

Virtual 3D Jigsaw Puzzles: Studying the Effect of Exploring Spatial Relations with Implicit Guidance

Felix Ritter¹, Bettina Berendt², Berit Fischer², Robert Richter², and Bernhard Preim³

¹O-v-G University of Magdeburg,  / ²Humboldt University of Berlin / ³MeVis gGmbH

Abstract

This paper investigates the engaging concept of virtual 3D jigsaw puzzles to foster the understanding of spatial relations within technical or biological systems by means of virtual models. Employing an application in anatomy education, it answers the question: How does guided spatial exploration, arising while composing a 3D jigsaw, affect the acquisition of spatial-functional understanding in virtual learning environments (VLE)? In this study, 16 physiotherapy students were interviewed before and immediately after using either a virtual 3D jigsaw puzzle enabled VLE or a simplified version without the interaction specific to the 3D jigsaw concept. Results indicate that students using the jigsaw-enabled VLE achieved a significant better understanding of the spatial and functional correlations illustrated by the model. These findings suggest that the concept of a 3D jigsaw puzzle, with its implicit guidance, facilitates and advances learner's understanding of spatial correlations and related functionality.

1 Introduction

In many different areas, learning involves the understanding of complex spatial phenomena. In engineering, the construction of machines has to be mastered as a prerequisite for maintenance. To replace a part of a complex engine, some of its parts have to be decomposed in a well-defined sequence. Medical students must imagine the spatial and functional relations within the human body to master anatomy. Since the human body is probably the most complex system known to mankind, understanding spatial relations between anatomical structures causes considerable difficulties. With interactive 3D computer graphics, based on computerized 3D models, these spatial relations can be explored.

Existing medical education software providing 3D viewing capabilities can be subdivided, on the one hand, into systems imparting factual knowledge in a way inspired by a 3D atlas, e.g. VOXEL-MAN (Höhne et al., 1996) and DIGITAL ANATOMIST (Brinkley et al., 1999). They allow to explore the voxel or polygon data of a virtual model and offer textual information. Models can be viewed from arbitrary positions and parts can be hidden. However, such systems rely solely on visualization, whereas people are used to touch and directly interact with the objects they want to explore.

On the other hand, specialized training systems, such as the one presented in (Sourin et al., 2000) and (Poston et al., 1996) for virtual surgery, foster the development of physical skills (Romiszowski, 1999). They provide interactive manipulations to a certain extent, which, however, are mostly restricted to special tasks required for this kind of intervention.

Understandings are learned differently from the way factual knowledge and skills are learned (Perkins and Unger, 1999). None of these systems offers guidance with the exploration of spatial relations between structures of the model nor do they comprise an engaging and reflective concept. In (Ritter et al., 2000) a method for actively investigating 3D models based on the concept of 3D jigsaw puzzles was proposed. 3D jigsaw puzzles provide a familiar and motivating concept for inter-

action with complex 3D models. In such a puzzle, a set of elementary objects is composed to form a specific model. The shape of these objects indicates which parts belong together. Guidance is introduced by restricting the composition to a subset of the model consisting of the structures and correlations to learn. By intensively manipulating the objects in order to connect them to an already composed subset, the user is expected to gain a deeper understanding of these structures or, with other words, to deploy already obtained knowledge with understanding.

This study explores the effectiveness of the virtual 3D jigsaw approach in the context of anatomical education. A prototypical implementation of a virtual 3D jigsaw is compared with a simplified version of the same application offering none of the interaction specific to the 3D jigsaw concept.

2 Implicit Guidance

Guidance is given implicitly by the interaction concept of the virtual 3D jigsaw. It requires the user to inspect the 3D model in detail at object level. Objects (structures in which the model has been subdivided) must be selected in order to place them upon an already composed subset of the model, missing objects must be identified. Moreover, the concept also requires the user to do so in a well-defined sequence. Inner objects must be assembled before outer objects can be connected. Thus, the user must imagine which structures lay in front of other structures when seen from a certain direction. Further implicit guidance is introduced by restricting the composition to a subset of the model consisting of the structures and correlations for which an understanding should develop. By manipulating these objects to connect them to the given, already assembled subset, we *hypothesize* the user to gain an understanding particularly about these structures and the surroundings.

