

A Fisheye Follow-up: Further Reflections on Focus + Context

George W. Furnas

School of Information, University of Michigan

Ann Arbor, Michigan, USA

furnas@umich.edu

ABSTRACT

Information worlds continue to grow, posing daunting challenges for interfaces. This paper tries to increase our understanding of approaches to the problem, building on the Generalized Fisheye View framework. Three issues are discussed. First a number of existing techniques are unified by the commonality of *what* they show, certain fisheye-related subsets, with the techniques differing only in *how* they show those subsets. Then the elevated importance of these subsets, and their generality, is used to discuss the possibility of non-visual fisheye-views, to attack problems not so amenable to visualization. Finally, several models are given for why these subsets might be important in user interactions, with the goal of better informing design rationales.

Author Keywords

Information Visualization, Focus+Context, Fisheye Views, DOI, Zoom, Bifocal lens, ZUI, View+Overview

ACM Classification Keywords

H. Information Systems H.5 Information Interfaces and presentation (I.7) H.5.2 User Interfaces (D.2.2, H.1.2, I.3.6)
Subjects: Theory and methods

INTRODUCTION

Twenty years ago, CHI86 published a paper [11] addressing a growing problem -- that information worlds were getting large, while our windows into those worlds were quite small. The paper proposed a general approach to making useful small views of large information worlds, called *Generalized Fisheye Views*. These views provide detail at the current focus of the viewer's attention, but show only increasingly important features further and further away. This formulation was inspired by Fisheye Lens views in photography which date back to 1906, when Wood [28]

investigated how the aerial world above the water's surface would appear to a fish submerged below. Because of refractive effects governed by Snell's law, light rays more directly overhead, being almost normal to the surface, are largely unbent, while those coming in at a shallow angle are substantially bent. As a result the center of the view is largely undistorted, shown full size, while the edges of the view are increasingly compressed. The metaphor captures something very intuitive, that people need a way to pay attention to particular details they are focused on, yet also need some surrounding context.

Since CHI86¹ many dozens of specific computer-based techniques have been proposed to provide a balance of "focus and context", particularly in the field of Information Visualization. While the accumulated body of work is held together by a diffuse idea that a balance of focus+context is needed, there has been relatively little unified understanding. We do not have good answers to questions like: Just what do we mean by focus and context? What information do we really need to show in these views? Why do users really need that information? This lack of understanding is a problem because in the last 20 years the situation has only gotten worse, with multi-terabyte databases, multi-mega-line software systems, and the multi-giga-page web. It is arguably time for a deeper look.

Three principal issues will be examined in this paper. The first will focus on the distinction between content and presentation, that is, *what* these techniques are trying to show as opposed to *how* they try to show it. Several classes of visual techniques will be analyzed in terms of what information they actually show. That they seem to reveal a specific common set of information leads to an instructive decomposition of the space of techniques in terms of variations of *how* they show that common information, and the corresponding trade-offs.

The *what vs. how* discussion will concentrate on examples drawn from the information visualization literature, because that is where most of the work regarding focus+context has been done. However, the conceptual separation clarified there sets the stage for remembering that visual presentations are only one option for *how* we convey the

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee.

CHI 2006, April 22–27, 2006, Montréal, Québec, Canada.
Copyright 2006 ACM 1-59593-178-3/06/0004...\$5.00.

¹ ... and a few before, notably [2] and [9]

information, and that, moreover, the problem of making interfaces to large worlds has important forms where the concrete optical metaphor, and visual approaches in general, may not be appropriate. Going “*Beyond Visualization*” will therefore be the second issue discussed in this paper.

The third issue to be addressed is “*Why?*”, that is, pursuing a deeper understanding of why users might actually need a balance of focus and context information. Possible answers provide critical insight for task analyses in the upstream stages of design: If we know what user-needs such focus+context balances satisfy, then we will know better what user-needs to be on the look-out for and how to shape our designs accordingly. Towards this end, the *Why* section will present a series of possible theories, using existing literature from several fields.

In addressing these three issues, this paper will draw heavily on the original generalized fisheye formalism [10] [11], for reasons that require some familiarity with the basics of the formalism. Furnas suggested that useful small views can be generated by simply presenting the most “interesting” subset that limited resources will allow of the large world. That suggestion turned the problem of generating a small view into one of estimating a user’s *Degree of Interest (DOI)* in various features of the world, given their current activity. It was then proposed that the DOI take into account both the *A Priori Importance (API)* of features in the world, and their *Distance (D)* from the user’s current focus. In its most general form [10], the Generalized Fisheye Degree of Interest function at some point, x , given the current focal point, “.”, was defined to be

$$DOI_{FE}(x | .) = F(API(x), D(., x)), \quad \text{Eqn.1}$$

where F is some combining function that is monotone increasing in the first argument, and decreasing in the second. That is, the degree of interest in x increases with its global importance and decreases with its distance from the current focus. Using the general DOI strategy, an interface would present all points, x , that at least meet some minimal criterion, c , of interestingness, i.e., the set of all points, x , such that

$$DOI_{FE}(x | .) > c. \quad \text{Eqn.2}$$

We will call such a set of points a *Fisheye-DOI subset* (FE-DOI subset). The threshold, c , would be chosen to be restrictively high if few resources were available, and generously low if resources were plentiful. The result is a presentation that shows even minor details near the point of focus, and only increasingly more important things that are further away.

This formulation differs in three important ways from its inspiring namesake, a fact often overlooked. First, while the optical Fisheye Lens is all about distortion, the Furnas version was about selection -- what is to be included (the more interesting stuff) in the view and what is to be left out (the less interesting stuff). This selection/distortion contrast

will be helpful in discussing *what* is shown vs. *how* it is shown. Second, while optical Fisheye Lenses work in the familiar world of low dimensional Euclidean space, the DOI version was deliberately agnostic about geometry, to allow it to generalize to other kinds of worlds. This abstract-general vs. concrete-optical contrast will be particularly useful for moving *Beyond Visualization*. Finally, the DOI formalism redefined the problem using terms like “interest”, “importance” and “distance”. These are concepts with specific relevance to users and tasks, thereby providing a useful orientation for the final, “*Why?*”, discussion.

