

Toward Disambiguating Multiple Selections for Frustum-Based Pointing

Greg Schmidt^{1,2}

Dennis G. Brown¹

Erik B. Tomlin¹

J. Edward Swan II^{1,3}

Yohan Baillot^{1,2}

¹3D Virtual & Mixed Environments Lab, U.S. Naval Research Laboratory*

²ITT Industries, Advanced Engineering & Sciences[†]

³Department of Computer Science & Engineering, Mississippi State University[‡]

ABSTRACT

Selection is a fundamental user operation in 3D environments. These environments often simulate or augment the real world, and a part of that simulation is the ability to select objects for observation and manipulation. Many user interfaces for these applications depend on six-degree-of-freedom tracking devices. Such devices have limited accuracy and are susceptible to noise, giving an imprecision that makes object selections difficult and hard to repeat. This difficulty is amplified when the user's viewpoint is also tracked, meaning the user must compensate for noise from both the head tracker and the pointing device when performing object selection. Also, users may experience fatigue when using handheld pointing devices for extended periods, creating error even if the tracking technology were perfect.

This paper presents a pointing-based probabilistic selection algorithm that addresses some of the ambiguities associated with tracking and user imprecision. It performs multiple selections by considering a frustum along the user's pointing direction and the hierarchical structure of the database. It assigns probabilities that the user has selected particular objects using a set of low-level 3D intersection-based selection techniques and the relationship of the objects in a hierarchical database, and makes the final selection using one of several weighting schemes. We performed several experiments to evaluate the low-level selection techniques, tested several weighting schemes for the integration algorithm, and we show that the algorithm is effective at disambiguating multiple selections.

CR Categories: H.5.2 [Information Interfaces and Presentation]: User Interfaces—Interaction styles; H.5.3 [Information Interfaces and Presentation]: Group and Organizational Interfaces—Computer-supported cooperative work;

Keywords: interaction, selection, algorithms, augmented reality, virtual reality, hierarchical databases

1 INTRODUCTION

Many virtual and augmented reality systems present the user with a rendering of a 3D world containing distinct objects that the user can query or manipulate. To perform these actions on objects, the user usually must first select the object. While there are many ways to select objects, pointing at the desired object is a common and natural way to select. Selection by pointing can happen using a range of devices, from a common 2D mouse controlling a cursor on a 2D projection of the 3D world, to a full six-degree-of-freedom (DOF) hand-held tracking device. Selection can also happen without using the hands at all, by allowing the user to select using head orienta-

tion (assuming the head is tracked) or gaze direction using an eye tracker.

We assert that all user selection operations are susceptible to error. First, there is human error: the imprecision that comes from lack of experience, not enough motor control to do fine grained selection, or fatigue developed during a session; Wingrave et al. [13] studied a number of correlations between certain attributes of users and their ability to perform selections. Second, there is equipment error, which could be noise, drift, and lag in a 6DOF tracking system, or simply not enough resolution on a wheel-based device to perform a fine selection. Finally, there are ambiguities associated with the scene itself, such as when the user tries to select one object occluded by another object. In our main application area, mobile augmented reality, this is a common problem because users have “x-ray vision” and can see spatial information, such as the position of a collaborator, that may be occluded by real or virtual objects—in this example, the collaborator may be behind a building. These errors can lead to selections that are totally incorrect, such as when using a ray-based selection that chooses a single object, or to ambiguous results when using multiple selection techniques that can choose many candidate objects.

We designed a pointing-based probabilistic selection algorithm that alleviates some of the error in user selections. This technique takes into consideration the hierarchical structure of the scene objects (e.g., a door is a child of a wall, which is a child of a building, and so on). It assigns probabilities that the user has selected particular objects, within a frustum along the user's pointing direction, using a set of low-level 3D intersection-based selection techniques and the relationship of the objects in a hierarchical database, and makes the final selection using one of several weighting schemes. We implemented this algorithm in our virtual and augmented reality application framework and performed several experiments to evaluate the low-level selection techniques, to evaluate several weighting schemes for the integration algorithm, and to show that the algorithm can effectively disambiguate multiple selections.

After describing related work, we describe the low-level selection techniques. We then present the design and discuss the results of the experiments described above.

2 RELATED WORK

Selection in 3D environments has been an active research topic since the first virtual environments were implemented. Hinckley et al. [4] presented a survey of, and a common framework for, techniques for 3D interaction, until that point in time. Liang and Green [8] developed the spotlight method of selection, which alleviated some issues with using ray-based selection for small and far objects, but introduced the problem of multiple selections, for which they set up rules to choose one of the possible selections. Mine [9] described a few techniques for selection, along with other interactions to be supported in virtual environments. Forsberg et al. [3] developed two novel selection techniques, aperture (an extension of spotlight) and orientation, to deal with the imprecision of ray-based selection using a 6DOF input device. Pierce et al. [11] introduced a set of selection techniques using tracked hands and the

*{gregory.schmidt, dennis.g.brown, erik.tomlin}@nrl.navy.mil

[†]baillot@ittid.com

[‡]swan@cse.msstate.edu

Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE 2006		2. REPORT TYPE		3. DATES COVERED 00-00-2006 to 00-00-2006	
4. TITLE AND SUBTITLE Toward Disambiguating Multiple Selections for Frustum-Based Pointing				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) 3D Virtual & Mixed Environments Lab,U.S. Naval Research Laboratory,Washington,DC,20375				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES 3D User Interface Symposium, Alexandria, VA, 12--15 March 2006.					
14. ABSTRACT					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 9	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

2D projection of a 3D scene. In his thesis, Bowman [1] gave a thorough survey of 3D selection techniques at the time, and introduced a novel technique for selection and manipulation.

In more recent work, researchers have dissected the task of selection even further, and have created novel selection techniques for specific application domains such as Augmented Reality (AR) and multimodal systems. Wingrave et al. [12] discovered that users do not have an internal model of how they expect the environment to behave (for example, in performing selections), but instead they adapt to the existing environment using feedback received when performing tasks. Olwal et al. [10] developed some of the first algorithms for selecting in AR. Their technique attaches a virtual volumetric region of interest to parts of a user’s body. When the user moves the body part, interacting with objects in the environment, a rich set of statistical data is generated and is used for processing selections. Kolsch et al. [7] developed a real-time hand gesture recognition system that can act as the sole input device for a mobile AR system. Kaiser et al. [6] developed mutual disambiguation techniques and evaluated their effectiveness for 3D multimodal interaction in AR and Virtual Reality (VR). They showed that mutual disambiguation accounts for over 45% of their system’s successfully recognized multimodal commands.

We acknowledge that multimodal systems are an effective way to alleviate some selection issues (for example, the user points to a group of objects that includes a window, and says the word “window,” giving the system the means to disambiguate the selection), and we have implemented multimodal (speech and pointing) processing in our system for disambiguating different types of objects (for example, windows, walls, and buildings). However, if there is more than one object of the same type (for example, four windows) in the selection space, then the system described above will fail since the utterance “window” is ambiguous and does not help correct a wrong selection. In this example, either more sophisticated speech semantics or pointing-based selection processing techniques are needed.