3 The Tested Virtual Learning Environments (VLE)

The most critical issue for the evaluation was the construction of an appropriate comparison condition. A VLE based on the virtual 3D jigsaw concept had to be compared with a second VLE similar in every major respect, except for those characteristic for the jigsaw approach. To meet this requirement, we employed two informationally equivalent learning environments differing only in configuration of the interactive features of the very same software (called SMARTJIGSAW).

In the following, techniques relevant in the context of this study will be discussed briefly. For a detailed discussion of the interaction and visualization tasks to be fulfilled by an interactive system based on the virtual 3D jigsaw approach see (Ritter et al., 2000).

3.1 The Jigsaw-Enabled VLE

The design has been guided by the concept of 3D jigsaw puzzles, but differs in some major respect from real jigsaws. It is intended to support learning for understanding (Perkins and Unger, 1999) rather than just providing entertainment. It offers additional possibilities in that the computer “knows” how the 3D model should be assembled. This is used to provide guidance to the user.

Figure 1 depicts the screen of SMARTJIGSAW in a typical learning session. The system provides a set of *views* where the user can place and group objects with drag-and-drop operations. Besides the main view where users compose the model, a detailed view shows the currently selected object. The system rotates it automatically to facilitate the perception of the shape. A final view shows the composed model like the photos on the packages of real 3D jigsaws, which help users find the right place for the puzzle pieces. Furthermore, an exploded view inspired by technical illustrations shows an exploded model revealing covered structures in a composed subset.

Because the 3D puzzle requires precise interaction in 3D, depth cues play an important role (Wanger et al., 1992). Each 3D view displays a light ground plane that is scaled such that all ob-

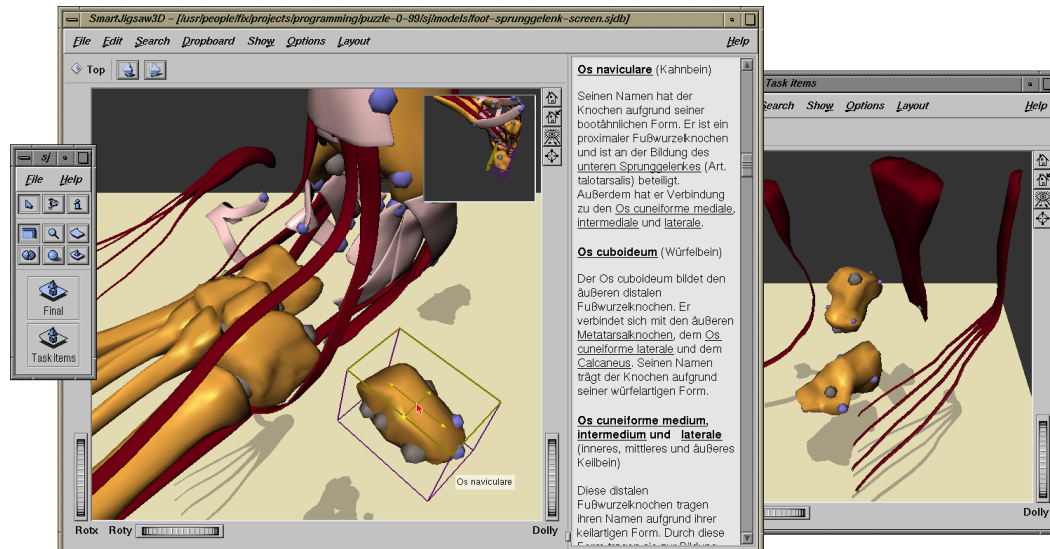


Figure 1. The virtual 3D jigsaw; the main view shows the so-far composed muscles, ligaments, and bones. The yet to be connected items are placed randomly on another view which is partly occluded on the right. The small panel on the left side lists all the 3D views. Additionally, the system offers textual information as seen at the right of the main view. In this example, the student has disconnected some occluding items prior to moving the 'Os naviculare' to the appropriate place.

jects cast a shadow on it. Motion parallax can be used most efficiently if the user has direct control over this effect (Hubona et al., 1997). Therefore, the system supports a 3D mouse in addition to the common pointing device (2D mouse), which is used exclusively to rotate the virtual camera around a user-defined point of interest (POI) and to control the distance to the POI.

