

Narrative Comprehension, Causality, and Coherence



**Essays in Honor of
Tom Trabasso**

Edited by

Susan R. Goldman
Vanderbilt University

Arthur C. Graesser
The University of Memphis

Paul van den Broek
University of Minnesota



1999
Mahwah, New Jersey

LAWRENCE ERLBAUM ASSOCIATES, PUBLISHERS
London

13

Content Integration and Source Separation in Learning From Multiple Texts

M. Anne Britt
University of Pittsburgh
Charles A. Perfetti
University of Pittsburgh
Rebecca Sandak
University of Pittsburgh
Jean-François Rouet
University of Poitiers, France

In recent years, we have been studying how students learn and reason from history texts, with the emphasis on the plural—multiple texts. We began with a simple assumption about the form of the typical history textbooks read by students, namely that they were essentially narratives. This implied that students' understanding of history texts would typically include a representation of temporal-causal event sequences. Work by Tom Trabasso and his colleagues on the causal analysis of narratives provided a fine starting point for our goal of describing this temporal-causal component in history stories. Building on early story-understanding research (Mandler & Johnson, 1977; Omanson, 1982; Rumelhart, 1975; Stein & Glenn, 1979; Thorndyke, 1977), Trabasso and colleagues proposed that causal connections play a critical role in establishing meaning and coherence in a story. By providing a detailed and explicit model of causal analysis and spelling out tests for necessary and sufficient conditions, Trabasso and colleagues were able to predict what readers recall (Trabasso & van den Broek, 1985),

include in a summary (Trabasso & van den Broek, 1985; van den Broek & Trabasso, 1986), and judge as important (Trabasso & Sperry, 1985; Trabasso & van den Broek, 1985). Although this type of causal analysis clearly captured an essential part of narrative understanding, an analysis of the use of multiple texts required more than this.

Our goal here is to explain the course our work on multiple texts has taken, given its initial foundation on the simple idea of history as narrative and understanding as taking account of temporal-causal event structures. First we describe our initial work to develop temporal-causal models that reflect the more complex demands of long multiple texts and summarize what we learned in one study. Then we describe a later phase of our work that required the development of an additional level of text representation to capture what readers understand about the source as well as the content of texts. In the last section, we briefly summarize the results of a recent study that tests possible models of how readers represent source information.

ADAPTING CAUSAL MODELS TO MULTIPLE TEXTS

From a student reader's point of view, a salient component of learning history is learning events and the causal-temporal connections among these events. Indeed, although historical scholarship goes well beyond the event narrative, it is fair to say that reconstructing the causes that explain events is a critical part of much written history. To mirror the narrative component of a typical text, a reader's representation must capture these temporal-causal relations, making history learning similar to general story understanding. However, for reasons we give later, it became clear that one cannot simply use existing analyses (e.g., Trabasso & van den Broek, 1985) to capture these representations. Some modifications were required.

Our first modification was to increase the grain size. Understanding the complex events related through actual history texts required us to use experimental texts that were much longer and more complex than the narrative texts previously studied. We intended to allow students to learn from texts that were between 2,000 and 13,000 words in length rather than the few hundred words typical in story-understanding research. A clause-by-clause analysis was impractical, whereas event units were both more practical and more appropriate for the level of understanding we wanted to capture. Increasing the grain size from the clause to the level of the event

solved the problem. Each event unit can be elaborated either in a single clause or in several paragraphs; each higher level event, in effect, is a substory, more or less elaborated depending on the knowledge and purpose of the author. Our assumption here is that causal structures are valid for macroevents (e.g., the June 6, 1944 Normandy invasion CAUSED the Western Front military defeat of Nazi Germany) just as they are for microevents (the ringing of the alarm clock CAUSED General Montgomery to wake up that morning).

Our second modification was to develop a shared template to capture common events described by different sources. Our assumption was, and is, that learning historical narratives requires the reader to construct a model of the situation from many sources and to integrate information from these sources rather than constructing a separate model of each text. To capture the kind of information that would be represented in an integrated model, we needed an abstracted representation of events and their surrounding elements that were shared among the various texts. To accomplish this, we read dozens of sources about the Panama Canal story and then constructed a single common template of the events and causal-temporal structures that comprise the typical Panama Canal story. The abstract causal-temporal event template developed for this story is shown in Fig. 13.1.

