

Iowa State University

From the Selected Works of Tammy Slater

2004

The evaluation of causal discourse and language as a resource for meaning

Bernard Mohan, *University of British Columbia*

Tammy Slater, *University of British Columbia*



Available at: https://works.bepress.com/tammy_slater/8/

The evaluation of causal discourse and language as a resource for meaning

Bernard Mohan and Tammy Slater

Systemic Functional Linguistics (SFL) has produced a volume of work that has made significant advances in the analysis of discourse, including discourse in education. Halliday's distinctive view of linguistics is concerned with texts in social contexts and is oriented to the description of language as a resource for meaning rather than with language as a system of rules (Halliday & Martin, 1993). SFL has major implications for the assessment of discourse, particularly in the important area of assessing learners of English as a Second Language in tertiary education, an area which test developers have a strong interest in improving. Halliday's rich perspective on discourse sets standards of adequacy for the assessment of discourse in general and of scientific discourse in particular. Do present language evaluation theories and practices meet these standards? To address this question, we examined the evaluation of causal explanations, an essential part of academic literacy, and found disturbing issues in both theory and practice which seriously question the adequacy of present models designed to assess competence in a second language.

We have chosen causal explanations for a variety of reasons. The topic of causal explanations in scientific discourse has been a continuing focus of SFL interest (Halliday & Martin, 1993; Martin & Veel, 1998), and its analysis illustrates features of the SFL approach in general. Causal discourse is a part of the metalanguage of science and an essential part of scientific literacy (Rowell, 1997). Causal explanations are not limited to science, however; since they occur across many academic disciplines,¹ they are part of academic literacy generally. If language assessment instruments at the university level are not capable of dealing adequately with causal explanations, we have reason to question their value and specifically their relation to scientific literacy in particular and academic literacy in general.

Our research strategy is a benchmark strategy: We elicited discourse of causal explanations, analyzed the texts to illustrate relevant contrasts between them, and used the texts and their analysis as benchmarks for the critical examination of models of testing and test procedures. Our aim was to illustrate a test case that illuminates some of the confusion surrounding the assessment of discourse. We have discussed this case with two different groups: those familiar with an SFL approach to discourse and applied linguists who have an interest in second language evaluation. It has been our experience that there are significant mutual misunderstandings between these groups, and for that reason we have written this paper for both audiences.

The two written explanations of the water cycle

The following two explanations of the water cycle were elicited using a visual prompt, a diagram of the water cycle. Explanation A was written by a secondary school teacher whose first language is English, and Explanation B was written by a university student who speaks English as a second language.

Explanation A:

The water cycle.

What are the processes that “water” goes through?

- 1) Initially, the water cycle begins as snow melts from the glaciers.
- 2) The water then meanders through various water sheds until it reaches rivers and lakes. Water eventually reaches the oceans.
- 3) Water, then, becomes water vapour (it evaporates into the air) and accumulates in what we call clouds.
- 4) The “clouds” then distribute water in the form of rain, snow, or sleet back to the mountains where the cycle begins again.

Explanation B:

The water cycle: The sun is the source of our water. The water, or hydrological, cycle begins when the sun heats up the ocean to produce water vapour through evaporation. This water vapour mixes with dust in the atmosphere and forms clouds. Cool air causes condensation of water droplets in the clouds, bringing about precipitation, or rain. This rain then falls into rivers, streams and lakes and eventually returns to the ocean, where the cycle begins again.

To analyze and therefore assess these causal explanations, Halliday’s work encourages us to ask three related questions. These questions and the concepts they include are not generally understood in the language assessment literature, yet without them we do not see how it is possible to assess the discourse of causal explanations adequately. We will deal with each question in turn.

(1) How do writers use language as a resource for meaning?

If we contrast Halliday’s view of linguistics with that of traditional grammar (see Derewianka, 1999) as in Table 1, we find very different implications for the evaluation of discourse. Where Halliday deals with discourse, considers functions of language and how they evolve in our culture, and explores how discourse varies with context, traditional grammar deals with the sentence, considers the form and structure of language, and offers a general description of language. Where Halliday sees language as a resource for making meaning in context and language learning as extending resources for making meaning, traditional grammar sees language as a set of rules and language learning as acquiring correct forms. In Halliday’s view, meaning and form are intrinsically related. In the traditional view, meaning and form in language are typically not seen as related but under a conduit metaphor: Language is a conduit through which meaning flows rather than being inherently associated with meaning. The obvious implication for the evaluation of discourse from traditional grammar and the language as rule perspective is to evaluate the correctness of form to see whether language rules are violated or not. Judgments about the meaning of discourse may be made at the same time, but they are usually holistic, impressionistic and, consistent with the conduit metaphor, are made independently of the evaluations of form. The implications for evaluation from Halliday’s view are much different. The emphasis shifts from what the

learner *cannot* do to what the learner *can* do. This view encourages us to evaluate discourse as making meaning using linguistic resources in context. How does the writer relate form and meaning? How do our writers use language as a resource for causal explanation? How do they express causality? Do they control a wide or narrow range of causal language?