WHAT VS. HOW

The theme of providing “Focus + Context” (F+C) views, has generated a large number of techniques, particularly in the information visualization literature. One group, more directly inspired by the geometric distortion of the Fisheye Lens metaphor, might be called “Distortion” views. Some of these have followed up explicitly on the Generalized Fisheye Degree of Interest formalism, like [23] who used a FE-DOI, not explicitly to select what should be shown, but to determine how much space should be given to what is shown, with more interesting things shown larger. Other distortion techniques have used various geometric approaches (perspective wall [19], document lens [22], hyperbolic tree browser [16]) to provide the F+C balance. Another set of approaches used explicit distortion, differential magnification [2][9][15][17] and stretching functions [24] to allocate space preferentially to the focus.

A second group has used non-distorting magnification techniques to make focus and context accessible. Some techniques are quite old, using separate display regions for different magnifications, e.g., maps that have a separate insert for either a close-up or an overview. Other methods of this group are quite new, like those using dynamic interactive zooming [5].

A third, very sparse group has simply used differential resolution for focus vs. context regions of a display [3][4], examples of a selection technique, without any size changes or distortion.

The list of good work given here for all three groups is by no means complete; see [6][7][8][17][26] for further review and discussion. Although covering a similar suite of cases to some of these other reviews, the goal here is somewhat different than, say, unifying them under a generalizing magnification function. Here the unifying emphasis is a claim that what is most important about all these approaches is *what* they convey. Users have tasks to do and need certain information to do them. *How* we provide that information is a secondary consideration compared to *what* we provide.

Fisheye DOI Subsets vs. Fisheye Lens Distortion

To get a better intuitive understanding of the difference between what is presented vs. how it is presented, consider

the famous cartoon example of a Fisheye Lens distortion presentation in popular culture, Joel Steinberg's much imitated *New Yorker* magazine cover, caricaturing a New York citizen's view of the world. It shows local details of parking garages and mail boxes on 9th Avenue while Chicago and LA are just shown as dots in the distance, with the Rocky Mountains as a few bumps somewhere in between. The Pacific Ocean is a stripe further off in the distance, and then China, Japan and Russia are featureless blobs away at the horizon. *What* this cartoon shows is, for example, details of mailboxes and parking from the local neighborhood. It has filtered out those details from regions further away, showing only much higher-level information from regions further away, like that there is a city called Los Angeles on the West Coast. The small subset of information actually included, i.e., *what* is shown, is actually fairly sensible for the New York resident – there is no need to know the location of mailboxes in LA. *How* this set of information was shown, i.e., the humorous distortion of sizes and distances, was a design choice of Steinberg as cartoonist.

A simple example will help further with this *what/how* distinction, and also allow us to bring the generalized fisheye formalism to bear. Consider trying to make a small view of a long, ordered list. Here, for example, we take the schematic case of a list of the letters of the alphabet. Figure 1 shows the successive construction of a FE-DOI for this list (a). In the A Priori Importance graph of (b), the first and last few items are given special importance, reflecting well-known primacy and recency effects. The graph in (c) simply shows the Distance from the focus. Combining these two components in an additive way creates a simple FE-DOI (d). The filtered list can be shown with items in their original locations (un-distorted, e) or can be compressed (eliding empty space, f). In the former case it is easy to see true distances between list items; the latter takes up less display space. These represent design trade-offs whose value depends on the demands of a task, but in both views, all but the “most interesting” information has been deleted. We will discuss such tradeoffs later. For now it is sufficient to understand that one can talk about *what* is shown – the information indicated to be of interest by the generalized fisheye DOI – separately from *how* it is shown.

The Fisheye Subset in Information Visualization

The underlying FE-DOI subset is central to a broad variety of F+C information visualization techniques. At a metaphorical level that is to be expected because the FE-DOIs are intended to capture just what the phrase “Focus + Context” is trying to capture. However the connection is deeper than that; at a formal level it provides a tighter unification of many of the methods. It is useful to begin exploring that unification by analyzing what some people think of as the canonical Fisheye views – those that, like Wood's photographic lens, use distortion to magnify the center and compress the surroundings.

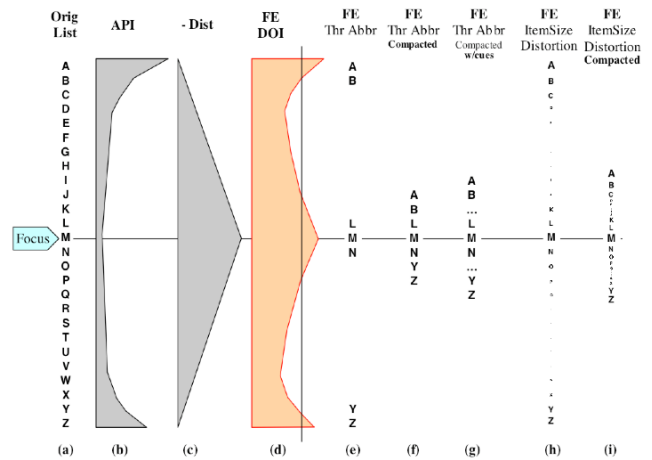


Figure 1. Selection vs. Distortion in a Fisheye View of a List. (a) Shows an ordered list. (b) Shows a hypothetical A Priori Importance function over the list, in this case reflecting that beginning and final items are more significant than random inner items. (c) shows a distance function from the focus at M. (d) shows the simple additive combination of the API and Dist components to get a Fisheye Degree of Interest function, from which nested FE-DOI subsets can be extracted by application of a threshold. (e) shows one such set, arising from the indicated threshold, resulting in information reduction. In (f) this FE subset is shown with a distorted geometry which decreases the space used (space reduction) but loses global distance information. (g) shows the use of explicit elision markers “...” to re-introduce some of the lost distance information (h) shows the items at a size related to the DOI value, and (i) shows the resized items moved together to achieve a more uniform density.