Objects, such as bones and muscles in anatomy, become highlighted and labeled when the user touches them with the pointing device. Labels are placed nearby to facilitate memorizing of names (Moreno and Mayer, 1999). Selecting an object yields an explanation of the object in the hypertext area of the main window. The structure of these explanations is inspired by maintenance manuals and anatomical atlases. Upon activation of a link, an integrated camera engine animates the movement of the current point of view smoothly to a close-up of the object cited.

Objects must be transformed precisely and finally be connected to solve a jigsaw. These tasks are supported by a snapping mechanism and by using the user's bimanual skills. The geometric model, which is to be prepared by an author before, has a number of docking points that indicate connections between parts. Objects are composed correctly if the docking points touch each other. Shape and color of the docking points give hints as to which objects can be connected. For further support, the user can point at a docking point yielding the corresponding object to highlight. Objects snap together if the distance between two docking points is below a given threshold, regardless of correctness. Once an object is attached, the same algorithm prevents the user from detaching it inadvertently. With a quick movement, however, separation is possible. Collision detection prevents objects from being moved through others. Thus, users are required to decompose objects from the model before placing an object inside. When objects collide they are highlighted for a moment to provide visual feedback. Objects also become highlighted when they are about to snap. A green flashing outline indicates correct connections, whereas a red flash signals a wrong combination.

The transformation of selected 3D objects is performed with a 3D widget that enables the user to translate the attached object with the 2D mouse. Rotation, even when constrained to angles of 45 degrees, had proven to be too difficult for composing complex parts. Since grabbing the backsides

of the 3D widget (recall Figure 1) for translation in the plane of that face requires the user to rotate the model, strong motion parallax is perceived while solving a puzzle. Users use the attached 3D mouse (preferably with the non-dominant hand) for the manipulation of the viewpoint and, in parallel, the pointing device to manipulate objects, a setup inspired by (Hinckley et al., 1994).

3.2 The Jigsaw-Disabled VLE

Recalling Figure 1, the jigsaw-disabled VLE comprises only the main view with the completely assembled model. The user may freely explore the model, for instance in the exploded view, and navigate through the accompanying hypertext but cannot detach objects. Indeed, the jigsaw-disabled version was configured to behave like an interactive 3D atlas (e.g. DIGITAL ANATOMIST).

4 Methodology

We are interested in the impact of the virtual 3D jigsaw concept on the students' understanding of spatial-functional relations by means of virtual models. As the human body is probably the most complex system to mankind, we chose to test students with the right knowledge in anatomy for their advances in understanding the spatial-functional relations between anatomical structures.

4.1 Experimental Design

Utilizing an independent group design, half of the students worked through several exercises employing the jigsaw-enabled VLE, whereas the other half used the VLE without jigsaw functionality. The dependent variable in this study was the score in an interview examining the spatial and functional understanding of spatial relations between anatomical structures. Beside measuring:

- *understanding of spatial and functional relations* (interview)

This interview was run before and immediately after the exercises and taped on video. Beside being asked well-defined questions by the interviewer, subjects also received the questions on paper as part of the accompanying questionnaires. We did not use a multiple choice test, because we felt it would not be sensitive enough. The understanding of spatial and functional relations is best assessed by an oral interview, where subjects may use gestures to describe shapes and suchlike. Furthermore, the video allowed the answers to be scored afterwards by two anatomy course instructors independently, thus maximizing objectivity.

