect on the choice of authentic symbols or control symbols also supports the claim that we measured some sort of collective knowledge, and furthermore that the results cannot be attributed to an individual mere exposure effect. It could be argued that the original stimuli were only selected more frequently because the participant had perceived them at some point prior to the experiment without being able to consciously recall it, and thus this was not detected by the control question. If that was the case, the test should have shown significant differences between those respondents who originated in the region where the stimuli have been very popular and participants from other regions, who were thus less likely to have perceived the symbols before. This was not the case. The effect measured in the experiment can moreover be referred to as independent of the locality.

We would also like to mention an alternative explanation for the effect we measured: the "evolutionary" approach says that symbols that spread successfully and that prevail over a substantial period of time undergo a modification process that makes them easier to memorize, even for people who have never seen them before. However, as already mentioned above, this explanation does not hold for our experiment involving Russian words, for as writing was invented much later than language, we have no reason to assume that there was a "natural selection" of words regarding more or less appealing word faces.

Which one of these approaches proves true, or if one of them proves true at all, is, however, of no importance for the implications of our experimental results for designing and handling brand logos, packaging, product parts, and design elements. What is important is that we found quasi priming effects for stimuli that were or have been popular for a long time, and these priming effects cause these stimuli to be perceived as more familiar and appealing. Furthermore, it is very intriguing that this effect seemingly neither is culture-specific nor depends upon variables such as age, gender, education or the degree of extraversion of a person. An improved favorability of 20% for original symbols versus comparable control symbols can be regarded as a solid competitive advantage.

The approach at hand used only entirely original symbols. It would be furthermore interesting to investigate whether recognition is particularly determined by specific graphic elements of a given symbol; which elements actually determine this recognition process would be very valuable information for trademark policy. Such elements could then be used as raw material for graphic design.

However, it has to be stated that these facilitatory effects are only tendencies or increased probabilities, but not mono-causal explanations. An absolutely ugly image for instance will not be perceived as being beautiful only because of collective knowledge, priming or mere exposure. The positive effects are minute, but they increase the probability of faster recognition or of a more positive perception regarding a certain stimulus. In the case of brand logos, packaging elements or other aspects that serve for recognizing products, a small positive difference can be crucial for them to be seen and purchased. This is especially important when objects are densely arranged on shelves, competing for the favor of the consumers. Often their optical appearance differs only marginally. Similar is the situation in the case of brand logos, store signs, product offers, and all types of signs that can be found in pedestrian zones, shopping malls, waiting areas in airports, train stations, etc. where the glances of costumers are or have to be attracted by one of the countless stimuli.

In our opinion, the presented results justify that one investigates with an increased interest the implications of yet unproved theories; but even though these theories may be hard to fathom, should we not rather hold on to the principle of falsifiability than to a dogmatic mechanistic world view?

References

1. Aaker, D.A. (1996). *Building Strong Brands*. New York: Free Press.
2. Anderson, J.R. (2000). *Cognitive Psychology and its Implications*. (5th ed.) New York: Worth.
3. Baltagi, B.H. (2001). *Econometric Analysis of Panel Data*. New York: Wiley.

4. Bonnano, G.A. & Stillings, N.A. (1986). Preference, familiarity, and recognition after repeated brief exposures to random geometric shapes. *American Journal of Psychology*, 99 (Fall), 403-415.
5. Bornstein, R.F. (1989). Exposure and Affect: Overview and Meta-Analysis of Research, 1968-1987. *Psychological Bulletin*, 106 (September), 265-289.
6. Hsiao, C. (1999). *Analysis of Panel Data*. Cambridge: Econometric Society Monographs.
7. Dienes, Z. & Berry, D. (1997). Implicit learning: Below the subjective threshold. *Psychonomic Bulletin & Review*, 4 (March), 3-23.
8. Eysenck, H.J. (1964). *The Eysenck Personality Inventory*. London: University of London Press.
9. Henderson, P.W. & Cote, J.A. (1998). Guidelines for Selecting and Modifying Logos. *Journal of Marketing*, 62 (April), 14-30.
10. Janiszewski, C. (1990). The Influence of Print Advertisement Organization on Affect toward a Brand Name. *Journal of Consumer Research*, 17 (June), 53-65.
11. Janiszewski, C. & Meyvis, T. (2001). Effects of Brand Logo Complexity, Repetition, and Spacing on Processing Fluency and Judgment. *Journal of Consumer Research*, 28 (June), 18-32.
12. Jung, C.G. (1981). *The Archetypes and The Collective Unconscious*. Collected Works of C.G. Jung Vol. 9 Part 1, Princeton: Princeton University Press.
13. Keller, K.L. (2002). *Branding and Brand Equity*. Cambridge, MA: Marketing Science Institute.
14. Kuehn, A.A. (1962). Consumer brand choice as a learning process. *Journal of Advertising Research*, 2 (December), 10-17.
15. Kunst-Wilson, W.R. & Zajonc, R.B. (1980). Affective discrimination of stimuli that cannot be recognized. *Science*, 207 (1 February), 557-558.
16. Lefrancois, G.R. (1972). *Psychological Theories and Human Learning: Kongors Report*. Belmont, CA: Wadsworth.
17. Mandler, G. (1980). Recognizing: The judgement of previous occurrence. *Psychological Review*, 87 (May), 252-271.
18. Mandler, G., Nakamura, Y., & van Zandt, B.J.S. (1987). Nonspecific effects of exposure on stimuli that cannot be recognized. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 13 (October), 646-648.
19. Mazur, J.E. (2002). *Learning and Behavior*. (5th ed.) London: Pearson Education Limited.
20. Moreland, R.L. & Zajonc, R.B. (1977). Is stimulus recognition a necessary condition for the occurrence of exposure effects? *Journal of Personality and Social Psychology*, 35 (April), 191-199.
21. Perrig, W., Wippich, W., & Perrig-Chiello, P. (1993). *Unbewusste Informationsverarbeitung*. Bern: Huber.
22. Polanyi, M. (1966). *The Tacit Dimension*. New York: Doubleday.
23. Ratchford, B.T. (2001). The Economics of Consumer Knowledge. *Journal of Consumer Research*, 27 (March), 397-411.
24. Reber, A.S. (1996). *Implicit Learning and Tacit Knowledge: An Essay on the Cognitive Unconscious*. New York: Oxford University Press.
25. Richardson-Klavehn, A. & Bjork, R.A. (1988). Measures of Memory. *Annual Review of Psychology*, 39, 475-543.
26. Seamon et al., (1983). Affective discrimination of stimuli that are not recognized: Effects of shadowing, masking, and cerebral laterality. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 9 (July), 544-555.
27. Seamon, J.G., Marsh, R.L., & Brody, N. (1984). Critical importance of exposure duration for affective discrimination of stimuli that are not recognized. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 10 (July), 465-469.
28. Sheldrake, R. (1981). *A new Science of Life*. London: Blond & Briggs.
29. Sheldrake, R. (1988). *The Presence of the Past*. London: Collins.
30. Williams, K.C. (1990). *Behavioural Aspects of Marketing*. Oxford: Heinemann.
31. Zajonc, R.B. (1968). Attitudinal effects of mere exposure. *Journal of Personality and Social Psychology, Monograph Supplement*, 9 (2, part 2), 1-27.