Using this template, we were able to describe the story contained in a selected set of texts on the Panama Canal and study a small group of students who learned from them (Perfetti, Britt, & Georgi, 1995). Over a 4-week period, the students read a lengthy (approximately 15 pages) scholarly or popular text describing the U.S. acquisition of a canal in Panama at the turn of the century. After reading at home, students came into the lab to write both a short and a long summary of the events as they understood them so far. They then proceeded to answer a series of comprehension questions and reasoning probes. Applying our causal-temporal event template to the information provided by students in their summaries and in answers to our questions, we obtained three general results. First, during their initial reading, students rapidly and extensively learned *core events* (i.e., those that were central to the story), and a limited number of the supporting details for these events. Second, over subsequent readings, students gradually learned the *noncore events*, that is, events that were not necessary to a coherent telling of this story. Students also continued to learn supporting details for both core and noncore events with each reading. Third, although they learned more of the complete story (core and noncore events elaborated with supporting details) over time, the composition of the students' sum-

maries remained rather stable. Their summaries were predominantly a synopsis of the core events and the causal-temporal connections of these events.

These two developments (larger grain size and a common template for shared events) enabled us to represent temporal-causal narrative structure across several texts of substantial size and to effectively describe the time course of learning across readings in the Perfetti et al. (1995) study. These developments also enabled us to represent the integrated accounts that the students provided in their summaries and answers to questions.

Further analysis of the student interviews revealed several occasions on which students accurately attributed information to a specific author. This need to represent the distinct contribution of a given text to the story is apparent in many real multiple-text learning situations, and it is particularly salient in history. History texts comprise multiple, and frequently contradictory, causal models. A representation of the texts must include not only what is shared among them, but also what is not shared and what is discrepant. When several possible causal connections exist between events, coherence of the representation cannot be estimated solely by the causal connectedness of a single text. For instance, two authors may propose incompatible causes for a single event. By allowing events or a series of events to be associated with the document's source, readers can maintain overall coherence in a mental model while still allowing for otherwise incoherent causal connections when integrating information across text. The reader can remember that "according to X, A caused B"; whereas "according to Y, C caused B." Thus a comprehensive model of multiple-text representation must include the potential for source information, and it becomes an empirical question as to the conditions that promote the encoding of source information.

This interesting combination of integration and separation of material from different sources led us to consider the range of possibilities for how to represent the kind of complex integrated-yet-separated representation that learners may develop from multiple texts. Do learners keep track of source information while they are learning the shared story? Surely, we thought, some learning of source information is likely, but it is unlikely that every piece of information is tagged for its source. Theoretically, the question is how the causal-temporal event model might include source tagging for its events. Our armchair considerations of the possibilities led us to four classes of models. Before considering these models, we first examine studies that appear to address these issues of integration and separation.