Distortion views and the FE-DOI

The list example in Figure 1(a)-(f) uses the logic of filtering first, then distorting the placements so as to conserve space. This distinction between filtering or selecting information, and distorting the information deserves further examination. There has been much research over the past two decades on various “fisheye” and other F+C distortion presentation techniques that never make explicit any notion of filtering. However, such *distortions* do indeed *filter* information, because any real transmission medium (the display, the retina) has finite resolution. As a result distortion, with its associated differential magnification, implies a filtering of information in the spatial frequency domain. Most obviously, as the rendered size of features of a world decrease below the pixel size in the display, they are filtered out. In fact, the ability to resolve those features gets worse way before that. Thus, while distortion based F+C techniques certainly change the displayed position of various items in the world, the associated magnifications and demagnifications are also really altering what information is available about those items. One could take the magnification function of any distortion transformation and use it as a variable blurring function, a space-varying spatial frequency filter, and run it over an undistorted version to get a clearer understanding of the filtering going

on. Places given high magnification in the distorted view get a higher spatial frequency cutoff (essentially a smaller blur radius). The distortion-only and blurred-only renderings contain the same subset of localized information about items throughout the space.

The information that they make available is exactly a FE-DOI subset. The distortion fisheye views decrease magnification as a function of distance from the current focus. The relationship of this decreasing magnification to a DOI can be understood formally by referring back to the definition of the FE-DOI. After introducing the most general form of the FE-DOI, Furnas [10][11] went on to explore a simple *additive* class of examples, letting $DOI(x, \cdot) = API(x) - D(\cdot, x)$. Distortion fisheyes can be seen as a very similar, multiplicative class. Their distortion rule renders an object or visual feature, x , according to a rule like: $RenderedSize(x) = TrueSize(x) * Mag(D(\cdot, x))$, where $Mag(D)$ is a monotone decreasing function of distance. Interpreted as a FE-DOI formula, $RenderedSize$ becomes DOI and $TrueSize$ becomes API . That $RenderedSize$ corresponds to the *Degree of Interest* would not be surprising to Sarkar and Brown [23], who set up their distortion FE explicitly in this way (“render with size proportional to interest”), but it applies to the broader class generally. More instructive is that $TrueSize$ corresponds to *A Priori Importance* – that is, the distortion techniques can be interpreted as acting as though big things are more important, *a priori*. If the assumption, that low spatial frequency information is more important than high spatial frequency information, is correct, then the distortion techniques are exactly Fisheye DOI *filters*, operating on spatial frequency. Insofar as the assumption is *wrong*, then they will not be reasonable FE-DOI filters, and probably correspondingly less useful. For example, the Document Lens [22] used reduced magnification to show pages of a document surrounding the current focal page. One of the drawbacks was that the large-scale features visible after such reduction were things like paragraph breaks and associated indents, features *not* more important, *a priori*, than a few small key phrases in the paragraph. In this way the technique was not showing a meaningful FE-DOI subset, and was probably less useful accordingly.

All magnification-based presentation techniques, including the zoom techniques of the next section, essentially constitute *design as if* size is what really matters: as if larger things are a priori more important and as if things of any a priori size should be shown larger if you are interested in them, and neither assumption is always true.² One virtue of the FE-DOI formulation is that it casts these visual techniques back in terms of user/task parameters, like importance and interest, a deconstruction that reminds us of such mismatches.

² Semantic zooming, in which the appearance of an item changes non-geometrically with size so as to stay meaningful, was devised (e.g., [5]) to solve exactly this suite of problems.

Zoom views

The FE-DOI also captures what goes on in the variable zoom techniques, often called Zoomable User Interfaces (ZUI) [5]. In a ZUI, there is typically a viewing window of fixed width, v , and fixed resolution, r (e.g., in pixels per cm). This is used to show a world region of variable width, W . In zooming, the magnification, $m=v/W$, is varied to look at larger or smaller world regions in the same fixed-size viewing window.

As the magnification changes the view’s extent, it makes a corresponding change in the world-size of the smallest details that are visible. If you zoom out to see more breadth, you lose details. Formally, the world-width, W , seen in a given view is rendered by rv pixels, and so nothing smaller than size $w = W/rv = 1/rm$ can be resolved. Therefore, in zooming, these two will always move together: if w is the world-size of the smallest feature visible when viewing a world-region of width W , then αw will be the world-size of the smallest feature when viewing a region of world-width αW .

Consider the following question: When can a feature i of size w_i and location x_i , be seen in a view of world-size W , centered at a point, “.”? The location x_i will only be included when a view centered at c gets large enough: $W \geq 2D(x_i, \cdot)$. But then, to be visible, the feature must be of a size:

$$w_i \geq W/rv = 2D(x_i, \cdot)/rv. \quad \text{Eqn.3}$$

That is, if you have a fixed center-point, the further away something is, the more zoomed out you will have to be to see it, and hence the larger it must be to be seen. Put another way, consider the information available from a series of concentric zooms (i.e., views related by changing the magnification, m , only, without panning the center, “.”, of the views). In aggregate, the whole set of concentric zooms will show features whose size must increase with distance from the center.

