

Teasing Apart Self-Explanations: How the Types of Utterances Generated While Self-Explaining May Impact Learning from Text



Alex M. Tseng, Norielle Adricula, Diane P. Lam
University of California, Berkeley

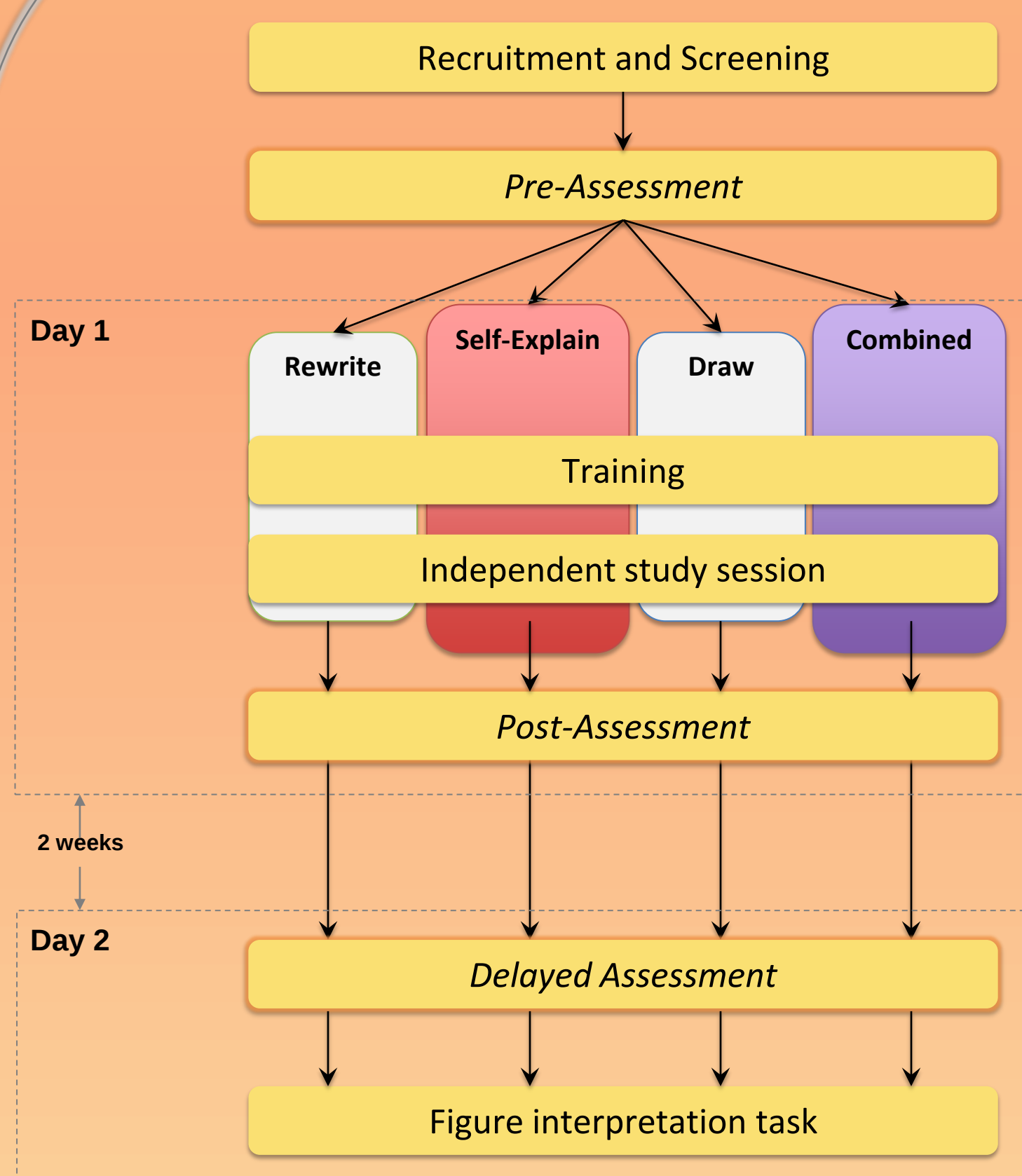
Abstract

The self-explanation (SE) effect describes the phenomenon where learners who explain ideas to themselves aloud after reading text, are more likely to learn the material than those who do not or do so less frequently. This study investigates the types of utterances that students produce while learning from a text about the cardiovascular system, and correlates the frequency of these different utterance types to their performance on free-response assessment questions aimed at measuring their understanding of various aspects of the cardiovascular system. Preliminary analyses reveal that self-explanations that are goal-driven or elaborative are predictive of performance on assessment questions about function; while incorrect statements are negatively correlated to performance on questions about structures. These findings inspire further analyses into the types of explanations that students may generate while learning text to improve learning.