In this paper, we have chosen to focus solely on improving pointing-based selection. Our selection algorithm introduces the concept of executing multiple selection techniques in parallel and choosing the final selection from the results of those techniques using a weighting scheme. The low-level techniques and the integration algorithm are described in the next section.

3 PROBABILISTIC POINTING-BASED SELECTION ALGORITHM

Selection by pointing in 3D environments is inherently imprecise when the user is allowed to select occluded objects—the user may have the impression of pointing to a specific object, for example, but the system may not know for sure which object in the pointing direction is meant to be selected. Users often make pointing errors, especially when selecting small objects, objects at a distance, or when trying to make a selection quickly. Furthermore, pointing provides the object’s direction, but not distance, so when several objects lie in the direction the user is pointing, it remains unclear which object the user intended to select.

To deal with selection ambiguity, we designed a probabilistic selection algorithm that generates lists of candidate objects the user may have meant to select, and probability estimates of how likely it is the user meant to select each object. The algorithm combines several intersection algorithms and the hierarchical structure of the dataset, and then integrates the resulting candidate selections. The processing steps of the algorithm are shown in Figure 1, which we describe in each of the following sections.

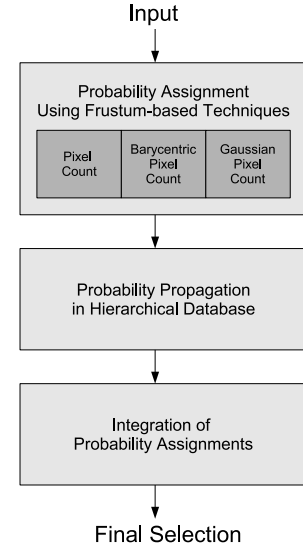


Figure 1: Flow of the pointing-based selection algorithm.

3.1 Frustum Intersection Algorithms

We have designed a set of three algorithms which attempt to mitigate the ambiguity associated with ray intersection. Each algorithm is based on the concept of rendering the scene into a small *selection frustum* (see Figure 2); the rendered scene is viewed by a camera pointing coincident to the selection ray, and the frustum is rendered into an off-screen image buffer. The algorithms then count and classify the pixels in this buffer, and use these counts to create a list $(o_1, p_1), \dots, (o_n, p_n)$ of potentially selected objects o_i and associated selection probabilities p_i .

As described in more detail below, each of these algorithms has differing utility depending on the user’s preferences for making selections, on what type of object the user is trying to select, and on its relationship to other objects in the scene. We have designed the three intersection algorithms such that each has a different user preference for selection. These preferences are: (1) select the item nearest the central pointing ray; (2) select the largest item in the viewing frustum; and (3) select using a combination of the two other approaches. We wanted to find out if having several algorithms available based on different user preferences increases the chances for correctly selecting objects. These algorithms could either be used individually or executed in parallel and their results integrated together.

We describe each intersection algorithm in more detail, and then show how their output lists of candidate selections are integrated when the algorithms are run in parallel.

Pixel-Count The PIXEL-COUNT algorithm preferentially orders objects according to their projected size in the selection frustum. PIXEL-COUNT simply counts the number of pixels occupied by each object, and weighs the objects accordingly. This pixel-counting technique is a very fast way of implementing ordering of objects by projected size. A similar technique has been reported by Olwal [10].

PIXEL-COUNT

Input: 3D direction

Output: list $(o_1, p_1), \dots, (o_n, p_n)$ of candidate objects o_i and associated probabilities p_i

- 1 calculate a small frustum about 3D direction
- 2 **for** each object o_i in the frustum
- 3 | render o_i into the frustum

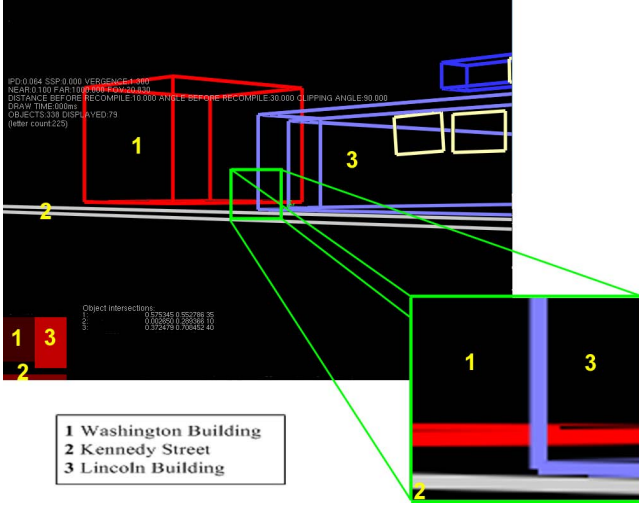


Figure 2: The operation of the PIXEL-COUNT algorithm. The scene is rendered into the small selection frustum shown in the center of the image.

- 4 $\text{pix}_i \leftarrow$ number of pixels covered by o_i
- 5 weights $w_i \leftarrow \text{pix}_i / \text{total-frustum-pixels}$
- 6 assign probabilities p_i from weights w_i
- 7 sort (o_i, p_i) list by decreasing probabilities p_i

Figure 2 demonstrates PIXEL-COUNT. The green square in the center of the image demonstrates one size of the selection frustum; in the lower-right the frustum contents are enlarged. Note that the frustum size may be adjusted by the user. The square in the lower left corner shows a low-resolution re-rendering of the frustum contents.

The PIXEL-COUNT algorithm is robust to noise and pointing ambiguity. However, it inherently assumes the user is attempting to select larger objects, and it does not work well for selecting small objects near larger objects.

Barycentric-Pixel-Count This algorithm was motivated by our observation that users tend to point toward the center of the visible part of the object they wish to select. Figure 3 describes how BARYCENTRIC-PIXEL-COUNT operates. The algorithm calculates the center point of the visible portion of each object (O_1^c and O_2^c), and then determines the distance to the center of the selection frustum (d_1 and d_2). It then weighs each object's pixels with the inverse of this distance; so in Figure 3 object 1's pixels are weighted by $1/d_1$ and object 2's pixels by $1/d_2$. Since it is assumed that the user is intending to look at one object, probabilities are estimated by normalizing the weights across all of the weighted pixels.