If this sounds like a subset dictated by a fisheye DOI, that is exactly what it is. Recall the use of a threshold to generate the DOI subsets (combining Eqns. 1&2), and what it would look like in the multiplicative case:

$$\text{General: } f(API(x), D(x, \cdot)) \geq c$$

$$\text{Multiplicative: } API(x) / D(x, \cdot) \geq c$$

Now, returning to the analysis of zoom, if we make the assumption that the size of a feature is its A Priori Importance (the implicit visual assumption in any magnification based technique), then rearranging Eqn.3 yields the formula for what can be seen in a set of con-focal zooms:

$$\begin{aligned} w_i / D(x_i, \cdot) &\geq 2 / (rv) \\ API(x_i) / D(x_i, \cdot) &\geq c \end{aligned}$$

That is, the union of concentric zooms yields exactly a FE-DOI subset. Thus the family of general fisheye DOI

functions describes what features are visible both in the fisheye distortion presentations and concentric-zooms presentation of the information. This equivalence of the techniques extends even to a detail of design control. Recall that the DOI threshold, c , is often accessible to the interface designer for adjusting the view to the existing display resources. As the formulation above indicates, the zoom designer, like the fisheye distortion designer in general, can naturally control how deep to go in the fisheye DOI subsets: showing smaller things further away from the center of focus by either increasing the pixel resolution, r , or the size of the viewing window, v .

View+Overview and View+Closeup

Two closely related magnification-based F+C techniques, “view+overview” and “view+closeup”, emerged long ago in the world of paper displays. Both show just two discrete levels of focus and context. The “view+overview” variant has a large detail view and with a small overview, inserted in a box typically off in one corner. The “view+closeup” variant, in contrast, has a large and extensive view with the small insert showing a close-up of some special region, e.g., a map of France with a smaller close-up of Paris. As in a ZUI, both variants show different “zoom” views, but simultaneously, instead of presenting them over time.

By presenting only two levels, focus and context, these differ from the richer range of trading off one against the other represented in the canonical FE-DOI. This difference must ultimately prove problematic for truly large worlds where there is important structure at many scales. There the user will need more than one layer of context.

This point can be quantified fairly simply. Let B be the scale bandwidth of the presentation technology, defined as $B = \text{Extent} / \text{Grain}$ (e.g., pixel-width of the display = size of display / size of pixel). Let R be the scale range of your information world, defined as $R = \text{WorldSize} / \text{SmallestDetailSize}$. Clearly, if the scale range of the world does not exceed that of the presentation technology, i.e., $R \leq B$, there is no problem. However, if $R > B$, special techniques are needed to show any point of interest in its full context. Using zoom, for example, we would need $a \cdot \log R / \log B$ views to show the focus and context of any specific detail point.³ For example, using a 1K display to view a point in a 1Meg world, would require at the absolute minimum 2 views: one 1K view showing a close-up of the details, a second 1K view showing an overview with a highlighted pixel showing where the detail view fits into the overview. A 1Gig world would require at least 3 views. The point is that while a single focal resolution and a single “context” overview may be enough for moderate sized worlds, large worlds would require several layers of context. This realization was the motivation behind the

many layers of simultaneous zoom used by Lieberman in his *Powers of Ten Thousand* technique [18] for showing very large worlds.

The lesson is that, for truly large worlds, even the View+Overview methods will need to show a fuller FE-DOI subset, requiring multiple levels of overviews.

Multi Resolution Displays

The last group of techniques has no zooming or distortion, but simply uses multiple resolutions, high in a center region of the display, and lower around it (e.g., [3][4]). This approach is explicitly a FE-DOI filter, working in exactly the spatial frequency domain that is only implicit in the distortion and zoom techniques. (The use of only two level means it is subject to some of the same world-size limitations as the two-level magnification approaches).

Same *What*, different *How*: Design Tradeoffs

The claim here has been that *what* these techniques show is much the same – a subset of the original information, as dictated by a FE-DOI. While this unifies the class, the techniques clearly differ. Much insight into the differences between them comes from considering *how* they show that same subset. Ideally, we want a technique that only does the FE-DOI filtering, with no undesired side effects. The only way to do this is with a resolution-based presentation – full sized but with increasing blur away from the focus. On the other hand, often the point is to make a presentation of smaller physical size – and immediately there are choices about what to give up: view-size reduction, shape preservation, topological continuity, simultaneity, etc. In any of the choices of *How* to display the FE-DOI there are sacrifices, though ways have been devised to mitigate the problems. Specifically,

- Distortion techniques present the information simultaneously, with topological continuity, but introduce geometric distortion, changing aspect ratios in regions outside the focus, and altering shapes of large patterns in general. The user must understand that there are distortions in shape and position, and be able to mentally undo them if needed. One approach is to give distortion maps that overlay the correspondingly distorted version of a regular grid.
- Dynamic zooming techniques (ZUIs) do not introduce geometric distortions, but present the information spread out over time. One problem is that it “uses up” the temporal dimension – making it poor for giving a F+C rendering of a dynamic, animated world. Another cost, even for static worlds, is that to apprehend the whole FE-DOI set the user must integrate over time, requiring not just memory for previous views, but a reasonable temporal calibration to keep track of where in the temporal sequence (and hence scale) they are at the moment. This problem can be mitigated, for example, by an auxiliary indicator of scale (e.g., a “scale thermometer”).

³ The constant a goes up proportionally if you want to use more than one pixel to show the smallest “detail” and the position of small views within the larger ones.

- Multiple simultaneous views at different scales, e.g., View+Overview and View+Closeup displays, or overlaid transparent views (Powers of 10,000) present the information simultaneously, without geometric distortion, but with topological discontinuity at the edges of the views (“Where does the road that runs off the edge of this close-up view, appear in the overview?”) Users must find correspondences between features at the edge of one view, and internal to the other. The problem is mitigated by an indicator rectangle diagramming where the close-up region fits in the overview, or by showing the mouse cursor in both views.
- The F+C multi-resolution displays present the information simultaneously, without topological or geometric distortion, but without display size reduction. Here the user needs large display resources – perhaps expensive and less portable. This problem is mitigated by a large equipment budget and a personal porter ☺.