Systemic Functional Linguistics	Traditional Grammar
Discourse Functions of language and how they evolve in our culture to enable us to do things How discourse varies with context Language as a resource for making meaning Language learning as extending resources for making meaning in context Evaluate discourse as making meaning with resources in context	Sentence and below Form and structure of language General description of language Language as a set of rules Form unrelated to meaning (“conduit”) Evaluate correctness. Judge meaning independently from form.

Table 1. Assumptions of SFL and traditional grammar (Adapted from Derewianka, 1999, p. 19)

(2) How do writers build causal meanings from a range of different features of language and discourse?

Halliday and Martin (1993) argued that the language of science—discourse that can be challenging and alienating to both children and adults alike—reflects the evolution of scientific knowledge itself. The historical development of scientific English discourse offers more powerful ways of talking, writing, and thinking about cause and effect relationships. Causal discourse, the authors suggested, has taken the following developmental path:

from A happens; so X happens
 because A happens, X happens
 that A happens causes X to happen
 happening A causes happening X
to happening A is the cause of happening X
 (Halliday & Martin 1993:66)

Thus, some of the language features that are important for causal explanation will be verbs of doing and happening, causal conjunctions such as *so* and *because*, causal dependent clauses, lexical verbs of causation such as *cause*, and nominalizations such as *happening*.

(3) How do writers construct causal (and other) lines of meaning?

Both writers were presented with a visual of the water cycle, a series of events that are linked together causally, and their written explanations each construct an account of the water cycle. Each contains a line of meaning that runs through the discourse. The notion of lines of meaning was put forth by Longacre (1990), who proposed that there exist in discourse lines of meaning which are constructed using various language features. His discussion, which revolved around the storyline in narratives, illustrated how, for example, verbs are used in clauses to distinguish the main storyline from other strands of development. Existential verbs, he claimed, serve to establish the setting whereas verbs in the past continuous refer to actions which are occurring alongside the main storyline. He further suggested that other forms of discourse, such as procedural or hortatory discourse, have similar lines of meaning which can be constructed, visualized, and analyzed in similar ways. An adequate analysis of a text, therefore, needs to show how grammar and meaning interact intimately to construct the line of meaning in the discourse.

A competent reader of the water cycle texts should be able to reconstruct the line of meaning (i.e., she should be able to draw the cycle of the water cycle), and a competent evaluator should be able to recognize how the line of meaning is constructed by the relevant language features. Clearly the line of meaning need not be the only aspect of discourse that is of interest in causal texts, but it is central for many causal explanations and any evaluation which fails to deal with it is inadequate. The line of meaning provides a natural focus for the evaluator and suggests a holistic view of the discourse.

Comparison of participants' texts

Table 2 compares the language features of the two explanations with respect to time sequence and cause. The writer of Explanation A has constructed a *time* line of events in time sequence; in other words, event A comes before event B. She has constructed a line of meaning with time conjunctions which reflects an explicitly sequential explanation of the water cycle. For example,

- 1) *Initially*, the water cycle begins as snow melts from the glaciers.
- 2) The water *then* meanders through various water sheds until it reaches rivers and lakes.
Water eventually reaches the oceans.

The writer of Explanation B has constructed a *causal* line: a series of actions (and some events) in a cause and effect relation; in other words, A causes B. He has constructed a line of meaning with causal language features that reflects an explicitly causal explanation of the water cycle. For example, consider his use of lexical verbs of causation, closely linked to nominalized processes:

- ... the sun heats up the ocean to produce water vapour through evaporation.
- Cool air causes condensation of water droplets in the clouds, bringing about precipitation, or rain.

Unlike the writer of Explanation A, the writer of Explanation B uses a range of language resources for

causal meaning. Where Explanation A uses one explicitly causal feature ('clouds distribute water'), Explanation B uses a number of causal features of different kinds. Actions, causal dependent clauses, lexical verbs such as *cause*, nominalizations, and causal metaphors are all features which play a part in constructing the causal line of Explanation B. Linked together through cohesion, they contribute to the line as a structure of ideational meaning.