Appendix

Table A-1

All symbols of the experiment

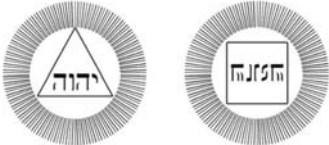
Symbol name	Symbols
Om	
Shiva Lingam	
Islam	
Key of Petrus	
Trinity	
PX	
Alexander the Great	
Roma	

Table A-1 (continuous)

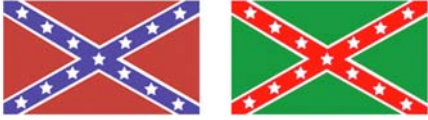
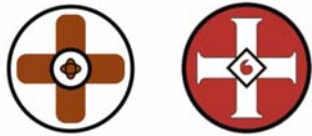
Jerusalem cross	
China flag	
Bangladesh flag	
Confederate Battle Flag	
SS-bolts	
Fasces	
Ku Klux Klan	
Sun Yat-sen	
Anatolij Chubais	

Table A-1 (continuous)

Coca-Cola	
Doordarshan	
Tata	

Table A-2

All Russian words in Cyrillic characters

Translation	Presented at the left side	Presented at the right side
Aggressive	еггссринова	агрессивно
War	война	ойвна
Sun	оцнлсе	солнце
To live	жить	тъжи
Pain	бльо	боль
To kill	убивать	автъбиу
Peace	мир	рми
Beautiful	исакров	красиво
Death	смерть	етьрмс
Water	адво	вода
Good	хорошо	оошхро
Mourning	ачплъе	печаль
To die	аримуть	умирать
Love	любовь	бювьло
Help	ощъпмо	помощь
Bad	плохо	опохл
Life	жизнь	зньжи
To give	адтъва	давать
Hate	ванинестъ	ненависть
Disease	болезнь	езбньло

Table A-3

The logistic regression analysis for each pair of symbols

Symbol	Hit ratio	
	Model	By chance (prior probability)
Islam	66.33%	66.67%
Fasces*	68.80%	63.84%
Coca-Cola	71.61%	71.45%
Om	74.62%	74.62%
SS-bolts	77.13%	77.35%
Key of Petrus	60.98%	61.99%
Tata*	61.84%	53.26%
Roma*	58.21%	50.40%
Shiva Lingam	63.99%	61.58%
Confederate Battle Flag	67.31%	66.73%
Sun Yat-sen	71.96%	71.96%
PX	74.70%	74.49%
Bangladesh flag	59.39%	60.03%
Doordarshan	89.10%	89.10%
Jerusalem cross	64.84%	64.03%
Chubais	61.09%	56.59%
Ku Klux Klan	65.97%	66.13%
China flag	60.19%	58.88%
Alexander the Great	62.00%	57.49%
Trinity	68.64%	68.48%

* indicates a significant model ($p < .05$)

Table A-4

The logistic regression analysis for each pair of Russian words

Word	Hit ratio	
	Model	By chance (prior probability)
Aggressive	56.32%	57.14%
War	67.97%	67.97%
Sun	58.48%	58.64%
To live	56.45%	52.74%
Pain	55.41%	50.89%
To kill	68.07%	68.23%
Peace	68.46%	68.29%
Beautiful	66.17%	65.84%
Death	54.15%	51.87%
Water	63.62%	63.29%
Good	74.05%	74.05%
Mourning	56.82%	55.38%
To die	57.75%	54.16%
Love*	56.91%	51.97%
Help*	56.01%	50.49%
Bad	53.85%	54.17%
Life	58.66%	57.84%
To give	57.33%	55.54%
Hate	60.26%	58.32%
Disease	68.77%	66.83%

* indicates a significant model ($p < .05$)