These design tradeoffs, even when mitigated, are serious enough that the techniques are not interchangeable -- their appropriateness depends on the tasks to be done. A truck driver might use a very-wide-angle Fisheye mirror (distortion) to get a small view of cars coming up from behind. This works well because he needs simply to see if there is anything there, and see its continuity of movement towards him. If his job were to identify the exact make and model of the car (a shape sensitive task) this would not be a good interface. Similarly a network engineer typically cares more about topological continuity than geometric accuracy, so FE Distortion views are quite fine. When geometry matters more, perhaps in situations using Geographic Information Systems, zooms and overviews may be more important. The point is that these are at their core all trying to hang on to the *What* – namely the FE-DOI subsets – while scrambling to figure out *How*.

BEYOND VISUALIZATION

The examples in the preceding analyses were all drawn from the field of information visualization, where there has been much work. As a result the examples all have a strong visual, often almost optical character. However, the FE-DOI formulation helped emphasize the invariant concern for the *what* implicit in these techniques, their content, instead of the visual/geometric particulars of the *how*. The independence of the FE-DOI from visual renderings can be taken further. The whole point of the FE-DOI formulation was to define *generalized* fisheye views, where familiar 2 or 3 dimensional Euclidean geometry may not be relevant at all, where the FE-DOI *what* can still be defined, but there may be no direct lens-like analogy to give a *how*. One can generalize the *geometry*, beyond 2D/3D, to list structures (as in Figure 1), trees, graphs, DAGs, multitrees [12], tables, etc. One can generalize the notion of *a priori importance* beyond geometric size, to major vs. minor conceptually important aspects, like Vice President vs.

laborer, or high-level directories vs. low-level files. Finally, one can generalize the presentation resource, using the DOI to allocate perceptual attributes like color or sound (instead of size or resolution).

The generalization can go even further, to purely conceptual domains and presentations. Much of the information overloading us is **not** particularly visual – consider the thousands of news stories, or millions of books, or billions of web pages, only a few of which we might be interested in. The information world is large, and our resources, in time, attention, effort, etc. are small. We need to explore the value of a FE-DOI in these cases as well.

A *generalized* FE strategy is possible whenever three requirements are met: (1) some reasonably static structure with a notion of distance, (2) some notion of independent level of detail or *a priori* importance for different parts of the structure, and (3) interaction can be considered as focused at a point (or small region, or small number of points) in the structure. There are many ways to define the kind of proximity structures over information objects needed for requirement 1. For example, text objects can be placed in high dimensional term spaces using standard techniques in vector-based information retrieval. These can be used to provide notions of distance (in fact these are central to vector IR). Independent notions of *a priori* importance can come from user-community popularity data, for example. A user’s focus can come from a query, using standard IR techniques to map it to a point in the high dimensional term space. Then, using the FE-DOI we can return, not just points close to the query, but a Fisheye DOI Subset – returning also points somewhat further from the query focal point if they are highly endorsed by the community. Indeed, the Google PageRank break-through is essentially a generalized Fisheye – with recursive linked-to weight defining a kind of API, which is weighted against closeness to the query.

Similar implementations should be possible in recommender systems, so that you do not just get things close to your personal favorite, but things further away if they are of compensatingly greater global popularity. (E.g., the system says, effectively: “I know you don’t usually like horror movies, but this one is particularly highly rated in the genre, so you might want to check it out.”) Such a non-visual fisheye strategy would not only give people more things they might really like (even though some are outside their usual preference locality), but also prevent the balkanization of the world into self-selected, narrow interest groups. People would end up getting more overlap with the rest of the world, more context to go with their focus.

There are many other possibilities. Context awareness, both in CSCW and in portable and embedded applications might benefit from representing a FE-DOI subset of their respective worlds. Text generation systems might use a FE strategy to select what to say: A question like “What is a

‘FE-DOI?’” probably deserves a few layers of nested conceptual context in the answer.

WHY THE FISHEYE DOI SUBSET MATTERS

This paper so far has had a lot of emphasis on the FE-DOI. It was shown to be central to quite a few intuitively created and effective visualization designs. Moreover, the various empirical studies done in the mid 1980’s on naturally occurring FE Views (see [11]) indicated that what mattered were the FE subsets characterized by the FE-DOI; the studies were agnostic about variations in the specific internal presentation/representation of these subsets.

Just why are these FE-DOI subsets so important? This is much more than an academic question. Good HCI design requires understanding what real users need when doing their real tasks, and then trying to select and shape technology options accordingly. This in turn requires understanding the function of various design attributes – what are they good for – so that the options can be selected and tuned appropriately. Thus, understanding why FE-DOI subsets might be important for various purposes should help designers know what to consider when tuning their designs to varied aspects of users’ tasks.

Before presenting hypotheses about the import of FE-DOI subsets, there are two caveats. First, the hypotheses are untested. Testing them would be fascinating but non-trivial. They are presented as hypotheses in the hope that, even if the specific ones here are false, they will encourage designers to try to think more deeply about what is going on in their F+C designs, how and why they work for people.