Four additional measures were employed to refine the analysis and to gather survey information:

- *factual knowledge*

Obtained using the VLEs in a special mode, each subject was randomly presented the same ten questions about names of anatomical structures employing the 3D illustration. Half of them had to be answered by choosing the name from a list (three or four choices) and the other half by entering the name in a fixed period of time using the keyboard. Subjects got immediate results by coloring correctly named structures green and otherwise red. In the latter case, also the correct name was displayed.

- *spatial visualization skills*

A subset of an intelligence structure test (German I-S-T 2000) consisting of figurative classification, mental rotation, and figurative recognition was used. This subset measures the ability to establish relationships between both planar and spatial geometric objects as well as the ability to memorize and recognize figurative information in short term.

- *learning confidence*

Measured before and after the experiment as part of the questionnaire, it assesses the subjects' confidence in attaining and in having attained insight of the spatial and functional relations with the system or, with other words, having advanced their understanding.

- *usability and attitude*

By observing subject utilizing the systems and by logging interactions, we tried to gather information about usage patterns and interaction difficulties. Furthermore, a section of the questionnaire asked participants for their likes and dislikes.

4.2 Participants

A total of 16 physiotherapy students took part in the study. They were all first or second year students in an undergraduate program and had a basic knowledge of anatomy. Since studying physiotherapy does not comprise the dissection of cadavers, students had to rely almost exclusively on anatomical atlases, textbooks or videos to build up a spatial image of the human body.

None of the subjects had experiences with computer-based 3D atlases before. Students were equally assigned to both groups according to their course instructors' rating of their performance in anatomical courses. The subjects ranked in age from 23 to 26 and comprised 11 women and 5 men, whereas the group using the jigsaw-enabled VLE consisted of 6 women and 2 men.

4.3 Test Environment and Materials

Both VLEs were installed on two LINUX-PCs with Logitech Magellan 3D mice. Put up side by side, the experimenter had enough space to sit in the middle of the two subjects.

We chose the anatomical structures of the right human foot, particularly the ankle joint, as the topic for this study. Following the classification of an anatomy atlas, a 3D polygonal model consisting of 45k Δ was split into 53 objects (28 bones, 11 muscles, 14 ligaments). Nerves and blood vessels were omitted for the sake of simplicity. Each pair of objects shared a number of docking points represented by small, colored spheres of different scale. Three distinguishable colors were used to indicate the possible object combinations (bone-bone, bone-muscle, and bone-ligament).

A hyperlinked text explaining the objects and describing regions to which they belong (e.g. bones of the lower and upper parts of the ankle joint) accompanied the model. References to spatial and functional relations tested in this study had been carefully removed. If possible, the objects and textual information were linked bidirectionally. Thus subjects could receive additional information either in graphical or textual form for a selected item.

4.4 Exercises

The tests and learning exercises were developed in close collaboration with two anatomy course instructors, who also volunteered to interpret the taped interviews with the participants of the study. The exercises were designed to be as similar as possible for each group to avoid any effect on the dependent variable. They also had to match with the questions posed in the test to maximize learning gains. Both groups were given 30 minutes to work on the exercises.

4.4.1 Subjects Using the Jigsaw-Enabled VLE

Subjects using the jigsaw version had to solve three puzzles:

1. to compose two bones of the ankle joint upon the otherwise complete foot model (5 min.);
2. to attach the three bones of the second toe including metatarsus to the model (10 min.);
3. to compose three bones of the ankle joint together with three muscles of this region upon an already assembled subset of the model (15 min.).

The partially assembled model of the foot was displayed in the main view. All missing objects, being part of the task, resided on a separate view (see Figure 1). Subjects were required to drop objects on the main view and to move the objects to their appropriate place.

4.4.2 Subjects Using the Jigsaw-Disabled VLE

The control group was given two different models of the foot, because they could not decompose objects from the model. They also received additional encouragement to look at the model from different sides, which this group was less often required than the other group that was guided implicitly by the jigsaws' interaction design. The first task (10 min.) using only the skeleton of the foot was divided into three subtasks:

1. Turn the foot and look at the bones from different angles. Can you make sense out of the arrangement?
2. Look at the bones of the second toe including metatarsus. What have all toes in common and what differentiates the large toe?
3. Look for the bones of the ankle joint and consider their interrelations with other bones.