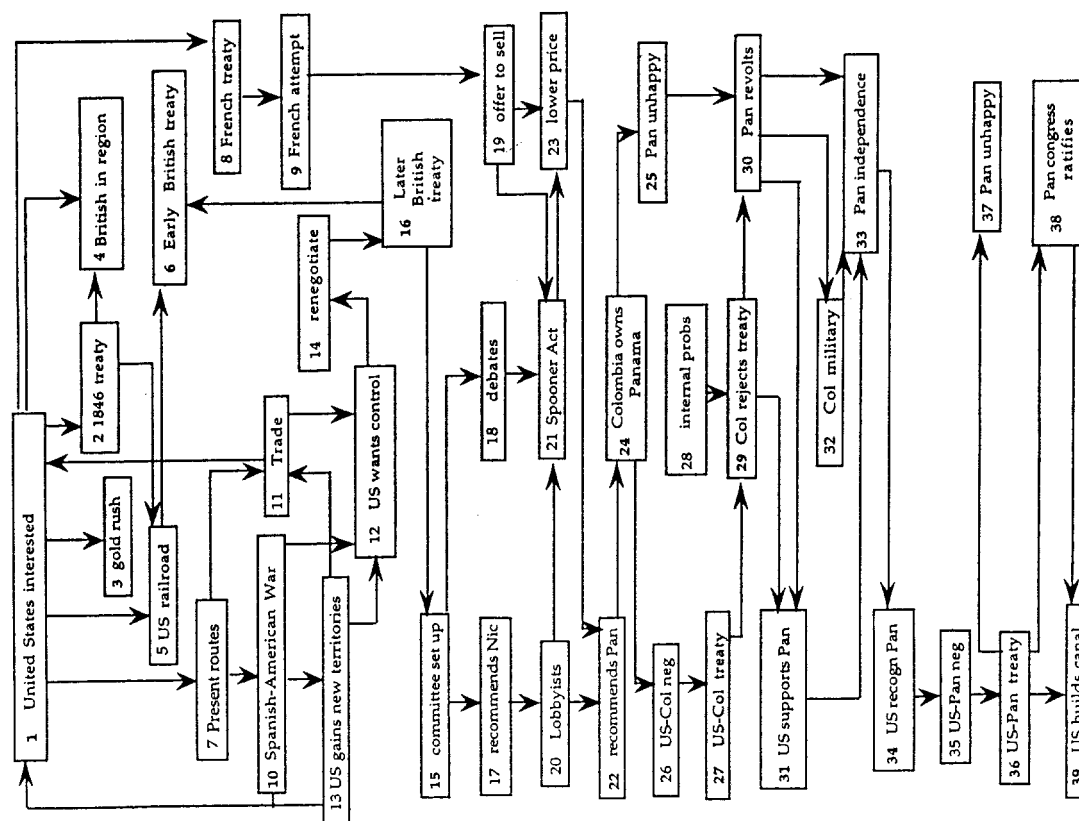


FIG. 13.1. Causal-temporal event template of the history of the acquisition of a U.S. canal in Panama. From Perfetti, Britt, and Georgi (1995). Copyright 1995 by Lawrence Erlbaum Associates. Adapted by Permission.

Evidence of Integration

Wineburg (1991) had historians and advanced high school history students read several documents and think aloud during problem solving. He found that historians employed a "corroboration" heuristic, which involves comparing important information against what one knows from other sources before integrating the information. Wineburg showed that experts, but not high school students, integrate information derived from multiple sources and they do this integration while reading.

Further evidence of integration comes from a series of studies by Wiley and Voss (1996; Voss & Wiley, 1997). College students read about the Irish Potato famine from eight separate sources or from a textbooklike chapter that weaved a story out of the same information. Wiley and Voss observed that students who read separate sources used more connectors in their essays and judged possible inferences from the texts better, whereas the single source lead to better verbatim recognition of information. These studies imply integration of information across separate sources leads to greater knowledge transforming rather than knowledge telling.

Evidence of Separation and Source Tagging

Several other studies of multiple document reading in history suggest that at least some source tagging is conducted as the reader goes through the texts. Wineburg (1991), in the study described previously, found an additional strategy employed by the expert historians. They used a "sourcing" heuristic, the encoding and evaluating of source information before reading the contents of a document. Because historians learn about the source of a document before reading, they have the information necessary to make connections between source and content during reading. Recall that experts also corroborate while reading, which requires them to use this source tagging to reactivate previous authors' versions of the events and causes. Both the sourcing and the corroboration heuristics are precisely the type of processing that would facilitate the construction and use of a representation that includes source tagging.

One of our previous studies provides further evidence that readers' model of the information contained in multiple texts is augmented with source tagging. Rouet, Britt, Mason, and Perfetti (1996) gave 24 college students several problem statements about the U.S. acquisition of a canal in Panama

(e.g., To what extent was the U.S. military intervention in the Panamanian revolution justified?). For each problem statement, they read short excerpts from different types of documents (e.g., textbook, historian essays, participant accounts). Students wrote essays to answer each problem and then ranked each document in the set as to its trustworthiness and its usefulness. Two findings support the hypothesis that students had encoded a link between the author and its content. First, 56% of student essays contained at least one explicit reference to the source of a statement. Most of these (91%) were, in fact, correct. Second, in analyzing the justification of their trustworthiness and usefulness ranks, we found that many had identifiable origins (e.g., this document was useful because it "cites precedent of guarding the railroad, but not interdicting Colombian troops"). Of the 110 justifications that were specific enough to trace the origins to a single document, 88% mentioned at least one correct statement of content linked to a source. Furthermore, 92% of these justifications were accurate. Even without specific instructions, students did source tag at least some content.