Thus our analysis shows how Explanation B, but not Explanation A, draws on causal resources of English to create an explicit causal explanation. Is second language assessment at the university level capable of recognizing such differences and able to deal adequately with causal explanation? If not, we may question how well assessment deals with scientific literacy and academic literacy in general.

It is worth noting here that a language as rule perspective is not able to address these texts as causal explanations at all. The language as rule perspective does not address how meanings are constructed by a discourse; it can only describe the forms occurring in a discourse. For those who see language acquisition as the learning of correct form, there is nothing to be said by comparing our two texts, since neither contains grammatical errors.

Language features	Explanation A	Explanation B
Time sequence		
Events (sequence of happenings)	water cycle begins, snow melts, water meanders, water reaches lakes and rivers, water reaches oceans, water becomes water vapour, it evaporates, cycle begins again	water cycle begins, rain falls into rivers, returns to ocean, cycle begins again
Time conjunctions	1), 2), 3), 4), initially, as, then, until, eventually	then, eventually, and
Dependent clauses of time	as snow melts	when the sun...
Cause		
Actions	clouds distribute water	sun heats up ocean, produces water vapour, evaporation causes condensation, bringing about precipitation
Causal conjunctions	Ø	Ø
Causal dependent clauses	Ø	to produce water vapour through evaporation
Cause/means as circumstance	Ø	
Cause as process	Ø	produces, causes, brings about
Nominalization	Ø	evaporation, condensation, precipitation
Causal metaphor	Ø	The sun is the source of our water.

Table 2: Language features of the two explanations

Models for evaluating communicative competence and language as rule

We have now analyzed our causal explanation texts, showing a contrast between them. We now move to the next stage of our study and use the texts and their analysis as a benchmark for the critical examination of models of testing and test procedures. In this we have shown how the analysis of these causal explanations relies on the assumptions of SFL (language as resource), and we have noted that a language as rule perspective is not able to address these texts as causal explanations. When we examine the assumptions that underlie current theory and practice in the assessment of discourse in the second language, we find that they are closer to language as rule than they are to language as resource.

At first glance, it would appear that currently accepted models for evaluating communicative competence or communicative language ability would be consistent with a language as resource perspective since these models were influenced by early work in functional linguistics. According to Canale and Swain (1979, 1980), discourse can be assessed in an integrated manner using their theoretical framework of communicative competence. They defined communicative competence as the relationship and interaction among three main competencies: grammatical competence, or the knowledge a speaker has about the rules of grammar; sociolinguistic competence, or the knowledge of the rules of language use in social situations; and strategic competence, which refers to the knowledge of strategies, both verbal and non-verbal, that a speaker may have which can compensate for deficiencies in the other two areas of competence. Although Canale and Swain stated that each of these competencies can be examined separately, they emphasized that for second language teaching and testing purposes, a communicative approach “*must integrate* aspects of both grammatical competence and sociolinguistic competence” (1980, p. 6). In their view, an integrative theory of communicative competence may

be regarded as one in which there is a synthesis of knowledge of basic grammatical principles, knowledge of how language is used in social contexts to perform communicative functions, and knowledge of how utterances and communicative functions can be combined according to the principles of discourse. (p. 20)

However, models for evaluating communicative competence, while appearing to accommodate causal explanation and the systemic analysis of discourse generally, contain deep differences of assumptions that set the stage for mutual misunderstanding. This was brought home to us powerfully by the remark of a language assessment specialist who commented: “Both Explanations A and B are competent. I don’t see why you have to bother to say anything further than that.” In other words, from this specialist’s standpoint, there was no need for the analysis of causal explanations or for the discourse analysis that lay behind it. Both writers were “competent” in the sense that they did not violate the rules of the language. The notion of competence did not extend further than that, and did not need to.