A second task (20 min.) employed the full model of the foot with all muscles and ligaments. Subjects received two additional hints:

4. Study the bones and muscles of the ankle joint.
5. Look again at the whole foot and try to understand the interaction between muscles and bones.

4.5 Interview and Questionnaires

Two questionnaires were designed, one to obtain pre-test knowledge and another to gather results immediately after the subjects used the VLEs. In the first question the subjects were asked to rate their previous experiences with interactive 3D illustrations, such as interactive anatomical atlases, on a six point scale, where 6 indicated regular use and 1 no experiences. This question was only asked in the pre-test questionnaire.

The second section had been included to assess the participants' subjective rating of confidence that the system facilitates the learning of spatial relations. In case of the pre-test, subjects were already introduced to the system and had a short glimpse on the jigsaw-disabled VLE. Confidence was also ranked on a six point scale (very high to no gain).

The primary section contained six questions carefully chosen and formulated by an anatomy course instructor to reflect spatial-functional knowledge about the human foot's anatomy. Matching with the exercises, the questions were related to distinct structures of the ankle joint (tarsus) and toes. The fourth question, for example, was as follows:

Explain the interaction of bones and muscles with the dorsiflexion (upward motion) of the foot. Which relationship exists between the movement axle of the joint and the muscles?

In contrast to the other sections of this questionnaire these six questions were asked and answered in an *oral interview*. A marking scheme with key terms for the analysis of the answers had been prepared together with the questions. The last section, only on the post-test questionnaire, asked the subjects to comment on the systems and to express their likes and dislikes.

4.6 Procedure

In the beginning all participants were tested for their spatial visualization skills using the I-S-T 2000. Having finished this test, they were introduced to the studies' goal on an informal basis by the experimenter. As part of this introduction, subjects were using the VLEs in a mode implemented to test factual knowledge (questions about names of anatomical structures). In those 5 minutes, subjects got a short glimpse of the VLEs and the 3D illustration without getting to know any differences between both learning environments.

After the short quiz, one subject at a time was led to a separate room where he or she was given the pre-test questionnaire from the interviewer. After rating his or her confidence, the interviewer

turned on the camera and started with the oral interview about the spatial and functional relations between certain anatomical structures of the human foot. Each question was read to the subject. If an answer was unclear, the interviewer tried to clarify it. The estimated 30 minutes maximum for the interview were not exceeded.

The experiment started as soon as one subject of each group had finished the pre-test. We decided for this setup, because in that way the experimenter could observe two subjects by sitting between them. Since both subjects worked on exercises similar in contents but different in process, independent results have been ensured. The two subjects were first given an introduction to the systems using a 3D model of the right human knee. After the introduction of approximately 12 minutes and answering of queries, the subjects completed a 7-minutes' training session with a model of the human knee. Thus being comfortable with how to operate the systems, the subjects were given the set of exercises. They were told to finish all tasks within time. During the exercises subjects could ask the experimenter for advice on how to use a certain feature of the system.

When they had finished, the first subject was interviewed again, this time from a different interviewer. We used two interviewers alternately to avoid remarks like: Didn't I tell you this already? Before asking the same six question, subjects rated their confidence about having attained insights.

4.7 Statistical Analysis

Given the nature of the data we had collected (independent samples, non-parametric data) we decided to use a Mann-Whitney Test to analyze the results. Because it was hypothesized that the gain in understanding would show a significant higher increase for the jigsaw group, the results were analyzed for a one-tailed test. A two-tailed Kendall rank correlation was employed to disclose dependences between spatial visualization skills and learning gain. We report medians (M), interquartile range (IQR), means (X), and standard deviations (s) for descriptive purposes.

5 Results and Discussion

It took the subjects an avg. time of 2.2h to complete the session. They had a break after completing the pre-test, before they were introduced to the system and started working on the exercises.