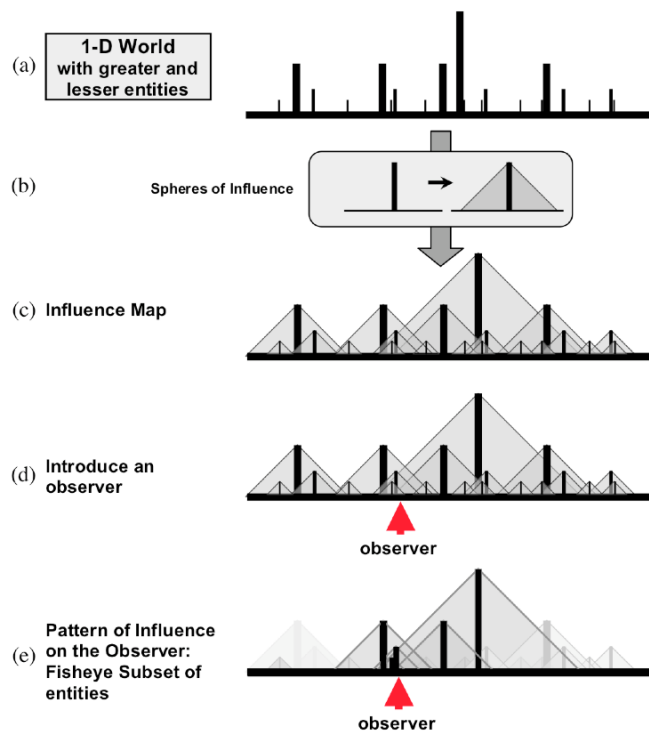


Figure 2. The Spheres of Influence Model. (See text.)

These hypotheses, therefore, are meant to be plausible starting places for such discussions. Second, a single untested conjecture can be seductively dangerous: too easy to accept as true, and blinding as a result. It is useful therefore to have multiple hypotheses, to shake oneself out of simple naïveté in the absence of data. It is in this spirit that several hypotheses are given. They are not necessarily mutually exclusive or logically inconsistent; in various real world situations, several of them may be operating.

The Spheres of Influence Model

The first model for understanding the possible importance of the fisheye subset begins with a world populated by entities of different magnitudes, and considers which of those would have influence on an observer. Figure 2(a) shows a 1-dimensional version of such a world. Let us suppose that some of these entities are "greater", others "lesser", in the sense that they have spheres of influence of greater or lesser size. For example, in a geographic world for navigation, these might be landmarks that can be seen, and hence serve to guide, from greater or lesser distances. Or the entities might be radio stations with transmitters of varying strengths and corresponding reaches. They might be commercial centers with trading radii of more or less size. They might be feudal castles with more or less power. They might be predator threat zones, or food sources of varying magnitudes.

Figure 2(b)&(c) shows an *influence map* for the world, created by drawing a "sphere of influence" of appropriate size around each entity. Here they are drawn as simply falling off linearly with distance, though of course they might have more arbitrary form (e.g., flat up to a fixed radius, hyperbolic, exponential, monotone). These spheres of influence overlap, leading quite naturally to the central question of the model: What is the set of things that influence any given point? To answer that we introduce an "observer" at some arbitrary point (Figure 2(d)), and then ask what spheres of influence fall over it. The result, illustrated Figure 2(e), is exactly a Fisheye Subset of entities.

Thus in any such world, if the notions of influence are meaningful to the activity of our observing agent, then any well designed agent would have to take into account a fisheye subset of its world. The various empirical results reported in [11] about people representing fisheye subsets of their worlds may just reflect an adaptive response to the need to represent those things which matter in a world with entities operating at different scales. If HCI techniques are going to augment human intelligence successfully to help users deal with large worlds containing entities having spheres of influence of varying sizes, they must help the user represent this FE subset.

Figure 2 illustrated a 1-dimensional spatial structure, but clearly the analogy holds for other sorts of structures as well, including non-spatial dimensional structures, like time, and non-dimensional structures like trees (e.g.,

corporate hierarchies, file systems, structured programs) and graphs - as long as there are notions of structure at different scale that have domains of influence of different sizes. All these would require agents to have FE-DOI representations of some sort to deal with those parts of their world that influence them.

There is an important duality here. Figure 2(b) shows the *has-influence-upon* set for a given entity in the world. Figure 2(e) shows the *is-influenced-by* set for a given observer. The former helps answer the question, "What views should this object be present in?", the latter, the question, "What should be in an observer's view from a given point?" The FE subset results, by this duality, from the simple fact that entities have influence at different scales. The dual questions are useful to remember in design. When making a focus+context viewer, think about what it implies about which views any given feature of the domain should be seen in, and ask if it makes sense.

Nested Nearly Decomposable Systems Model

In the Spheres of Influence model, more remote items made it into the observer's view only if they had a larger sphere of influence. Items were otherwise not relevantly distinguished. It is not that some were aggregates of smaller ones, or some were abstractions of sets of others. They were just bigger, in some sense.

A fisheye interest set can arise in a multiscale world of aggregates via an extension of a process described by Herb Simon in *Sciences of the Artificial*, in his discussion of Nearly Decomposable Systems [25]. Simon described a system as a set of elements that interact with one another. A fully decomposable system is one whose elements can be divided into disjoint subsets where there are no interactions between elements in different subsets. In such a system, by definition, the behavior of any particular element is only affected by those in its subset.

Simon [25] analyzed a somewhat more complicated situation, where the system can be divided up into subsets such that couplings within subsets are strong, and couplings between subsets, while not non-existent, are weak. Such *Nearly Decomposable Systems (NDS)* obey a theorem that, (1) the behavior of an individual element is determined in the short run only by those within its own component, and (2) elements in other components **do** have impact, but at a slower time scale and only in aggregate. If one defines, as Simon does, a hierarchical nesting of such NDSs with successively looser couplings at higher levels, the NDS theorem implies that the behavior of any element is influenced by a fisheye subset of the whole. That is, details matter nearby, but only successively more aggregate behavior matters further away.⁴

Meaning from Context

One of the basic purposes of any presentation of an information structure is to help a user extract meaning, to *understand* something about the structure. But the meaning of any token in the structure almost always depends on the context in which it appears. Thus an individual letter has no meaning except by virtue of the letters around it that together form a word. An individual word gets much of its meaning by virtue of the other words in its sentence. A sentence gets meaning from the paragraph, a paragraph from the surrounding section, and so on. Similarly in computer code, the significance of a variable depends on where it sits in the nested contexts around it.