These fundamental differences can be illuminated by a closer analysis of the development of Canale and Swain’s theory. Bachman (1990) developed Canale and Swain’s work, but he accomplished this primarily by expanding the taxonomy of competencies: From his work we now have grammatical,

textual, illocutionary, and sociolinguistic competence, plus strategic competence. Interestingly, he included Halliday's work in a number of these taxonomic categories. However, Bachman assumed the standard view of language as rule with its concomitant of language error, and extended it via the notion of convention. Grammatical competence is spoken of in terms of rules, textual competence is stated in terms of "conventions for organizing discourse," illocutionary competence is described as "the pragmatic conventions for performing acceptable language functions," and sociolinguistic competence is said to encompass "the sociolinguistic conventions for performing language functions appropriately in a given context" (pp. 88, 90). There is no evidence of any consideration given to language as a resource for meaning in Bachman's work. Without the concept of language as resource for meaning, a taxonomy of communicative language ability is too easily interpreted as a series of categories in which a learner can be judged right or wrong. At its most basic, it can be seen as a simple extension of the practice of assessing the language learner's discourse for grammatical errors. Bachman's development now presents us with an extended taxonomy in which the learner can be incorrect in a variety of ways, not only in grammar but also in discourse conventions, sociolinguistic appropriateness, and so on.

Canale and Swain's emphasis on integration of the parts of their taxonomy seems to call precisely for the kind of question that we were pursuing in our systemic analysis of causal explanation: "How does the speaker/writer create a line of meaning which is causal, using the resources of discourse semantics and lexico-grammar, taken together?" From the standpoint of discourse analysis, it is very hard to see how discourse evaluation could avoid facing a question like this, yet integration, in that sense, has not been developed at all by Bachman. He has interpreted it in quite a different way. Bachman considered his model of strategic competence to provide the notion of integration which Canale and Swain presented. He redesigned strategic competence to provide a model that performs assessment, planning, and execution functions in determining communication goals. This model, however, is a psycholinguistic model of the generalized processes involved in communication that is vastly different from a discourse analysis model of the kind offered through systemic functional linguistics, a model which has the ability to address questions such as how causal lines can be constructed from the resources of language. Questions such as these are ignored, since they become irrelevant when models of communicative competence based on the assumption of language as rule are used as frameworks for the assessment of error, because the assessor does not need to consider how categories of the taxonomy are related to each other. A judgment of error in each category on the checklist is all that is required, as we will see in the next section.

Using the assessment instruments

In order to see whether assessment tools based on the taxonomic model of communicative competence discussed above could address the question of how causal lines differ in their construction, our water cycle texts were explored using two assessment instruments. The first was a locally developed tool for evaluating the communicative competence of potential second language teachers, and the second was the scoring guide for the Test of Written English (TWE). As can be seen from the discussion below, neither

instrument was able to distinguish between the two texts, thereby suggesting that neither could adequately assess causal explanations.

The initial assessment was done by two raters who were experienced in using the scoring criteria of a test locally developed from Canale and Swain's Theoretical Framework of Communicative Competence (1979, 1980). Previous statistical analyses had determined that the overall reliability for this instrument was high. The actual test had ten tasks, and in the assessment of each, errors were classified in one of five categories of competence: linguistic, sociolinguistic, discursive, strategic, and receptive. Linguistic competence referred to the knowledge an individual has of the lexis, phonology, morphology, and syntax of the language, and included the ability of the individual to use this knowledge to create words and sentences. Within this category, the raters distinguished among vocabulary, mechanics (including spelling and punctuation), and morphology and syntax. Sociolinguistic competence in this instrument was concerned with whether the writer or speaker could carry out linguistic functions in particular social contexts. Having discursive competence meant that the individual was able to demonstrate an ability to connect sentences into cohesive and coherent discourse using devices such as transitional phrases and conjunctions, or in the raters' words, whether the text made sense or not. The definition of strategic competence for these raters was whether or not the writer got the message across, but the category primarily referred to the use of survival strategies such as paraphrasing and guessing to compensate for a lack of other competencies. Finally, receptive competence was the ability to understand what was being asked or said. An error in any of the above competencies could be assigned only to one category, determined by the raters. The scoring was binary within the categories for each of the ten tasks; in other words, the raters could give one point to a category if there were no problems exhibited and zero if there were one or more examples of errors.

Our writing task paralleled one from their test that the raters could identify with to some extent, yet there were still important differences to consider. Our water cycle task did not assign a specific role to the writer, and nor did it give the writer a context or an audience. Their test stated these explicitly for the writers. These differences prevented the raters from addressing the category of sociolinguistic competence because they could not judge whether the writer was communicating appropriately for the assigned context, role, and audience. This comment was brought up very quickly in the assessment discussion.