BARYCENTRIC-PIXEL-COUNT

Input: 3D direction
Output: list $(o_1, p_1), \dots, (o_n, p_n)$ of candidate objects o_i and associated probabilities p_i

- 1 calculate a small frustum about 3D direction
- 2 let F^c be the center of the frustum
- 3 **for** each object o_i in the frustum
- 4 let O_i^c be the center of the visible portion of o_i
- 5 $\text{bary-weight} \leftarrow 1 / \|F^c - O_i^c\|$
- 6 render o_i into the frustum
- 7 **for** each pixel a generated by o_i
- 8 $\text{pix}_i \leftarrow \text{pix}_i + a * \text{bary-weight}$
- 9 weights $w_i \leftarrow \text{pix}_i / \text{total-frustum-pixels}$

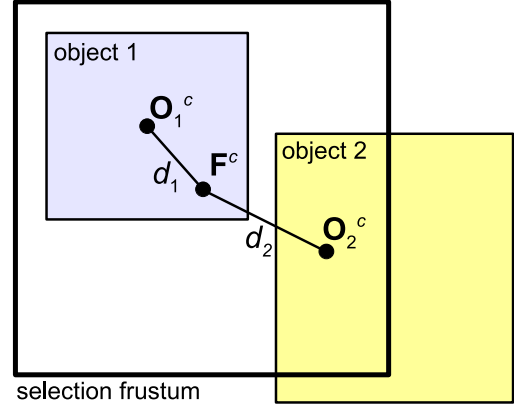


Figure 3: The operation of the BARYCENTRIC-PIXEL-COUNT algorithm. Because $d_1 < d_2$, each pixel of object 1 will be weighted more heavily than the pixels of object 2.

- 10 assign probabilities p_i from weights w_i
- 11 sort (o_i, p_i) list by decreasing probabilities p_i

BARYCENTRIC-PIXEL-COUNT works very well for selecting small objects near larger objects, but it does not work well if the user points away from the center of an object, or if the object has a shape such that the Barycentric center does not lie within the object itself.

Gaussian-Pixel-Count The GAUSSIAN-PIXEL-COUNT algorithm is also motivated by the general observation that users tend to center the objects they want to select. However, this algorithm tries to address the failing of the BARYCENTRIC-PIXEL-COUNT algorithm, which occurs when the Barycentric center does not lie within the object itself. GAUSSIAN-PIXEL-COUNT operates by applying a Gaussian mask, centered in the selection frustum, to each object's pixels. The mask operates, in effect, by assigning weights to each pixel based on its distance from the center ray according to a Gaussian bell curve. Figure 4 describes how GAUSSIAN-PIXEL-COUNT operates. The filtered output for each individual object is combined in an accumulation buffer. Probabilities are assigned, again, assuming one object is intended to be selected, by normalizing across the weighted pixels.

GAUSSIAN-PIXEL-COUNT

Input: 3D direction
Output: list $(o_1, p_1), \dots, (o_n, p_n)$ of candidate objects o_i and associated probabilities p_i

- 1 calculate a small frustum about 3D direction
- 2 calculate a Gaussian filter G centered in frustum
- 3 **for** each object o_i in the frustum
- 4 render o_i into the frustum
- 5 **for** each pixel a generated by o_i
- 6 $\text{pix}_i \leftarrow \text{pix}_i + a * G$
- 7 weight $w_i \leftarrow \text{pix}_i / \text{total-weighted-frustum-pixels}$
- 8 assign probabilities p_i from weights w_i
- 9 sort (o_i, p_i) list by decreasing probabilities

The algorithm is less susceptible to being biased by large visible objects and it favors selecting objects near the central viewing ray.

3.2 Probability Propagation in Hierarchical Database

The probability estimates generated by the ray intersection algorithms assume that a single object occupies a given space in

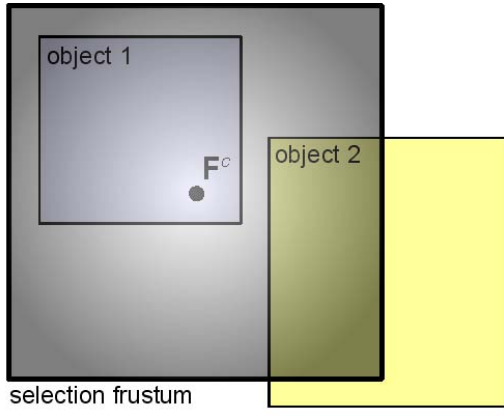


Figure 4: The operation of the GAUSSIAN-PIXEL-COUNT algorithm. Pixels in objects 1 and 2 are weighted by a circularly symmetric Gaussian function centered at F^c .

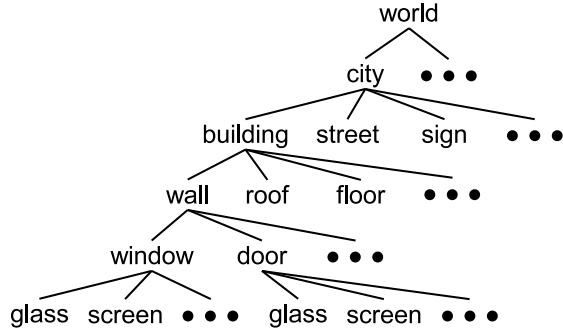


Figure 5: A portion of a hierarchically-organized database that contains urban data such as buildings, streets, signs, etc.

the viewing frustum. This assumption is not the case for a hierarchically-organized database, which contains objects composed of smaller objects, which in turn, may be composed of even smaller objects, and so forth. This type of database can have several objects occupying a given space, though there will always be a relationship between the occupying objects. An example of a hierarchically-organized database, which we use in our testing, and the inter-relationships between objects, is illustrated in Figure 5.

In order to assign probabilities properly for a hierarchically-organized database, we (1) set the ray-intersection algorithms to probabilities for the lowest level structure for each pixel, and (2) propagate the probabilities up the tree hierarchy from the leaf nodes. Since our ray intersection algorithm returns the lowest-level structures, we use the following algorithm to propagate probabilities up the database hierarchy:

PROBABILITY-PROPAGATION

Input: old list $L_O = (o_1, p_1), \dots, (o_n, p_n)$ of objects o_i and associated probabilities p_i

Output: new list $L_N = (o_1, p_1), \dots, (o_m, p_m)$

```

1 create empty list  $L_N$ 
2 for each object  $o_i$  in  $L_O$ 
3   if  $o_i$  not in  $L_N$  then
4     add  $o_i$  to  $L_N$ 
5      $o_i.weight \leftarrow p_i$ 
6   for each recursive parent of  $o_i$ 
7     if  $o_i.parent$  not in  $L_N$  then
8       add  $o_i.parent$  to  $L_N$ 
9      $o_i.parent.weight \leftarrow o_i.parent.weight + p_i$ 

```

```

10 for each object  $o_i$  in  $L_N$ 
11   normalize  $o_i.weight$ 
12   assign probability  $p_i$  from  $o_i.weight$ 
13 sort  $L_N$  by decreasing probabilities

```

One subtlety to note is in line 5, where we consider the probability for each pair as a “weight” for that pair, since we perform operations with these values that do not strictly consider the values as probabilities. The resulting probability assignments estimate likelihood that any of the occupying objects for a given space is the desired selection. For example, using the hierarchy in Figure 5, for a ray intersecting a window, its probability of selection is equal to the probability of selecting the wall, building, city, and world, at the intersecting pixel. Another property to understand using this probability propagation approach is the probabilities for the parents are always at least as much as the probabilities for any child. Keep in mind, another reasonable and equally valid manner of assigning probabilities is to consider the hierarchical nature of the database at a lower level, when the ray intersection algorithms estimate the probabilities. This approach may lead to different but still similar assignments of probabilities.