This becomes a Fisheye phenomenon via the Spheres Of Influence and the NDS mechanisms: the surrounding entities at different scales of aggregation exert a semantic influence on any given item of interest. Thus the meaning of a word depends on the words around it in a detailed way, but only in an aggregate way on the words in other paragraphs.

Navigational Support

Context is not only needed to interpret a static view of an item, providing meaning. It is also a critical for moving around effectively.

In [13] it was pointed out that one basic need for moving efficiently through large information worlds is the ability to get from any one point to another with a small number of steps, each chosen from a small set (i.e., presentable in a small view). The judicious use of a relatively small number of long distance links can be of great advantage here, decreasing the traversal diameter of the structure. The result is that the set of choices available at any moment tends to involve some short local links as well as some increasingly long distance ones -- essentially a FE subset. A FE-DOI subset thus tends to give good traversal: Successively larger distances can be traversed efficiently by following a direct link to successively more remote things. Various traversal schemes -- FE Lens movement, ZUI, etc shorten the time to get to remote things by presenting direct access to a FE subset.

Actually the pattern of increasingly long distance links that enables efficient traversal does not itself yield a generalized fisheye view because nothing has been said yet about what information is provided about those links. A link to a remote place is valuable because it provides quick access to a large set of otherwise hard to reach locations. In order for a navigating user to know that such a link exists and is worth following, there must be information associated with that link that efficiently indicate the set the link leads to. If the links are set up to provide efficient traversal, the further away the links lead the larger the set of things they provide access to. The information associated with these links will therefore have to indicate the content of sets that are increasingly larger as they get further away. Thus, to make navigation of the structure effective, the total navigational

⁴ The NDS theory actually results in a FE in both space and time, with lower spatial *and* temporal resolution for remote components.

information at a given node will tend to be a FE-DOI view. (See [14] [21] for some empirical validation.)

Analogues in Neurological Systems

Various lessons can also be drawn from looking at biological versions of FE processing. That the “design” processes of evolution created such examples is further testimony that the underlying principles are important. The benefits these “designs” confer in nature are instructive for understanding the benefits they might have for systems that humans design.

Fovea+Periphery in Human Vision

The fisheye DOI is implemented in human vision, though there is no distortion involved. Spatial resolution on the retina varies dramatically, by more than a factor of ten from the fovea to the periphery [20]. By garnering detail only in the fovea, basically extracting a FE subset, the information that must be transmitted to the brain is dramatically reduced, and the sensory apparatus made much lighter and more mobile.

Memory

Human memory mechanisms also implement a kind of FE design. Numerous models of human memory posit mechanisms that can be thought of as analogs of the distance and *a priori* components of the FE-DOI. Memory has a major semantic associational structure, where if a given concept is activated, associated concepts are also activated. For example, after reading the word *doctor*, associated words like *hospital*, *nurse*, *medicine*, *surgery*, and *disease* are more readily picked out of noise, distinguished from nonsense words, and spontaneously produced, compared to other, random words like *toast*, *flower*, *lamp*, or *nickel*. Such effects (and cognitive theories about them) relate to aspects of memory that reflect the current task needs of the user, specific to their current focus. (Ironically these are called “context” effects - though here they represent the focus part of the FE effect). There are also components that are more independent of the current focus - frequency, recency and importance effects. Items that occur more often in the world (car, tree), or are of greater salience or significance (fire-alarm) are more available to be evoked cognitively. The former effects have been modeled by spreading activation in semantic networks, the latter by changes in initial activation level. These two trade off against one another, so important items associated with the current context are even more likely.

The net result is that the set of items in memory that are most readily available are a FE subset of memory: items that are either close to the current semantic focus, or if not, of increasingly high *a priori* importance.

The relevance of this to HCI comes from a design rationale analysis. The “design” of human memory has been analyzed by Anderson [1], justifying the observed capabilities of human memory in terms of how they meet requirements of living in the world. He argues that, regardless of how it actually works, memory does a good

job on the task set before it, namely to make available things from the past that are needed in the present. To do this memory must, in effect, be constantly updating its estimate of the “Need Probability” for everything it has saved. Anderson decomposes the *need probability* function into two major components: temporal and contextual. One of the major temporal components is frequency – things needed often in the past are more likely to be needed in the next moment. The context component of *need probability* says things that are associated with other things a person is currently working with are also likely to be needed. These components were explicitly chosen to explain the memory phenomena cited earlier. For the purpose of this paper, the point is that Anderson’s rational analysis explains why **any** memory system *should* provide a kind of fisheye subset. Such a *need probability* analysis applies to many other circumstances of interest to HCI: visualization design, information architecture, etc. Anderson’s decomposition into associative and temporal components is likely to be quite general, and a FE aspect to the resulting design is to be expected.

Discussion of Sources of FE Importance

These analyses provided several reasons why the FE-DOI subsets might be important for rational action. The subsets show the full set of entities that have influence over your current position. They reflect the decreasing need for detail in hierarchical NDSs. They allow the appropriately contextualized interpretation of otherwise locally ambiguous items. They afford efficient traversal to, and effective navigational information about, other possible foci. And, not surprisingly, they mimic sophisticated biological design in both vision, where they provide data compression and smaller, lighter hardware, and in memory, where they mirror Anderson’s need probabilities.

The differences between these theories have important implications for design. The human eye provides a nice example. By the Nearly Decomposable Systems arguments, more remote things should require not just less spatial but also lower temporal resolution. In vision, however, the periphery gets lower spatial, but *higher* temporal resolution. Why might this be? The argument would be that retinal structure is actually *not* reflecting an NDS, but an information and hardware abbreviation strategy. In fact, objects in the external world peripheral to your current visual focus, an attacking predator or an oncoming car, for example, may in fact become tightly coupled to you. The peripheral sensitivity to change allows the low spatial frequency information there to be able to indicate where additional high spatial frequency info is needed, so that you can shift your focus appropriately (basically a view navigation argument).