5.1 Using the Virtual Learning Environments

Although the systems provided exactly the same information in terms of 3D model and hypertext, the usage was quite different.

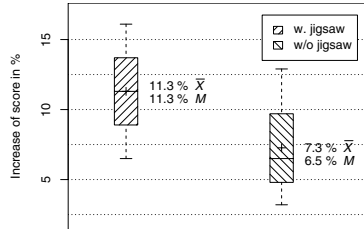
5.1.1 Jigsaw-Enabled VLE

As expected, subjects constantly interacted with the 3D model, looking at it from all angles. Users did decompose more objects than required to insert inside objects. Even when finished with the task they started to remove objects to explore covered structures. A structure's description was accessed by pointing at the 3D object rather than by scrolling the hypertext.

Analyzing the software logs of the third exercise revealed that this puzzle took the students an avg. of 9min 37sec ($s=3$ min 22sec) to compose the ankle joint. All subjects read the description of the six objects to be assembled. They changed the view on the model an avg. of 60 times ($s=39$) whereby an avg. of 15 rotations ($s=9$) with the 3D mouse were performed whilst composing the 'Talus' (a bone). Most rotations changed the view only by a small angle.

5.1.2 Jigsaw-Disabled VLE

Subjects spent most of the time exploring the model in the exploded view, which they used almost exclusively. Users of this group also utilized the hypertext more frequently, reading through the



	w. jigsaw	w/o jigsaw
Median	11.3%	6.5%
Mean	11.3%	7.3%
Interquartile Range	13.7% - 8.9%	8.9% - 5.2%
Range	16.1% - 6.5%	12.9% - 3.2%

Figure 2. Increase of overall score for the interview of spatial-functional understanding in percent by group. Given are median precision, interquartile range, mean (+), and extremes for both test groups.

text and navigating to the objects of the 3D model with the hyperlinks, an interaction enabled by the integrated camera engine. While working with the complete foot model for an avg. of 19min 32sec ($s=2\text{min } 10\text{sec}$) subjects changed the view an avg. of 21 times ($s=11$). Almost all rotations changed the view considerably.

5.2 Advances in Understanding Spatial Relations (learning gain)

Learning gain was measured as difference between the scores of the interviews done before and after the exercises. The scores were obtained using a marking scheme for the video-taped interviews. Two anatomy course instructors scored the answers of all subjects independently.

Analyzing the results of both groups, we found that the learning gain differed significantly. Subjects who used the jigsaw-enabled VLE had a higher gain than subjects using the other system without jigsaw-specific interaction (Mann-Whitney $U=8$, $p=0.037$). Accuracy overall score increase can be obtained from the Table and Figure 2.

As depicted in Figure 3, for almost every question the average gain of the jigsaw group is higher except for question 5. It turned out that question 5, although asking for spatial relations, was in fact assessing knowledge that could be obtained from the hypertext. This result was consistent with our observation that subjects using the non-puzzle version did much more work with the hypertext.

5.3 Factual Knowledge

A comparison of factual knowledge about anatomical names of certain structures obtained in the beginning revealed no significant differences between both groups. Subject assigned to the group using the jigsaw-enabled VLE did slightly worse ($X_w=64.6\%$, $s_w=17.2\%$) than the other group ($X_{w/o}=70.4\%$, $s_{w/o}=14.7\%$) when choosing the names from a list. When asked to enter the names with the keyboard the differences were roughly the same; ($X_w=62.3\%$, $s_w=15.8\%$) compared to ($X_{w/o}=66.1\%$, $s_{w/o}=11.9\%$). No significant correlation between results in the factual knowledge quiz and advances in understanding spatial relations was found (Kendall $\tau=-0.12$, $p=0.76$).