3.3 Integration of Probability Assignments

The three lists of objects and associated probabilities generated by the ray intersection algorithms and probability propagation algorithm need to be combined into one list. One caveat to this process is that each list may contain a slightly different set of object-probability pairs due to differences in the how each algorithm operates. Thus, the elements of the lists will have to be matched and like items combined. Another important note is that, in the process described below, just as in the probability propagation algorithm, we consider the probability for each pair as a “weight” for that pair, since we perform operations with these values that do not strictly consider the values as probabilities. A naive integration approach would be, for each object, to simply average the weights assigned to that object from the different algorithms. This approach, however, does not take into consideration the strengths and weaknesses of each of the three algorithms. A more appropriate way to integrate them is to assign a weight to each algorithm, W_i , based on how well each performs in comparison to the others. The lists are then integrated by the following WEIGHTED-INTEGRATION algorithm.

WEIGHTED-INTEGRATION

Input: 3 lists $L_G, L_B, L_P = (o_1, p_1), \dots, (o_n, p_n)$ of objects o_i and associated probabilities p_i

Output: new list $L_N = (o_1, p_1), \dots, (o_m, p_m)$

```

1 create empty list  $L_N$ 
2 for each list  $L_j$  in  $L_G, L_B, L_P$ 
3   for each object  $o_i$  in  $L_j$ 
4     if  $o_i$  not in  $L_N$  then
5       add pair to  $L_N$ 
6     else
7       find  $o_i$  in  $L_N$ 
8        $L_N.o_i.weight \leftarrow p_i * W_j$ 
9   for each object  $o_i$  in  $L_N$ 
10    normalize  $o_i.weight$ 
11    assign probability  $p_i$  from  $o_i.weight$ 
12 sort  $L_N$  by decreasing probabilities

```

The integration weights, $W_i, i = G, B, P$, corresponding to the intersection algorithms GAUSSIAN-PIXEL-COUNT, BARYCENTRIC-PIXEL-COUNT, and PIXEL-COUNT, respectively, are initially arbitrarily assigned to $\frac{1}{3}$ each (giving the same effect as the naive

method of averaging). However, we acknowledge that having proper weight assignments for data integration is important for optimizing performance and is a difficult task. We made several attempts to refine the weight assignments. We used the performance estimates of the intersection algorithms to influence the assignment of the weights in the integration. In one case, we normalized the algorithm performance estimates and used those as the weight values. For the other case, we used the normalized performance estimates as a guide to refine the weight assignments. More details about the assignment of weights is given in Section 4.

4 PERFORMANCE EVALUATION

We conducted several experiments to gain a better understanding of the effectiveness of the algorithms for disambiguating multiple pointing selections. We approached this task by first comparing empirically the three intersection algorithms head-to-head to learn the strengths and weaknesses of each algorithm for specific dataset cases. Second, we evaluated the integration algorithm, exploring several weighting schemes, to determine how best to utilize combinations of the intersection algorithms in parallel. Lastly, we conducted a short experiment to demonstrate and empirically evaluate how well the algorithms work for disambiguating selections.

4.1 Comparison of Intersection Algorithms

We conducted three experiments to compare the three intersection algorithms head-to-head. The first experiment was designed to test the experimental protocol, flushing out any design and testing issues, and used a simple real-world urban dataset and a few test cases. Each successive experiment increased the detail of the datasets and complexity of the test cases. Comparing any algorithm thoroughly typically requires running an extensive set of experiments with multiple datasets and conditions. The goal of these experiments was to get a general feel for the accuracy of each algorithm for performing selection using a real-world urban dataset.

Since our overall goal is to apply these algorithms for disambiguating multiple selections, we acknowledge that precision plays a role in how we evaluate the algorithms. In particular, the selection cases we address only become interesting when high precision is required, otherwise the selections could be easily performed using well-known ray-based pointing techniques. We ran a set of preliminary tests to evaluate the degree of precision needed for the experiments and tested the selection techniques on a range of sizes for objects and varying amounts of space between each. Some of the test cases are shown in Figure 6. We used the precision estimates to guide our choice of the test cases for the experiments, making sure that the required degree of precision was high, but low enough to collect meaningful data to compare the three algorithms. Later, when we tested the algorithms for disambiguating selections of smaller, distant objects, we increased the degree of precision required. We next describe the experiment preparations, experimental protocol, and each experiment.

Experiment Preparation We prepared for the experiments by implementing the algorithms and user interface in a combined AR and VR system our laboratory has been developing over the last several years. The two major building blocks of the combined system are the Battlefield Augmented Reality System (BARS) [5], which provides support for 3D in the form of VR and AR (with special emphasis on mobile AR and multi-wall VR), and Quickset [2], which provides support for 2D map-based interaction using an agent-based architecture that supports probabilistic, asynchronous input events. The system has evolved extensively over the last several years to support a large number of input and display devices, interaction and display algorithms, as well as the interface techniques and associated algorithms described in this paper. The

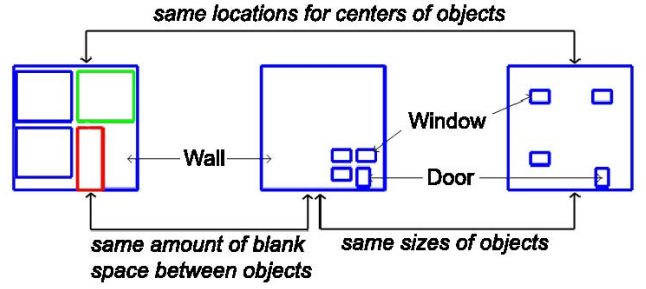


Figure 6: Three of the test cases used to determine the degree of precision needed by the experiments. We varied the size of several windows, and the spacing between each, and asked several users to perform selections of the windows for each test case. Statistics were collected to determine the precision required for our experiments.

major components we added are the pointing-based user interface, the associated selection algorithms, a speech recognition system, and the Quickset multimodal speech-and-gesture integrator.

Experimental Protocol For each of the three experiments, we recruited two to four subjects from our development team. Each subject wore a tracked, see-through head-mounted display (HMD) that included a microphone for the speech recognizer. The orientation of the head was used as the pointing direction. The system showed a cross-hair cursor at the center of the view frustum for the pointing direction. The view frustum’s borders were made invisible to the user and its size was kept fixed throughout the experiments. The experiments operated under the assumption that the user’s selection strategy was to point directly through the center of the object intended to be selected. Other strategies, such as selecting the largest item in the view frustum, are evaluated later in Section 4.3.

Each subject was given a training session to become familiar with the pointing and speech input interaction mechanism. The subjects were prompted by the experiment administrator to perform selections using head direction combined with a voice command. For example, the subject is asked to point at the upper right window of the left building, shown in Figure 6, while speaking “secure this window.” The system responds by changing the color, based on the voice command, of what the system determines to be the correct selection. In Figure 6, the verb “clear” triggers a change in the color of the upper right window to green.