These results demonstrate that readers can take note of specific sources among a set of related documents under conditions that promote problem solving. The question of the generality and spontaneity of sourcing remains to be determined. In a talk-aloud procedure, Kushner (1996) asked both undergraduate and graduate college students to verbalize their thoughts during the reading of history texts without any expectation of essay writing. Although the instructions did not refer to source information, all readers made comments about the source—as opposed to the content—of what they were reading, in one form or other—references to the author directly, or to document characteristics. Interestingly, graduate students made more comments about sources than undergraduates. Thus, Kushner's results suggest, at least in these explicit verbalizing conditions, both that readers attend to source information spontaneously during reading and that individual differences in this attending are related to education level.

The research thus far suggests that a reader's representation of multiple texts is partially integrated and at least sometimes includes some links between source information and its content. Beyond this, there seems to be little to constrain the theoretical possibilities. To give full recognition to the range of these possibilities, the next section outlines four models of how tagging may be represented. Then in the last section, we describe an experiment (Britt, Sandak, Perfetti, & Rouet, 1998) that tests the predictions of the four models.

MODELS REPRESENTING SITUATIONS AND SOURCES

The four possible models of how readers might represent information learned from two or more related texts vary along a dimension of integration from completely separate and nonoverlapping situation models of each text to a single highly integrated situation model of the events. The models also vary in the degree of source tagging from no tagging to complete tagging of all content.

Separate Representation Model

The *separate representation* model, illustrated in Fig. 13.2, assumes that readers form a unique representation of each text without connections among them. Although individual propositions are not tagged, the entire unique representation may be tagged if author information is prominently provided. Information from each text is not connected to information from any other text (i.e., not integrated). Consider the two texts shown in Fig. 13.2, which are excerpts from the Britt et al. (1998) materials that are described in the next section. The fictitious authors, Clark and LaCosta, are represented by shaded nodes in the diagram, which we refer to as *document nodes*. Document nodes represent information known about the source and the document itself and are connected to content through bidirectional links. The top part of the diagram represents the events (rectangles) and causal connections (solid lines) from Clark's text, whereas the bottom portion represents the events and causal connections from LaCosta's text. The defining feature of this model is that there are two distinct situation models, one for each text, and they are completely independent of each other. Here the reader fails to make a connection between the events that overlap between each author's version of the story (i.e., fails to integrate).

Under most circumstances, this model would be less than optimal because prior knowledge (i.e., the situation model of the first author's text) was not accessed during the reading of the second text. There are conditions, however, when readers may be more likely to form a separate representation model of multiple texts. For example, if the source of the text is distinct and elaborated (i.e., who the author is and what type of document it is), the reader may be more able to keep the representations clearly distinct. Furthermore, if the task is to learn author X's version of the story, a reader may inhibit the activation of prior knowledge during the encoding process,