5.4 Spatial Visualization Skills

The analysis of the test papers did not show significant differences between both groups ($F(1,14)=1.21$, $p=0.86$). There was no correlation between the results in mental rotation and learning increase (Kendall $\tau=0.04$, $p=0.85$) and also none between figurative recognition and learning increase (Kendall $\tau=-0.32$, $p=0.15$). There was, however, significant correlation between figurative classification and learning increase (Kendall $\tau=0.71$, $p=0.0004$). Since the task consisted of decomposed, planar, geometric objects to be matched with a set of composed objects, the result might suggest that users performing well in this part of the test will also have advantages in learning spatial relations in 3D.

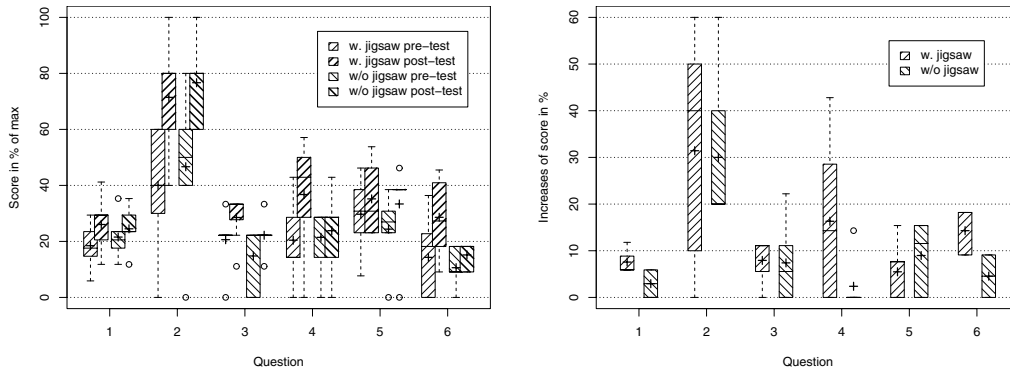


Figure 3. Absolute scores and increase of score for the interview of spatial-functional understanding in percent by question number. Given are median precision, interquartile range, mean (+), extremes and outliers (•) for both test groups.

5.5 Subjective Rating

We collected direct feedback from the subjects with the questionnaires.

5.5.1 Learning Confidence

Comparing the subjects' confidence in attaining and in having attained insight of the spatial and functional relations on a six point scale revealed a small decrease ($M_{pre}=5$, $s_{pre}=0$, $M_{post}=4.71$, $s_{post}=0.49$). There was no significant difference between both groups. As discussed in (Draper et al., 1996), this small drop in confidence is most likely due to realizing that they had still more to learn to master the topic. This assumption is supported by the fact that the subjects did only obtain an avg. of 51.8% of the maximum score in the spatial-functional knowledge pre-test and an avg. 61.1% in the post-test. However, the rating showed that subjects indeed liked the systems.

5.5.2 Attitude

Participants were quite enthusiastic about the overall system, but could not relate it to any other system because of missing experiences with 3D anatomy atlases. Subjects gave the learning content a high rating. All subjects liked the control of the view with the 3D mouse, although they had to be encouraged to utilize it during the initial training session. After only a short time, they were able to benefit from the 3D input device and used it in parallel with the 2D pointing device.

Subjects using the jigsaw-enabled VLE enjoyed the idea of using a jigsaw puzzle to learn anatomy. Comments we got from the questionnaire were:

"...foot easier to imagine by turning and composing; solving a jigsaw puzzle is fun." or
"By composing the parts, the correlations of the structures can be made clear better."

However, most of the users also remarked that the interaction with the objects was quite difficult. They would like to move the object straight to the place where it belongs instead of moving the object in several steps. This, however, was not possible most of the time because of the employed 3D widget we had configured in that way on purpose. Snapping was considered to be essential. Some of the subjects also would like the system to be able to compose automatically a subset of the model in an animation and then do the same task themselves.

Subjects using the jigsaw-disabled VLE liked the bidirectional links between model and hypertext. All participants enjoyed the name quiz we employed to familiarize them with the 3D illustration and to assess factual knowledge. They found the coloring of the items following their answer highly motivating, trying to dye green as much structures as possible.