The trials of each experiment asked the subjects to perform a sequence of selections over a range of test cases. The selection test cases presented were sets of windows, doors, walls, and buildings. The speech actions included “danger on this object,” “clear this object,” and “reset this object” — these commands are meaningful in the urban situational awareness domain of the BARS system [5]. Data were collected for the three frustum-based algorithms. Cases where the speech recognizer failed were thrown out, since we are not evaluating the speech recognizer nor the multimodal features of the system. The test cases were presented in a counter-balanced manner in order to eliminate any learning-effect biases.

Experiment 1 The first experiment used a small subset of the dataset shown in Figure 7. The dataset contains semantic information about the buildings surrounding our laboratory—it is a real-world database used by our system for various development, demonstration, and evaluation purposes. We ran the experiment, following the experimental protocol above, using two subjects. Each subject was asked to perform selections of windows and doors in three different buildings. We collected statistics for 12 different cases per user, for a total of 24 cases. We recorded the *accuracy* of the intersection algorithms; we recorded a boolean value of ‘1’ if an algorithm made a correct selection, and a ‘0’ if it made an incorrect selection. Figure 8 shows the accuracy performance of each of the



Figure 7: The dataset used in Experiments 1–3. Progressively smaller subsets of the dataset, with different test cases, were used in Experiments 1 and 2. (Left) Ground-level views; these are the views that subjects saw during the experiments. (Right) Elevated view; given to give the reader a feel for the dataset’s layout.

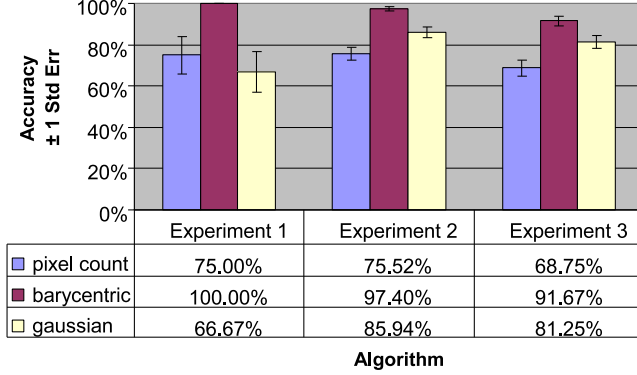


Figure 8: The accuracy (mean percentage of correct selections) given by the three intersection algorithms in Experiments 1–3.

three algorithms for Experiment 1.

Experiment 2 The second experiment expanded on the first experiment, this time using a larger dataset and a broader set of test cases. As before, the dataset chosen is a subset of Figure 7’s dataset. We only presented the ground-level view of the dataset to the subjects. The test cases were determined based on observed strengths and weaknesses of the three intersection algorithms. We designed the test protocol such that the subjects would make selections for situations where each of the algorithms is weak, and cases where each is strong; we tested an equal number of “good” and “bad” test cases for each algorithm. We used four subjects and collected the performance statistics for 192 selection cases (48 per subject). The users were asked to make selections of windows and buildings. We recorded the accuracy performance for each of the three algorithms; the results for Experiment 2 are shown in Figure 8.

Experiment 3 The third experiment used the largest portion of our test dataset (Figure 7), as well as the broadest set of test cases, which we believe better explored the strengths and weaknesses of the intersection algorithms. This experiment followed the same protocol as Experiment 2, with three subjects and 144 selection cases (48 per subject), and an equal amount of good and bad test cases for each algorithm. We again recorded the accuracy performance for each of the three algorithms; the results are shown in Figure 8.

Results and Discussion We analyzed the accuracy performance results with a one-way analysis of variance (ANOVA). In addition to the standard p -values, the standard measure of *effect significance*, we calculated and report ω^2 , a standard measure of *effect size*. ω^2 is an approximate measure of the percentage of the observed variance that can be explained by the effect.

As suggested by Figure 8, we found a strong effect of algorithm for each of the experiments (Experiment 1: $F(2, 69) = 5.07$, $p = .009$, $\omega^2 = 10.3\%$; Experiment 2: $F(2, 573) = 20.7$, $p < .000$, $\omega^2 = 6.4\%$; Experiment 3: $F(2, 429) = 12.7$, $p < .000$, $\omega^2 = 5.2\%$). The error bars in Figure 8, which show ± 1 standard error,

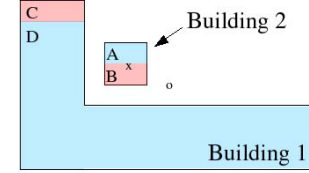


Figure 9: The effect of BARYCENTRIC-PIXEL-COUNT on selection involving a concave shape. The blue regions work properly, while the red fail. Empirical results show regions B and C fail 85% of the time.

indicate that BARYCENTRIC-PIXEL-COUNT had greater accuracy than both GAUSSIAN-PIXEL-COUNT and PIXEL-COUNT for all three experiments. For Experiments 2 and 3, GAUSSIAN-PIXEL-COUNT had greater accuracy than PIXEL-COUNT.

This analysis empirically validates that the choice of intersection algorithm makes a difference in selection accuracy. Approaches similar to what we are calling Barycentric have been presented as Liang and Green’s spotlight method [8] and the aperture method of Forsberg et al. [3], but here we present the first empirical evidence for the general effectiveness of this technique.

However, although the data show that the BARYCENTRIC-PIXEL-COUNT algorithm clearly outperforms the other algorithms (for this choice of dataset and degree of pointing precision), we can observe some interesting test cases if we allow the user to view the same dataset from above the buildings. For example, we observed that concave shape patterns show up more often looking from above than from the ground-level views—see Figure 9. One weakness of the BARYCENTRIC-PIXEL-COUNT algorithm is how it operates on concave objects. The algorithm relies on an estimation of the center of the objects trying to be selected. As seen in Figure 9, one estimation of the center of Building 1 is located at the position where the ‘o’ is shown. Due to the complex nature of the behavior of the algorithm’s weighting function, we decided to empirically evaluate the test case. We ran a quick study collecting statistics for how well different regions of the buildings could be selected properly using the BARYCENTRIC-PIXEL-COUNT algorithm. Pointing in the regions shown in blue properly select the correct building, while the red regions fail. Empirical results show regions ‘B’ and ‘C’ fail 85% of the time.

4.2 Evaluation of Integration

We conducted three experiments to evaluate different weighting schemes for the integration algorithm. Our objective was to determine how to best utilize different combinations of the intersection algorithms, including determining when to use an intersection algorithm by itself or when to combine the algorithms in parallel. We focused mainly on finding an optimal weight assignment for the typical use case and compared it against the best intersection algorithm (BARYCENTRIC-PIXEL-COUNT). In each experiment, we used the datasets and experimental protocol from the experiments in Section 4.1. We applied several different weighting schemes (de-

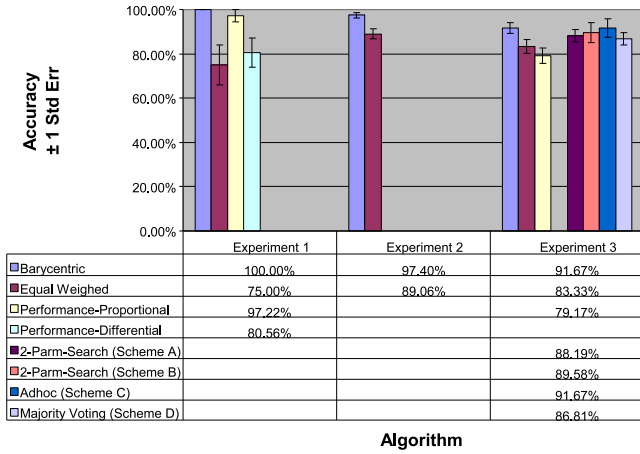


Figure 10: The accuracy of the integration schemes versus the BARYCENTRIC-PIXEL-COUNT intersection algorithm in Experiments 1–3.

scribed below) to the integration algorithm, and for each scheme, ran an experiment and assessed the performance of the integration. We show the weight assignments for each experiment in Table 1 and the performance results of the integrations in Figure 10.