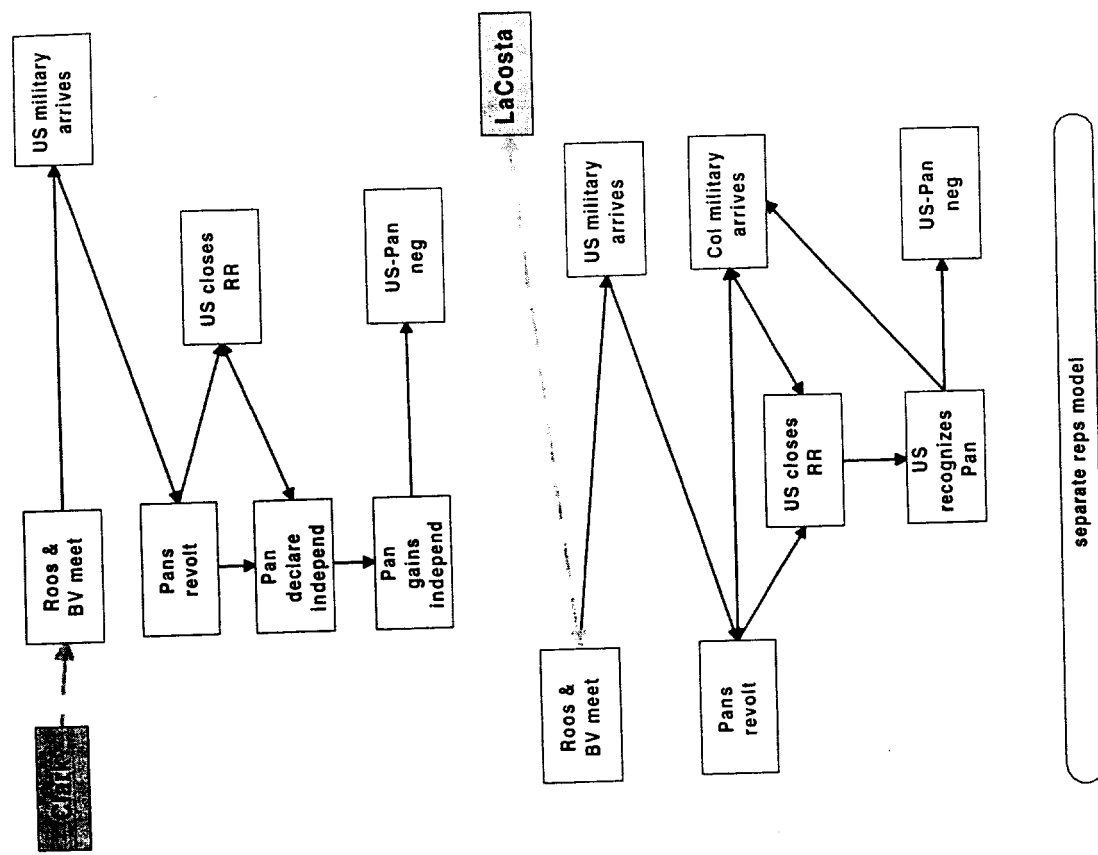


FIG. 13.2. Separate representation model representing a segment of text from two authors (Clark and LaCosta). Representation of situation (boxes and solid arrows) and intertext links (dotted lines) to a document node (shaded rectangles).

thus resulting in two unintegrated situation models. Anything that hinders integration will more likely lead to this type of representation. For example, if the texts do not highly overlap much in content then they may not be sufficiently related to trigger the activation of the other representation as the new situation model is created. The same might be expected if too much time has elapsed between readings. Either of these factors would result in a lessened ability to integrate events during the processing of the second text.

Mush Model

The *mush model* is essentially the opposite of the separate representation model. As shown in Fig. 13.3, the reader forms a combined representation that integrates information learned from all relevant texts. The essential element of this model is that there is no tagging or marking for where the information came from. All information is integrated completely, without attention to the source of the information. In the illustration, note that there