Introduction

- The self-explanation effect describes how learners gain deeper understanding while reading text if they make inferences about the material aloud as they are reading (Chi, Bassok, Lewis, Reimann, & Glaser, 1989).
- Previous studies have used experimenter prompting in order to encourage self-explanation but this essentially confounds the act of *self*-explaining, as participants are likely to respond *to* the experimenter (Chi, Deleuw, Chiu, & Lavancher, 1994).
- In this study, we provided participants with a brief training session on how to self-explain and mimicked an independent study environment by having subjects study in a cubicle.
- We investigated the types of utterances subjects generate in this environment and whether or not they were predictive of different types of understanding.

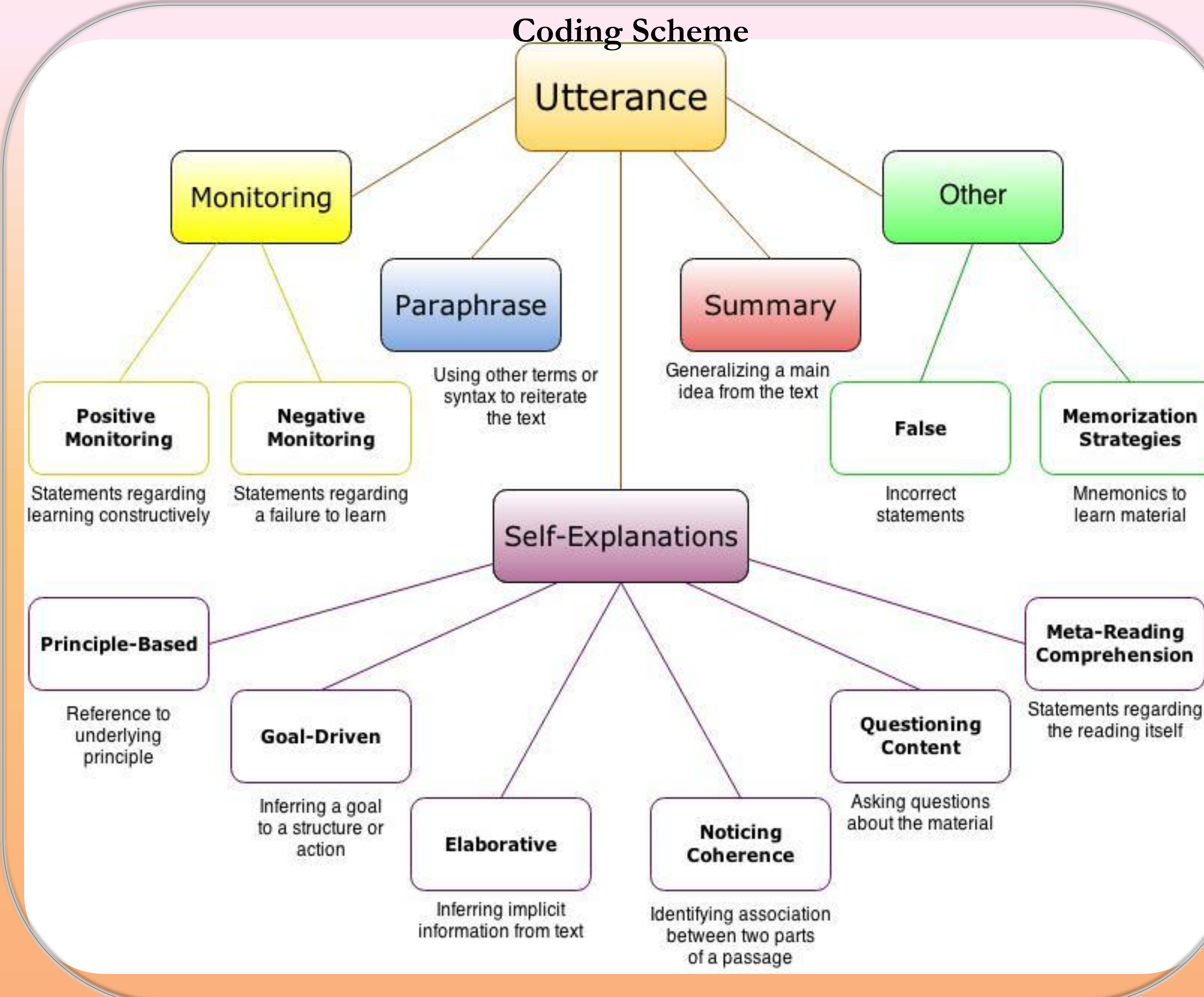
Methods



Guiding Questions for Self-Explaining

- 1. Main idea:** Do I understand the concept or main idea in the text? What kind of inferences can I make about these ideas?
- 2. Relate to other ideas:** How does this idea relate to other ideas in the text or from my prior knowledge?
- 3. My questions:** Does the self-explanation I have generated help to answer the questions that I am asking?

- This study is part of a larger research agenda that compares rewriting, SE, drawing, and a combination of SE and drawing.
- For the SE group, 25 participants completed a brief training session which emphasized three guiding questions for self-explaining.
- During the independent study phase, students read 11 passages describing the cardiovascular system and were asked to self-explain each passage.
- These self-explanations were video and audio recorded and transcribed
- Transcriptions were then broken down into utterances to be analyzed.
- The following coding scheme was applied to the utterances.