Table 1: Weight assignments for experiments.

Exp	Scheme	W_P	W_B	W_G
1	Equal Weighting	.3333	.3333	.3333
1	Performance-Proportional	.3103	.4138	.2759
1	Performance-Differential	.3031	.4238	.2731
2	Equal Weighting	.3333	.3333	.3333
3	Equal Weighting	.3333	.3333	.3333
3	Performance-Proportional	.28	.38	.34
3	2-Parm-Search (Scheme A)	.1894	.4688	.3418
3	2-Parm-Search (Scheme B)	.1138	.5400	.3462
3	Adhoc (Scheme C)	.1000	.5000	.4000
3	Majority Voting (Scheme D)	n/a	n/a	n/a

Weighting Schemes The weighting schemes listed in Table 1 and Figure 10 are:

- **Equal Weighting:** Simply assign the weights $W_P = W_B = W_G = \frac{1}{3}$.
- **Performance-Proportional Weighting:** Assign the weights $W_i = \text{norm}(A_i)$, for $i = P, B, G$, where A_i are the accuracy estimates for i -th intersection algorithm.
- **Performance-Differential Weighting:** Assign the weights $W_i = \text{sum}(p_{i,j}^{\text{correct}} - p_{i,j}^{\text{next-highest}}) / n$, for $i = P, B, G$, for $j = 1..n$, where p^{correct} and $p^{\text{next-highest}}$ are the probabilities of the correct and next highest selection, n is the number of trials used to estimate the performance of algorithm i . Negative differences can be assigned zero.
- **Two-Parameter-Search Weighting:** Start with an estimate for one of the weights (assume weight W_1). Next, use performance estimates of the algorithms to compute the ratio $r = \frac{A_1 - A_2}{A_1 - A_3}$, where A_1, A_2 , and A_3 are the respective algorithm accuracies. Compute the weights by solving the two equations: $W_1 + W_2 + W_3 = 1$ and $W_1 - W_2 = r(W_1 - W_3)$.
- **Adhoc Weighting:** We assigned the weights by hand, based on observed trends in the performance of the other weighting schemes.

- **Majority Voting:** Select the candidate object that has the majority of the votes across the intersection algorithms. Empirical data, when using the equal weighting scheme, shows: (a) if all three vote for one object, the integration always votes for that object; and (b) if two of the three vote for the same object, the integration algorithm selects the same object a majority of the time.

Results and Discussion We again analyzed the accuracy performance with a one-way ANOVA. We found a strong effect of algorithm/integration scheme for Experiment 1 ($F(3, 116) = 4.37$, $p = .006$, $\omega^2 = 7.8\%$) and Experiment 2 ($F(1, 382) = 10.8$, $p = .001$, $\omega^2 = 2.5\%$). In addition, we found an effect for Experiment 3 ($F(6, 809) = 2.17$, $p = .044$, $\omega^2 = .85\%$), and while this effect is significant, it is a substantially weaker effect than the others reported in this paper. The error bars indicate that in Experiment 1, performance breaks down into two groups: (1) the Barycentric intersection algorithm and the Performance-Proportional integration scheme, and (2) the Equal Weighted and Performance-Differential schemes, with group (1) performing better than group (2). In Experiment 2, the Barycentric algorithm performed better than the Equal Weighted scheme. We only tested one integration scheme in this experiment because we decided it would be more interesting to look at test cases where the best intersection algorithm (Barycentric) had a lower global effectiveness, which motivated the next experiment. In Experiment 3, the Equal Weighted and Performance-Proportional schemes appear worse than the others, but otherwise the performance of all the schemes (and the Barycentric algorithm) is comparable. This relative equality is why the effect significance and size is substantially smaller for Experiment 3.

4.3 Evaluation of Disambiguation

The overall goal of our research was to develop methods for disambiguating multiple selections. We made strides towards this goal by developing several intersection algorithms and an integration algorithm that can combine the algorithms in many ways. The combination can be automatic, by using one of the schemes described previously, or it can be manual, by allowing the user to explicitly set the weights in the integration algorithm—for example, to only use the BARYCENTRIC-PIXEL-COUNT algorithm, the weights are set to $W_B = 1$ and $W_P = W_G = 0$. However, we have not yet shown that the integration algorithm is effective at disambiguating multiple selections, and we will address that now.

We conducted Experiment 4 to demonstrate, and evaluate empirically, the effectiveness of the integration algorithm in a simple disambiguation scenario. The experiment was performed using two subjects. We developed a dataset that requires a higher degree of precision for selecting by pointing than was necessary in the previous experiments. The dataset consists of a wall with nine windows; we placed a large window in the center and arranged eight smaller rim windows surrounding the center window. We asked subjects to select each of the nine windows separately, and the order of selection was random.

We collected statistics for 18 selections of the middle window and 31 selections of the outer windows. We analyzed the data considering how different selection strategies could be used for the scenario. Specifically, we considered how the windows could be selected using four strategies. The first strategy is to simply point at the window—“select object by pointing at it.” The second strategy is to select the largest window, no matter where the user is pointing—“select largest object.” The third and fourth strategies are combinations of the previous two, called “select largest object while pointing at it” and “select largest object while not pointing at it.” We analyzed the data considering these four selection strategies and show the results in Figure 11.

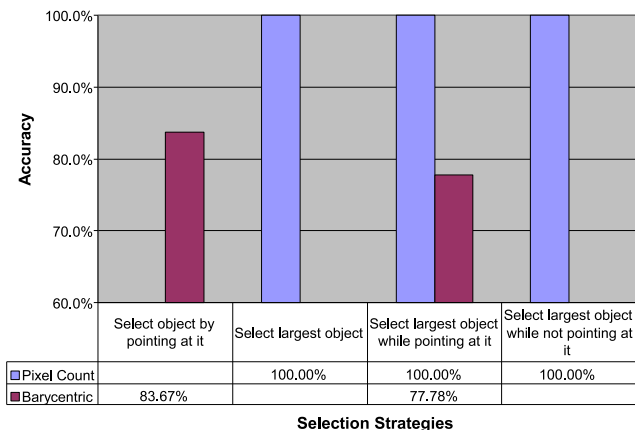


Figure 11: Comparison between the PIXEL-COUNT and BARYCENTRIC-PIXEL-COUNT algorithms over a range of user selection strategies. The dataset contains small, distant objects that are very difficult to select by pointing. The noise in pointing is high enough that BARYCENTRIC-PIXEL-COUNT fails often.

