

Using stories to increase sales at Pfizer

Description

[Tweet](#)

by [Nigel Edwards](#), Strategic Communications Management Vol. 15, Issue 2, Feb-March 2011. pages 30-33. [Available from Cognitive Edge website](#), and found via a tweet by [David Snowden](#)

[RD comment] This article is about the collation, analysis and use of a large volume of qualitative data, and as such has relevance to aid organisations as well as companies. It talks about the integrated use of two sets of methods: anecdote circles as used by a consultancy [Narrate](#), and [SenseMaker software](#) as used by [CognitiveEdge](#). While there is no mention of other story based methods, such as [Most Significant Change\(MSC\)](#), there are some connections. There are also connections with issues I have raised in the [PAQI page](#) on this website, which is all about the visualisation of qualitative data. I will explain.

The core of the Pfizer process was the collection of stories from a salesforce in 11 cities in six countries, within a two week period. With a further two weeks to analyse and report back the results. Before then, the organisers identified a number of “signifiers” which could be applied to the stories. I would describe these as tags or categories that could be applied to the stories, between one and four words long, to signal what they were all about. These signifiers were developed as sets of choices offered in the form of polarities and triads. For example, one triad was “achieving the best vs respecting vs people, making a difference”. A polarity was “worried vs excited”. In previous work by Cognitive Edge and LearningbyDesign in Kenya the choice of which signifiers to apply to a story was in the hands of the story-teller, hence Cognitive Edge’s use of the phrase self-signifiers. What appeared to be new in the Pfizer application was that as each story was told by a member of an anecdote circle it was not only self-signified by the story teller, but also by the other members of the same group. So, for the 200 stories collected from 94 sales representatives they had 1,700 perspectives on those stories (so presumably about 8.5 people per group gave their choice of signifiers to each of the stories from that group).

I should back track at this stage. Self-signifiers are useful for two reasons. Firstly, because they are a way by which the respondent can provide extra information, in effect, meta-data, about what they have said in the story. Secondly, when stories can be given signifiers by multiple respondents from a commonly available set this allows clusters of stories to be self-created (i.e. being those which share the same sets of signifiers) and potentially identified. This is in contrast to external researchers reading the stories themselves, and doing their own tagging and sorting, using NVIVO or other means. The risk with this second approach is that the researcher prematurely imposes their own views on the data, before the data can “speak for themselves”. The self-signifying approach is a more participatory and bottom up process, notwithstanding the fact that the set of signifiers being used may have been identified by the researchers in the first instance. PS: The more self signifiers there are to choose from, the more possible it will be that the participants can find a specific combination of signifiers which best fits their view of their story. From my reading there were at least 18 signifiers available to be used, possibly more.

The connection to [MSC](#): MSC is about the participatory collection, discussion and selection of stories of significant change. Not only are people asked to describe what they think has been the most significant change, but they are also asked to explain why they think so. And when groups of MSC stories are pooled and discussed, with a view to participants selecting the most significant change from amongst all these, the participants are asked to explain and separately document why they selected the selected story. This is a process of self-signification. In some applications of MSC participants are also asked to place the stories they have discussed into one or another categories (called domains), which have in most cases been pre-identified by the organisers. This is another form of self-signifying. These two methods have advantages and disadvantages compared to the Pfizer approach. One limitation I have noticed with the explanations of story choices is that while such discussions around reasons for choosing one story versus another can be very animated and in-depth, the subsequent documentation of the reasons is often very skimpy. Using a signifier tag or category description would be easier and might deliver more usable meta-data – even if participants themselves did not generate those signifiers. My concern, not substantiated, is that the task of assigning the signifiers might derail or diminish the discussion around story selection, which is so central to the MSC process.

Back to Pfizer. After the stories are collected along with their signifiers, the next step described in the Edwards paper is – looking at the overall patterns that emerged. The text then goes on to describe the various findings and conclusions that were drawn, and how they were acted upon. This sequence reminds me of the [cartoon](#), which has a long complex mathematical formula on a blackboard, with a bit of text in the middle of it all which says – then a miracle happens. Remember, there were 200 stories with multiple signifiers applied to each story, by about 8 participants. That is 1700 different perspectives. That is a lot of data to look through and make sense of. Within this set I would expect to find many and varied clusters of stories that shared common sets of two or more signifiers. There are two ways of searching for these clusters. One is by intentional search, .i.e. by searching for stories that were given both signifier x and signifier y, because they were of specific interest to Pfizer. This requires some prior theory, hypotheses or hunch to guide it, otherwise it would be random search. A random search could take a very long time to find major clusters of stories, because the possibility space is absolutely huge. It doubles with every additional signifier (2,4,8,16!) and there multiple combinations of these signifiers because 8 participants are applying the signifiers (256 combinations of any combination of signifiers) to any one story. Intentional search is fine, but we will only find what we are looking for.

The other approach is to use tools which automatically visualise the clusters of stories that exist. One of the tools CognitiveEdge use for this purpose (and it is also used during data collection) are triangles that feature three different signifiers in each corner (the triads above). Each story will appear as a point within the triangle, representing the particular combinations of three attributes the story teller felt applied to the story. When multiple stories are plotted within the triangle multiple clusters of stories commonly appear, and they can then be investigated. The limitation of this tool is that it only visualises clusters of three signifiers at a time, when in practice 18 or more were used in the Pfizer case. It is still going to be slow way to search the space of all possible clusters of stories.

There is another approach, which I have discussed with David Snowden. This involves viewing stories as being connected to each other in a network, by virtue of sharing two or more signifiers. Data consisting of a list of stories with associated signifiers can be relatively easily imported from Excel into Social Network Analysis software, such as Ucinet/NetDraw, and then visualised as a network. Links can be size coded to show the relative number of signifiers any two connected stories share. More

importantly, a filter can then be applied to automatically show only those stories connected by x or more shared signifiers. This is a much less labor intensive way of searching huge possibility spaces. My assumption is that clusters of stories sharing many signifiers are likely to be more meaningful than those sharing less, because they are less likely to occur simply by random chance. And perhapsâ?| that smaller clusters sharing many signifiers may be more meaningful than larger clusters sharing many signifiers (where the signifier might be fuzzier and less specific in meaning). These assumptions could be tested.

To recapitulate: Being able to efficiently explore large possibility spaces is important because they arise from giving participants more rather than less choice of signifiers. Giving more choice means we are more likely to hear the participantsâ?? particular views, even though they are voiced through our constructs (the signifiers). And larger number of signifiers means that any clusters of highly connected stories is more likely to be meaningful rather than random.

Social Network Analysis software has an additional relevance for the analysis of Pfizer data set. Within the 1700 different perspectives on the stories there will not only be a network of stories connected by shared signifiers. There will also be a network of participants, connected by their shared similar uses of those signifiers. There will be clusters of participants as well as clusters of stories. This social dimension opened up by the participatory process used to apply the signifiers was not touched upon by the Dawson paper, probably because of limitations of time and space. But it could be great significance for Pfizer when working out how to best respond to the issues raised by the stories. Stories have owners, and different groups of owners will have different interests.

Category

1. Uncategorized

Tags

1. Most Significant Change (MSC)
2. narrative
3. qualitative
4. SenseMaker

Date

13/10/2025

Date Created

18/02/2011

Author

admin