Findings

Goal	Gain (Function)
1	2
0	1
3	6
3	3
2	3
0	3
0	1
0	-1
1	0
4	3

Goal vs. Gain (F):
 $r^2 = 0.4351$
 $p = 0.0391$
 $\hat{y} = 0.863x + 0.892$
(where $\hat{y}$ is predicted gain in function and x is number of goal-oriented utterances)

Elab vs. Gain (F):
 $r^2 = 0.3795$
 $p = 0.0381$
 $\hat{y} = 0.669x + 1.297$
(where $\hat{y}$ is predicted gain in function and x is number of elaborative utterances)

Elab vs. Gain (P):
 $r^2 = 0.3157$
 $p = 0.0910$
 $\hat{y} = 0.399x + 0.622$
(where $\hat{y}$ is predicted gain in pathway and x is number of elaborative utterances)

Elab	Gain (Function)	Gain (Pathway)
0	2	0
0	1	2
3	6	2
1	3	0
3	3	3
0	3	1
0	1	2
0	-1	0
0	0	-1
5	3	2

Wrong	Gain (Structure)	Ret (Structure)
0	1	-1
0	0	2
0	3	0
0	4	-1
3	-2	4
2	3	-2
0	2	-4
0	2	-2
3	-1	5
1	2	1

Wrong vs. Gain (S):
 $r^2 = 0.3831$
 $p = 0.0564$
 $\hat{y} = -0.913x + 2.221$
(where $\hat{y}$ is predicted gain in structure and x is number of wrong utterances)

Wrong vs. Ret (S):
 $r^2 = 0.4620$
 $p = 0.0306$
 $\hat{y} = 1.490x - 1.141$
(where $\hat{y}$ is predicted retention in structure and x is number of wrong utterances)

Examples

	Goal-driven	Elaborative	Paraphrase
9199 Utterances (Passage 2)			
OK, so the septum is basically like a barrier that separates the heart into two hearts.		1	
Hm, I'm guessing it's sort of like a tissue barrier.		1	
Okay so the right side brings the blood to the lungs,			1
the lungs give us oxygen so the right side is probably- carries the blood that doesn't have the oxygen and that's why it has to go to the lungs to get oxygen.	1		
The left side is probably the side with the oxygen so, that's why it's bringing it to the other parts of the body.	1		
Totals	2	2	1

9199	Gain Pathway	Gain Function	Total Gain
Assessment Scores	2	3	8

	Goal-driven	Elaborative	Paraphrase
9613 Utterances (Passage 3)			
that prevent the blood flowing in the wrong direction so they separate, like I said, the atrium from the ventricles and yeah,			1
so the valves separate the ventricles from larger vessels through which the blood flows,			1
and, as the blood is pumped out of the ventricles the valves consist of flaps of tissue so that goes through the flaps of tissue,			1
what's interesting is as the blood returns to the heart it has a high concentration of carbon dioxide,			1
Totals	0	0	4

9613	Gain Pathway	Gain Function	Total Gain
Assessment Scores	0	-1	1

Discussion

- Goal-Driven explanations:** There is a moderately strong, positive correlation between the number of utterances that are goal-driven and score for function-related assessment items. We believe this increase can be explained by how goal-oriented explanations tend to answer “why” and “how” questions, prompting a subject to infer the function of structures and how they might be important to the overall system.
- Elaborative Explanations:** There is also a moderately strong, positive correlation between the amount of utterances that are elaborative and the increase in score with regards to function and pathway. The explanation of this increase is likely similar to those related to goal-oriented explanations. With regards to elaboration, it is essentially a translation from the text to the student’s own words or context of understanding, thereby increasing the score in both pathway and function.
- Wrong Explanations:** There is a moderately strong, negative correlation between the number of wrong utterances and the improvement of score on the structure portion of the post test. This is possibly due to the tendency of wrong explanations to impede learning.
As for the retention score on structure, there is a moderately strong, positive correlation that is significant ($p < 0.05$). By the regression line, it is predicted that with each wrong explanation, retention score for structure increases by 1.490 points. This positive correlation is possibly explained by the tendency of wrong explanations to be consistently wrong. If a student does not score well in the pre-test and the explanations are incorrect, any coincidental increase in delayed testing will be statistically significant.

Conclusions and Future Directions

- The present work suggests that when studying from a text, students should:
 - think about why and how structures work within a system (goal-oriented explanations)
 - explain beyond the explicit text and translate it to one’s own understanding of the material (elaborative explanations)
- In the context of the larger study, the self-explain group had the most retention about function after 2 weeks, very likely explained by these strong correlations with gain for function shown here.
- Breaking up utterances using a coarse grain versus a fine grain could yield different insights and analyses, providing insight into using coding schemes for utterances.
- More work and analyses still remain but these preliminary results provide more insight for student learning from text and how to use the self-explanation strategy more effectively
- Based on these present findings, another component for coding will be added to the scheme to see if relationships can be found between utterance types, content about structure, pathway, or function, and gains/retention on structure, pathway, function questions.



This research is partially supported by the National Science Foundation Graduate Research Fellowship Program