The PIXEL-COUNT algorithm worked perfectly for selecting the largest object across all the selection cases. The most interesting result is the overlapping case, “select largest object while pointing at it,” where the results show the success of the PIXEL-COUNT algorithm in comparison to the weaker performing BARYCENTRIC-PIXEL-COUNT algorithm. This result seems to indicate that as the precision required increases, the BARYCENTRIC-PIXEL-COUNT algorithm will not be as reliable as the PIXEL-COUNT for selecting the largest item.

The difference in the performances between the “select largest object while pointing at it” and the “select object by pointing at it” strategies indicate that the surrounding objects may be selected correctly more often than the middle object. This phenomenon may occur because the rim objects have free space on some of their sides. If this is the case, the middle object has a slight disadvantage using the “select object by pointing at it” strategy.

5 CONCLUSIONS AND FUTURE WORK

We presented three 3D pointing-based object selection techniques, PIXEL-COUNT, GAUSSIAN-PIXEL-COUNT, and BARYCENTRIC-PIXEL-COUNT, and applied them to the case of selection using a hierarchically-organized object database. We empirically demonstrated the general effectiveness of the BARYCENTRIC-PIXEL-COUNT technique. However, there are cases where that technique fails, so we developed an integration algorithm to try to leverage the strengths of each technique by combining their results using one of many weighting schemes. We evaluated these schemes and presented a careful analysis of the results. The bottom line is that different selection schemes work best in different scenarios, and the selection integration algorithm can disambiguate multiple selections.

There are several ways to improve this work. First, we need to take advantage of the semantic information in our database, for example, if it seems the user is selecting a window, it could be that the user is actually trying to select an object behind that window. Second, we can involve the user in the selection process beyond simple pointing. Perhaps the user could manipulate a dial or similar controller to scroll through the multiple selections until the correct object is selected. Third, we could determine the best sets of weights for certain common types of databases, or even for different areas within a single database, and establish “weight profiles” to be employed for those databases or areas of databases. A modification

of that idea is to classify weight assignments by the situations in which they work best, and then use that information for developing better performing automated selection techniques that adjust to the situation on-the-fly by swapping weight assignments. Another possibility is to have the algorithms learn from user indications as to whether the correct item was chosen or not. An implementation issue to address is how to properly handle ties in probabilities. For future studies, we would also like to add more subjects in the experiments.

6 ACKNOWLEDGEMENTS

We thank Matt Wesson, Dave McGee, Phil Cohen, Dennis Lin, Eric Burns, John Park, Elliott Cooper-Balis, Derek Overby, Aaron Bryden, Jason Jerald, Brian Hurley, Joshua Eliason, Jesus Arango, Eric Klein, Doug Maxwell, Simon Julier, Patrick Violante, and Mark Livingston for their contributions. This research was sponsored by the Office of Naval Research under grants #N00014-04-WR-2-0049, #N00014-02-1-0038 and #N00014-04-WX-2-0053 and by the Naval Research Laboratory Base Program grant #N00014-04-WX-3-0005.

REFERENCES

- [1] D.A. Bowman. *Interaction Techniques for Common Tasks in Immersive Virtual Environments: Design, Evaluation, and Application*. PhD thesis, Georgia Institute of Technology, Atlanta, Georgia, 1999.
- [2] P. Cohen, M. Johnston, D. McGee, S. Oviatt, et al. Quickset: multimodal interaction for distributed applications. In *Proc. of Intl. Multimedia Conference*, pages 31–40, New York, NY, 1997.
- [3] A. Forsberg, K. Herndon, and R. Zeleznik. Aperture based selection for immersive virtual environments. In *Proceedings of User Interface Software and Technology 1995*, pages 95–96, 1995.
- [4] K. Hinckley, R. Pausch, J.C. Goble, and N.F. Kassell. A survey of design issues in spatial input. In *Proceedings of User Interface Software and Technology 1994*, pages 213–222, 1994.
- [5] S. Julier, Y. Baillot, M. Lanzagorta, D. Brown, and L. Rosenblum. Bars: Battlefield augmented reality system. In *NATO Symp. on Information Processing Techniques for Military Systems*, October 2000.
- [6] E. Kaiser, A. Olwal, D. McGee, H. Benko, A. Corradini, X. Li, S. Feiner, and P. R. Cohen. Mutual disambiguation of 3d multimodal interaction in augmented and virtual reality. In *Proc. Intl. Conference on Multimodal Interaction-Perceptual User Interfaces*, Vancouver, BC, Canada, November 2003.
- [7] M. Kolsch, M. Turk, and T. Hllerer. Vision-based interfaces for mobility. In *Proc. of Intl. Conference on Mobile and Ubiquitous Systems*, 2004.
- [8] J. Liang and M. Green. Jdcad: A highly interactive 3d modeling system. *Computers and Graphics*, 18(4):499–506, 1994.
- [9] M. Mine. Virtual environment interaction techniques. Technical Report UNC Chancel Hill Computer Science Technical Report TR95-018, 1995.
- [10] A. Olwal, H. Benko, and S. Feiner. Senseshapes: Using statistical geometry for object selection in a multimodal augmented reality system. In *Proc. IEEE Intl. Symposium on Mixed and Augmented Reality (ISMAR'03)*, pages 300–301, Tokyo, Japan, October 2003.
- [11] J.S. Pierce, A. Forsberg, M.J. Conway, S. Hong, R. Zeleznik, and M.R. Mine. Image plane interaction techniques in 3d immersive environments. In *Proceedings of 1997 Symposium on Interactive 3D Graphics*, pages 39–43, 1997.
- [12] C.A. Wingrave, D.A. Bowman, and N. Ramakrishnan. Towards preferences in virtual environment interfaces. In *Proc. of Eighth Eurographics Workshop on Virtual Environments*, pages 63–72, 2002.
- [13] C.A. Wingrave, R. Tintner, B.N. Walker, D.A. Bowman, and L.F. Hodges. Exploring individual differences in raybased selection: Strategies and traits. In *Proc. of IEEE Virtual Reality Conference*, pages 163–170, 2005.