Instead of noticing a difference in the two texts and crediting Explanation B as being the better of the two, these raters judged them as equal, assigning both a score of five out of six. The writer of Explanation A was faulted in the area of mechanics because she spelled *watersheds* as two words. A supposed mechanics problem was the reason for an imperfect score for Explanation B: The raters felt that it was wrong to use an upper case letter after a colon. Although the assessment focused on the negative aspects of the writing, the raters had good things to say as well. One rater, for example, appreciated the writer's use of *meanders*, commenting that it was "a nice word." Explanation B, both raters agreed, made sense:

This one is textbook, unless no I don't see anything wrong unless ah but I'm not a scientist.
But as far as the English and the you know like the sequence, it makes sense.

Unfortunately, there was nothing in the assessment instrument that allowed the raters to add points for what they considered exceptional because it was assessing the errors made rather than the resources each writer was exhibiting.

Intuitively, however, these raters judged Explanation B as definitely “more advanced,” “scientific,” and at a “higher caliber than” Explanation A, yet they admitted that the assessment instrument would not account for this discrepancy:

Researcher: So [Explanation A] and [Explanation B] you consider to be equal?

Rater 2: Yeah, I would say. Yeah. Yeah.

Researcher: But intuitively?

Rater 1: To me [Explanation B] is a higher caliber than

Rater 2: Yeah, yeah, yeah.

Researcher: It's more sophisticated or?

Rater 2: Yeah, yeah.

Researcher: But there is nothing in your instrument that would make that difference?

Rater 1: No.

Rater 2: No.

They stated that the only way they could distinguish levels with this assessment instrument was when there were errors, and because these two explanations did not contain a great number of errors, they were difficult texts to assess.

Because the instrument above was developed to assess general communicative competence in a second language rather than the ability to write academically, we decided to see if a difference between the explanations could be detected by using a more appropriate instrument, the scoring guide for the Test of Written English. This guide breaks competency into six levels, with a score of six being the highest. At each level there is a brief general description of what the rater can expect at that level, followed by a somewhat more detailed list of points to look for in the writing. Although these details are in point form, our exploration revealed that the points can be addressed quickly by the rater by turning them into yes/no questions, and the score which is arrived at will probably be the score which contains the greatest number of positive responses to these questions. A paper which is rated at a level five, for example, suggests in the general description that it should demonstrate that the writer has a good level of competence in syntax and rhetoric regardless of occasional errors. Among the list of points to consider for a level five paper is that it should be quite well developed, although not as well developed as a level six paper and with fewer details. It will also likely contain more errors than a level six paper. On the other end of the scale, a level one paper demonstrates incompetence, may be severely underdeveloped, and will contain serious and persistent errors.

Our initial step was to find individuals who had been trained to use this assessment tool. As our interest was primarily exploratory, we asked the one trained individual we knew to hold a training session for eight volunteers from our university's language and literacy department. Two of the volunteers had

undergone the training session two years previously. After the trainer held the session, the participants were asked to score various samples of writing to establish whether they were using the scoring guide correctly. After examining the results of the training, we asked the participants, in pairs, to work together to assess the two water cycle explanations using the TWE scoring guide, and these interactions were audio-taped. The individual who had conducted the training session was also asked to rate the two explanations. The results of the nine raters paralleled that of the two raters who used the communicative competence instrument above; in other words, all eleven raters judged the two explanations to be more or less equal regardless of the instrument they were using. Whereas the previous raters had assigned scores of five out of six to each of the two explanations using their locally developed tool, the raters using the TWE scoring guide gave Explanation A an average score of 5.67 and Explanation B a score of 5.89, the difference being non-significant. According to the TWE raters, the reasons for the lower score given to Explanation A was that it resembled an outline more than a written explanation and that it overused the word *then*.

As with the raters who used the locally developed instrument, these raters had an intuitive judgment of the explanations that did not always match the TWE score they assigned. Moreover, just as with the previous raters, there was general agreement among the TWE raters that when two texts are relatively well written and contain few or no errors, they are difficult to assess because both can receive only the highest grade despite any apparent differences between them. A good example of this difficulty surfaced in the discussion of one pair of raters who had just finished discussing Explanation A and had assigned a level six to it, justifying their decision by discussing each point on the scoring guide. When they read Explanation B, they began to laugh. Compared to Explanation A, they claimed, Explanation B was a much better piece of writing, yet they had confidently given Explanation A the maximum score based on the scoring guide. Laughing, one of the raters concluded, “So much for criteria-referenced marking!” In general, all raters agreed that intuitively Explanation B was a better example of academic writing. It was judged as sounding more scientific and academic than the “more elementary school” writing of Explanation A, and was credited as being able to show “how it’s all working.” In Explanation A, one rater claimed, “the phrases are there but not the process,” and according to the intuitive judgments of these raters, the process was explained much more clearly in Explanation B.