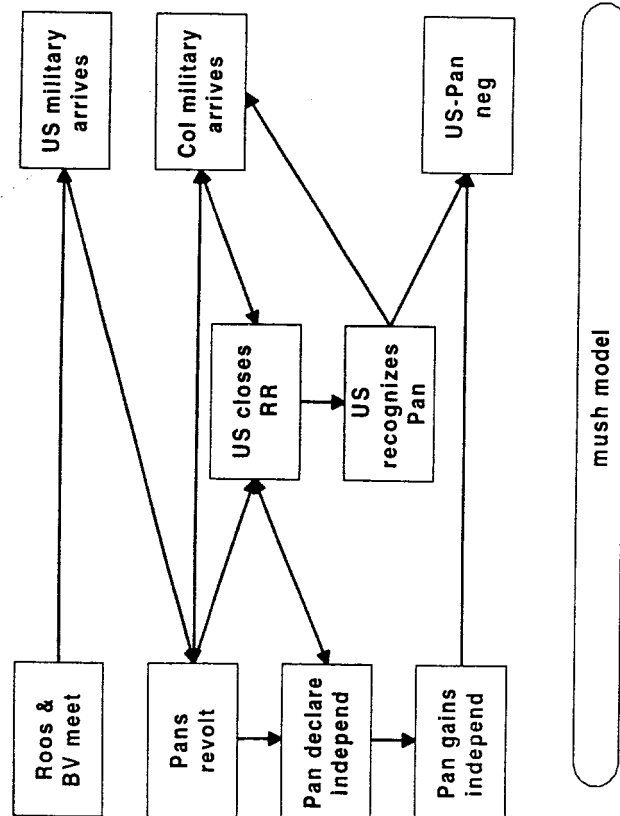


FIG. 13.3. Mush model representing a segment of text from two authors (Clark and LaCosta). Authors are not represented nor are the links between the author and the situation.

are no document tags on the events to reflect the model's assumption that author information is no longer available in memory.

Like the separate representations model, this representation is less than optimal. Whereas the first model lacks integration, the mush model lacks source marking. Readers may be more likely to form the mush model when the texts' sources are either never mentioned or not elaborated. A directed focus on simply learning the story, as opposed to solving some problem that involved text contradictions, may also increase the probability that the reader will form an integrated representation. Anything that encourages integration will increase the probability that this type of representation will be formed. This includes a high degree of overlapping events mentioned among texts and a short time span between reading of the multiple texts. We suspect that this form of representation is most common in students learning in school settings where they often see their task as one of learning about the topic, and because all sources are presented with relatively equal status, there is no reason to pay attention to where information comes from. The mush model, although promoting a situational understanding, leaves the learner short on source understanding and makes it difficult to verify information and evaluate it.

Tag-All Model

Whereas the mush model might be said to characterize minimalist learning of the kind that professors complain about in their undergraduates, the *tag-all model* is one for the scholar. As shown in Fig. 13.4, each and every event is tagged for its source. The shaded document nodes indicate the source information for that document and the dotted lines link the document node to the situation model derived from that author's text. There is also a link between document nodes to indicate the relationship between the two sources (e.g., support vs. oppose, agrees with vs. disagrees with, gives evidence for vs. gives evidence against, etc.).

The problem with this model is the high processing demands it implies. Each piece of information must be tagged for its source and integrated into a situation model. The only way out of this problem would appear to be a dubious assumption that links between the document node and each element in the situation model can be generated automatically, or at least with no cost. Such an assumption, we think, is unwarranted for typical student readers, but may be appropriate for content and author experts. Scholars who already have a rich knowledge base that includes the texts of a certain