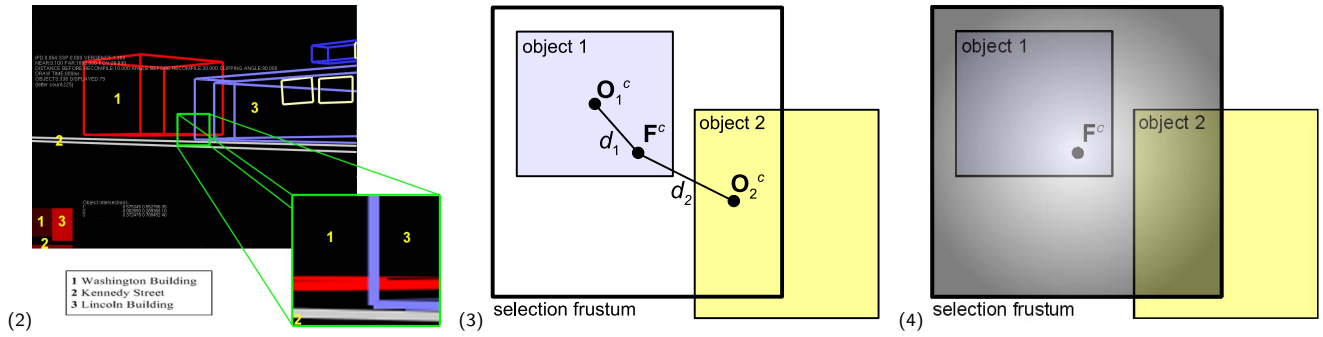


Figure 2: Operation of PIXEL-COUNT algorithm. The scene is rendered into the small selection frustum shown in the center of the image.

Figure 3: Operation of BARYCENTRIC-PIXEL-COUNT algorithm. Because $d_1 < d_2$, each pixel of object 1 will be weighted more heavily than the pixels of object 2.

Figure 4: Operation of GAUSSIAN-PIXEL-COUNT algorithm. Pixels in objects 1 and 2 are weighted by a circularly symmetric Gaussian function centered at F^c .

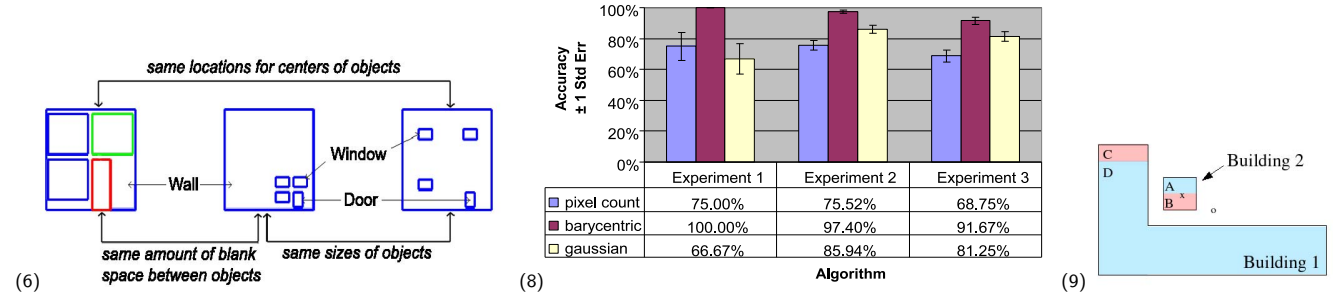


Figure 6: Three test cases used to determine the degree of precision needed by the experiments. We varied the size of the windows and spacing between each, and asked several users to perform selections of the windows for each test case. Statistics were collected to determine the precision required for our experiments.

Figure 8: The accuracy (mean percentage of correct selections) given by the three intersection algorithms in Experiments 1–3.

Figure 9: The effect of BARYCENTRIC-PIXEL-COUNT on selection involving a concave shape. The blue regions work properly, while the red fail. Empirical results show regions B and C fail 85% of the time.

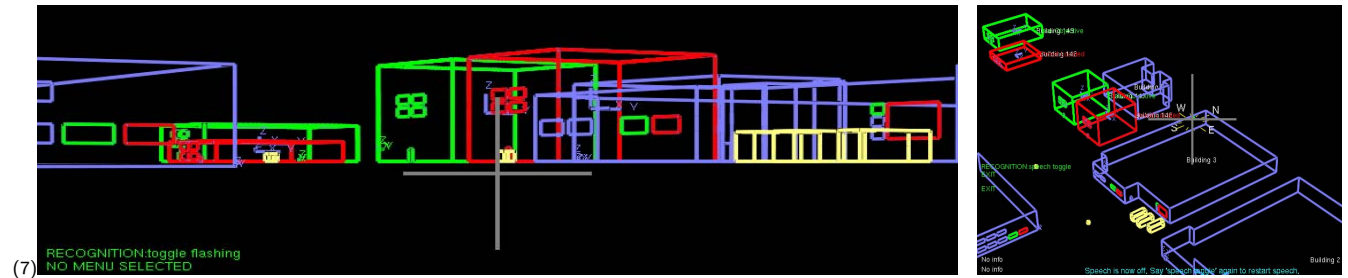


Figure 7: The dataset used in Experiments 1–3. Progressively smaller subsets of the dataset, with different test cases, were used in Experiments 1 and 2. (Left) Ground-level views; these are the views that subjects saw during the experiments. (Right) Elevated view; given to provide the reader a feel for the dataset's layout.

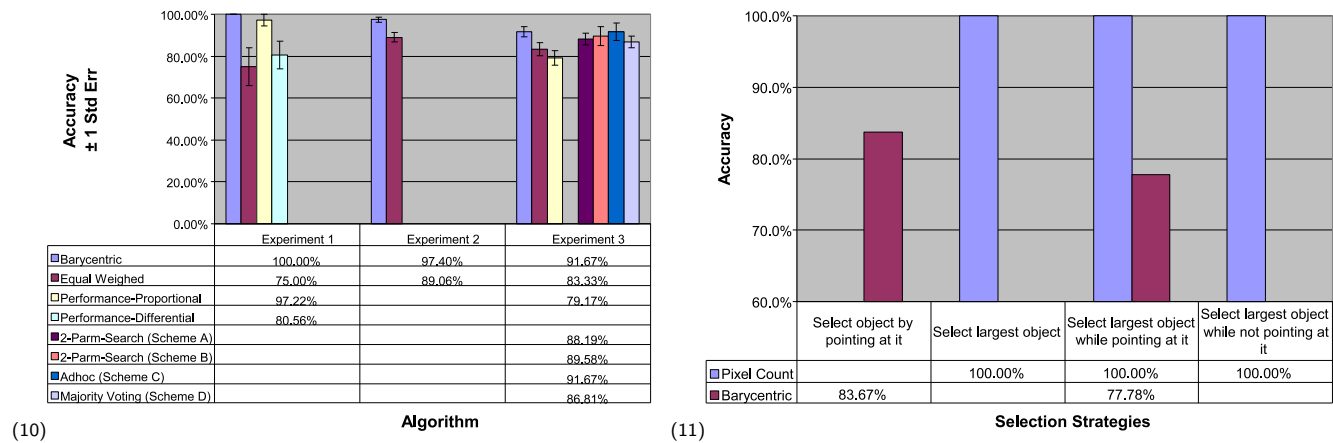


Figure 10: The accuracy of the integration schemes versus the BARYCENTRIC-PIXEL-COUNT intersection algorithm in Experiments 1–3.

Figure 11: Comparison between the PIXEL-COUNT and BARYCENTRIC-PIXEL-COUNT algorithms over a range of user selection strategies. The dataset contains small, distant objects that are very difficult to select by pointing. The noise in pointing is high enough that BARYCENTRIC-PIXEL-COUNT fails often.