The results of these exploratory assessments can be related to the traditional view of language as rule. Using both instruments, a broken rule—an error in other words—suggests that the writer is lacking the necessary knowledge to complete the task successfully. In these two explanations, however, the writers were successful and no rules were seriously broken, so the resulting scores were similar. Both instruments encouraged the raters to adopt a binary-style checklist method of assessing categories of error. Once the raters narrowed a text down to a particular category, they asked themselves the questions listed in the scoring guide. Are there mechanical errors? Does it contain some serious errors that occasionally obscure meaning? Yes or no? There were also questions which solicited vague impressions of the meanings of the text as a whole, but which did not ask the rater to consider the specifics of the text, which presumably was considered to be a mere conduit for these meanings: Did the writer understand the task? Is the paper adequately organized? The answers to all the above and other questions provided the

rationale for scoring the paper at a particular level or moving the score up or down. The weakness of this taxonomic checklist approach can be seen by contrasting it with the description that Halliday and Martin (1993) offered regarding the integration of features in scientific writing:

Whenever we interpret a text as ‘scientific English’, we are responding to clusters of features...But it is the combined effect of a number of such related features, and the relations they contract throughout the text as a whole, rather than the obligatory presence of any particular ones, that tell us what is being constructed is the discourse of science. (p. 56)

Although this binary, taxonomic, and error-based approach may be used to record in some sense whether the writer has or does not have the linguistic knowledge to complete the writing task, it is not able to show *how well* he uses that knowledge to construct the text as a whole, and therefore it is not able to record a difference between these two explanations. Furthermore, there was considerable evidence during both assessments that raters were intuitively able to recognize that Explanation B was superior to Explanation A, strongly suggesting that the assumptions built into these testing procedures actually suppress the reporting of differences easily noticed by the raters.

Conclusion

Our examination of the evaluation of written causal explanations seriously questions the adequacy of present models and practices that purport to assess competence in the second language. These models, and the testing practices which are associated with them, reflect underlying rule-based assumptions about the nature of language and the role of evaluation which at their worst degrade the evaluation of discourse to a reductive exercise in the detection of errors. If these underlying theoretical assumptions are not recognized and changed, it is hard to see how the major implications of the SFL analysis of discourse can be incorporated so as to change present practices. What is most likely to happen is that present practices will grind remorselessly on within their present goals.

This chapter has demonstrated that it is possible to assess causal explanations with due regard to the three important questions posed earlier and repeated here: How do speakers and writers use language as a resource for meaning? How do they construct causal (and other) lines of meaning? How do they build meanings such as causality with a variety of features from different components of language? It is imperative at this time that we extend this work to move towards more acceptable practices in the assessment of scientific and academic literacy.

References

- Bachman, L. (1990). *Fundamental considerations in language testing*. Oxford: Oxford University Press.
- Canale, M., & Swain, M. (1979). *Communicative approaches to second language teaching and testing*. Ontario: Ministry of Education.

- Canale, M., & Swain, M. (1980). Theoretical bases of communicative approaches to second language teaching and testing. *Applied Linguistics*, 1 (1), 1-47.
- Derewianka, B. (1999). *Introduction to systemic functional linguistics: Institute Readings*. The 26th International Systemic Functional Institute & Congress. Singapore: Department of English Language and Literature, National University of Singapore.
- Halliday, M.A.K., & Martin, J.R. (1993). *Writing Science: Literacy and discursive power*. Washington D.C.: The Falmer Press.
- Longacre, R.E. (1990). *Storyline concerns and word order typology in East and West Africa*. Los Angeles, CA: The James S. Coleman African Studies Center and the Department of Linguistics, UCLA.
- Martin, J.R., & Veel, R. (1998). *Reading Science*. New York: Routledge.
- Rowell, P.M. (1997). Learning in school science: The promises and practices of writing. *Studies in Science Education*, 30, 19-56.
- Tang, G.M. (2001). Knowledge Framework and classroom action. In B. Mohan, C. Leung, & C. Davison (Eds.), *English as a Second Language in the Mainstream* (pp. 127-137). Harlow: Pearson Education Ltd.

Footnotes

- ¹ Tang (2001) offers an example of ESL students in the elementary school writing causal explanations of the Fall of the Roman Empire.