The point is that different situations may call upon different aspects of these functions, and the design must be tailored accordingly.

CONCLUSION

To enhance efforts to deal with the problems posed by ever growing information worlds, this paper has used the *Generalized Fisheye View* formulation to put more depth behind the concept of focus + context. It was argued that the formulation's value has little to do with the specific distortion metaphor of a photographic fisheye lens. What really matter are the Fisheye DOI subsets; these are what guarantee *what* is shown (details near the current focus but only increasingly important features further away) regardless of *how* it is shown. It was demonstrated that, in a rigorous sense, a variety of focus+context techniques, including distortion, zoom, and closeup+overview, all show these FE subsets, differing mostly in how they make other ancillary design tradeoffs. The generality of the idea of the FE-DOI was then reiterated, highlighting its possibilities for non-visualization uses in dealing with our increasingly large information worlds. Finally, several theories were offered that rationalize the importance of the FE-DOI subsets. These theories provide possible substance for design rationales, as designers try to understand the user's needs in tasks that must deal with large information worlds.

Future work awaits to explore various non-visual fisheye presentations, and to validate empirically any of the theories about *why* the FE-DOI is useful in various circumstance of importance to HCI.

ACKNOWLEDGMENTS

The author would particularly like to thank Maria Slowiaczek for her invaluable help in rescuing early drafts of this manuscript from expository chaos, and the thoughtful comments of anonymous reviewers.

REFERENCES

1. Anderson, J.R. (1989) Human memory: an adaptive perspective. *Psych. Rev.* 96(4), 703-719.
2. Apperley, M., Tzavaras, I., Spence, R. (1982) A bifocal display technique for data presentation, *Proc. EuroGraphics'82*, 27-43.
3. Ashdown, M. & Robinson, P. (2005) Escritoire: a personal projected display, *IEEE MultiMedia*, 12(1), 34-42.
4. Baudisch, P., Good, N., Stewart, P. (2001) User Focus plus context screens: combining display technology with visualization techniques. *Proc. ACM UIST 2001*, 31-40.
5. Bederson, B. & Hollan, J. (1994) Pad++: a zooming graphical interface for exploring alternate interface physics, *Proc. ACM UIST 1994*, 17-26.
6. Card, Stuart, Mackinlay, Jock, Shneiderman, Ben (1999) *Readings in Information Visualization: Using Vision to Think*, Morgan Kaufman.
7. Carpendale MST, Montagnese C, (2001) A framework for unifying presentation space, *Proc ACM UIST 2001*, p61-70.
8. Chen, Chaomei (2004) *Information visualization: beyond the horizon*, 2nd Ed. New York: Springer.
9. Farrand, W. (1973) *Information Display in Interactive Design*, Doctoral Thesis, Department of Engineering, University of California at Los Angeles.
10. Furnas, G. (1982) The FISHEYE view: A new look at structured files. *Bell Laboratories Technical Memorandum*, #82-11221-22, Oct 18, 1982. (Reprinted in [6].)
11. Furnas, G. (1986) Generalized fisheye views. *Proc. ACM CHI'86*, 16-23.
12. Furnas, G., Zacks, J. (1994) Multitrees: enriching and reusing hierarchical structure. *Proc. ACM CHI'94*, 330-336.
13. Furnas, G., (1997) Effective View Navigation, *Proc. ACM CHI'97*, 367-374.
14. Gutwin, C. & Skopik, A. (2003) Fisheyes are good for large steering tasks. *Proc. ACM CHI 2003*, 201-208.
15. Keahey, T., (1998) The Generalized Detail-In-Context Problem, *Proc. IEEE InfoVis 1998*, 44-51.
16. Lamping, J. & Rao, R. (1994) Laying out and visualizing large trees using a hyperbolic space, *Proc. ACM UIST 1994*, 13-14.
17. Leung, Y., & Apperley, M. (1994) A review and taxonomy of distortion-oriented presentation techniques. *ACM TOCHI*, 1(2), 126-160.
18. Lieberman, H. (1994) Powers of ten thousand: navigating in large information spaces, *Proc. ACM UIST 1994*, 15-16.
19. Mackinlay, J., Robertson, G., Card, S. (1991) The perspective wall: detail and context smoothly integrated. *Proc ACM CHI'91*, 173-179
20. Navarro, R., Artal, P., Williams, D. (1993) Modulation transfer of the human eye as a function of retinal eccentricity, *J.Opt.Soc.Am.A.*, 10(2), 201-212.
21. Pirolli, P, Card, SK, Van Der Wege, MM (2003) The effects of information scent on visual search in the hyperbolic tree browser, *ACM TOCHI*, 10(1), 20-53.
22. Robertson. G. & Mackinlay, J. (1993) The Document Lens Visualizing Information. *Proc. ACM UIST 1993*, 101-108.
23. Sarkar, M., & Brown, M. (1994) Graphical fisheye views. *Comm.ACM*, 37(12):73-84
24. Sarkar, M., Snibbe, Scott S., Tversky, O., Reiss, S. (1993) Stretching the Rubber Sheet: A Metaphor for Visualizing Large Layouts on Small Screens Visualizing Information. *Proc. ACM UIST 93*, 81-91.
25. Simon, H., (1996) *The Sciences of the Artificial*, 3rd Ed., Cambridge: MIT Press. 197ff.
26. Spence, Robert (2001) *Information visualization*. Harlow: Addison-Wesley.
27. Ware, Colin (2004) *Information visualization: perception for design*. San Francisco: Morgan Kaufman.
28. Wood, R. (1906) Fish-eye views, and vision under water. *Philosophical Magazine*, 12(6):159-162.