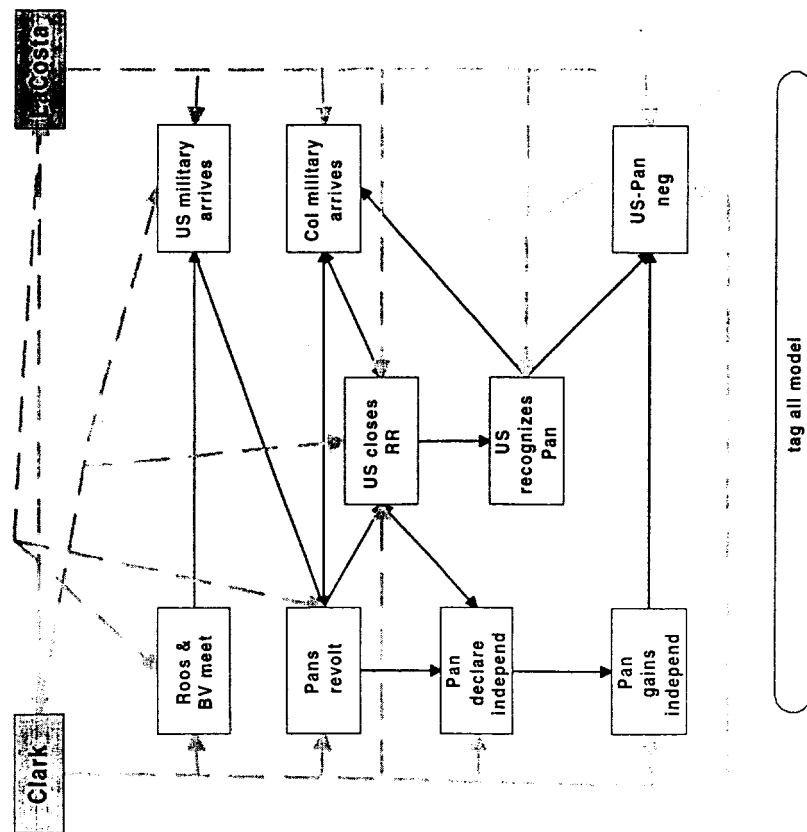


FIG. 13.4. Tag-all model representing a segment of text from two authors (Clark and LaCosta). Representation of situation (boxes and solid arrows) and intertext links (dotted lines) to a document node (shaded rectangles).

writer may come to "automatically" connect content and source for this author under some conditions. For a less expert reader, the effort may be too great.

Documents' Model

The model we proposed as most typical of a good reader's model of multiple-text learning in history is the *documents' model* (Perfetti, Rouet, & Britt, 1999), which has two interconnected components: a *situation model* and an *intertext model*. The intertext model is an added level of

representation that includes an interconnection of links among document nodes and links from these document nodes to one or more nodes in the situation model. The intertext model for our two experimental texts is shown in Fig. 13.5 with the dotted lines representing the links between document nodes and relevant situation model nodes. The documents' model is a highly integrated situation model of the events learned, with only the most important events (i.e., core events) tagged for the source, as represented in the intertext model. Comparing Figs. 13.4 and 13.5, the difference between the tag-all model and the documents' model should be immediately apparent. Figure 13.5 shows only limited marking of source information as indicated by the links from the shaded Document Nodes.

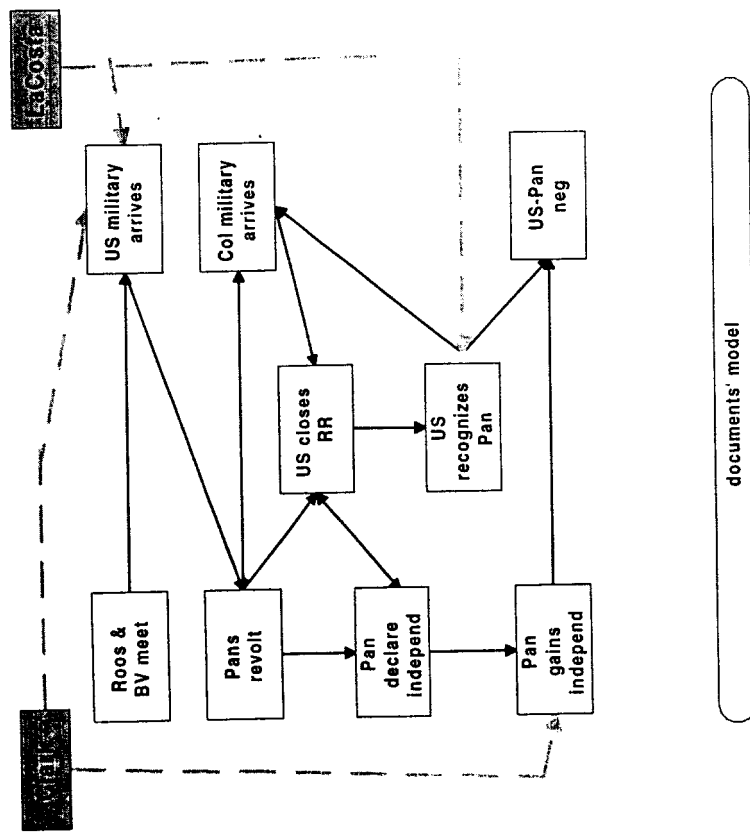


FIG. 13.5. Documents' model representing a segment of text from two authors (Clark and LaCosta). Representation of situation (boxes and solid arrows) and intertext links (dotted lines) to a document node (shaded rectangles).

Document nodes, which allow marking of events in the situation model, when fully elaborated contain information about the document's *source* (including information about the author, setting, and form of the document), the author's *rhetorical goals* (including the author's intents and perceived audience), and a synopsis of the *content* of the document. The document node for a particular text will be more or less elaborated depending on the quality of information provided to a reader about the source and also the reader's task (e.g., learning from more than one text written by the same author). This information can be used by the reader to connect a piece of information to its source (e.g., "according to author X ...") and to state the rhetorical relationships among texts (e.g., "based on Y, author X claims that ..."; "document X contradicts document Y ...", etc.).