6 Conclusions

There are two principal conclusions concerning the impact of the virtual 3D puzzle concept. First, the results indicate that the virtual 3D jigsaw puzzle does indeed improve the understanding of spatial relations from 3D illustrations. Hence it is of interest for users who need to imagine and understand spatial relations and their functional dependencies. The implicit guidance of the 3D jigsaw concept requires the user to interact with well-defined objects illustrating a chosen topic. Furthermore, it directs the user's attention to correlations between objects. The puzzle task that is defined by an author has direct impact on the objects and structures on which the user focuses.

Second, the jigsaw puzzle provides a level of motivation for learning, which is hard to achieve with other concepts. Users enjoyed the session having fun whenever they succeeded in attaching an object correctly.

For educational or maintenance purposes a wealth of textual information, such as about objects and their meaning, and about possible complications of an intervention, are required. Since users must recognize the objects, students benefit from the 3D jigsaw provided they have already obtained at least a basic understanding of the topic to be studied.

Although the employed virtual learning environment has been evaluated using an example in anatomy education, there are other areas (e.g. mechanical engineering) where it can be used to enhance the learners' understanding of spatial relations within a virtual model and correlated functionality.

7 References

- Brinkley, J. F., Wong, B. A., Hinshaw, K. P., and Rosse, C. (1999). Design of an Anatomy Information System. *IEEE Computer Graphics & Applications*, 19(3):38–48.
- Draper, S., Brown, M., Henderson, F., and McAteer, E. (1996). Integrative Evaluation: An Emerging Role for Classroom Studies of CAL. *Computers & Education*, 26(1–3):17–32.
- Höhne, K.-H., Pflesser, B., Pommert, A., Riemer, M., Schiemann, T., Schubert, R., and Tiede, U. (1996). A virtual body model for surgical education and rehearsal. *IEEE Computer*, 29(1):25–31.
- Hinckley, K., Pausch, R., Goble, J. C., and Kassell, N. F. (1994). Passive real-world interface props for neurosurgical visualization. In *Proc. of ACM CHI (Boston, April 1994)*, pages 452–458. ACM, New York.
- Hubona, G. S., Shirah, G. W., and Fout, D. G. (1997). The effects of motion and stereopsis on three-dimensional visualization. In *Proc. of the International Journal of Human-Computer Studies*, volume 47, pages 609–627.
- Moreno, R. and Mayer, R. E. (1999). Cognitive principles of multimedia learning. *Journal of Educational Psychology*, 91:358–368.
- Perkins, D. N. and Unger, C. (1999). Teaching and Learning for Understanding. In (Reigeluth, 1999), pages 91–114.
- Poston, T., Nowinski, W. L., Serra, L., Choon, C. B., Hern, N., and Pillay, P. K. (1996). The brain bench: Virtual stereotaxis for rapid neurosurgery planning and training. In *Proc. of Visualization in Biomedical Computing (Hamburg, September 1996)*, pages 491–500. Springer Verlag.
- Reigeluth, C. M., editor (1999). *Instructional-Design Theories and Models*, volume 2. Lawrence Erlbaum Associates, Publishers, Mahwah, NJ.
- Ritter, F., Preim, B., Deussen, O., and Strothotte, T. (2000). Using a 3D Puzzle as a Metaphor for Learning Spatial Relations. In Fels, S. S. and Poulin, P., editors, *Proc. of Graphics Interface (Montréal, May 2000)*, pages 171–178. Morgan Kaufmann Publishers, San Francisco.
- Romiszkowski, A. (1999). The Development of Physical Skills: Instruction in the Psychomotor Domain. In (Reigeluth, 1999), pages 457–481.
- Sourin, A., Sourina, O., and Sen, H. T. (2000). Virtual Orthopedic Surgery Training. *IEEE Computer Graphics & Applications*, 20(3):6–9.
- Wanger, L. R., Ferwerda, J. A., and Greenberg, D. P. (1992). Perceiving spatial relationships in computer-generated images. *IEEE Computer Graphics & Applications*, 12(3):44–58.