The conditions that promote or discourage source tagging need to be empirically determined. There is nothing intrinsic to the documents' model framework that constrains information tagging. Rather it is a model that allows, in principle, complete tagging or no tagging, with parameters to be established by considerations outside the model; for example, data of the kind we summarize next or a theory of text "importance" in which causal information (Tabasso & van den Broek, 1985) is more likely to be tagged would provide constraints on tagging. It is also possible that what gets tagged might be predicted by a computational text-comprehension model such as the construction-integration model (Kintsch & Welsch, 1991), the capacity-constrained construction integration model (Goldman, Varma, & Côté, 1996), or the landscape model (van den Broek, Ridsen, Fletcher, & Thurlow, 1996) that allows importance to arise from basic text-processing parameters.

Our initial approach was to reason that the core events of a situation are marked for source more than noncore events. This reflects our belief that readers who are motivated to attend to sources during reading manage the task complexity by marking what is important rather than what is not. We further thought that initially mentioned events may benefit from sourcing, whereas interference might affect tagging of later read events. The important constraint guiding these arguments is the assumption that processing limitations require nonexperts to tag only a small subset of both events and details.

Readers who are knowledgeable about a text's source will be more likely to form a documents' model. Similarly, it will be easier to form a document node when a text's source is salient and well elaborated. Even with unelaborated sources a reader could still code the two sources as "the first guy" and "the other guy," but a more semantically elaborated source would enable a reader to build a model of the knowledge and motives of the author, thus

supporting inferences about likely claims and evidence mentioned by the author. A reader would also be more likely to form a documents' model when the task is to learn the story while keeping texts separate and when the texts provide conflicting accounts of partially overlapping events. Characteristics of the reader, such as their knowledge and experience using source information, will influence whether a reader attends to source information during reading as well.

EVIDENCE FOR AN INTEGRATED SOURCE-TAGGED REPRESENTATION

A recent study (Britt et al, 1998) examined some of the factors that might affect whether readers mark sources and, in so doing, tested the relative plausibility of the models just described. First, the study tested the assumption that sourcing would be associated more with core events than with noncore events. Second, it examined two task conditions that we believe affect how likely readers are to build integrated versus source-separated models. Specifically, we assumed that instructions to attend to sources would encourage source tagging compared with more typical comprehension instructions. And we thought that information that discredited one of the sources might discourage the building of an integrated model that lacked source tagging.

The study presented 80 students with well-controlled pairs of texts that, although sharing some events and details, differed in others and expressed different interpretations. Both texts related a narrative version of the Panama Canal story wrapped in either a pro-U.S. or an anti-U.S. argument. We selected 16 complex event propositions (e.g., Panama gained its independence through a rebellion) and two details supporting each complex event (e.g., the rebellion occurred with no bloodshed and it was brief). To test the hypothesis that sourcing accuracy varies with the centrality of the event in the causal-temporal event model, half of the 16 complex events were core events (those events that are critically important to telling a summary of story) and half were noncore events (those events that elaborate the story but can be omitted in an adequate summary of the story). To test readers' knowledge of the source of a particular proposition, we manipulated its *uniqueness*, that is, whether the proposition was contained in both texts (shared) or only one (unique). To control for the number of times a proposition was read—which would otherwise be two times for shared information but just once for unique information—items that were unique

(i.e., to be mentioned only by one author) were mentioned exactly twice by that author and shared items were mentioned exactly once by each author. Each event had two critical details that were also mentioned precisely once or twice depending on whether the event was shared or unique. Whether within or across texts, the two mentions of a target proposition were paraphrases of the abstract proposition. Finally, in order to eliminate differences in memorability due to the specific propositions themselves, the study created two versions of each story. The two versions controlled the assignment of shared versus unique information, so that what was shared in one version was unique in the other. The result of this procedure was to unconfound shared versus unique with specific information within a text.

Groups of subjects participated in four different conditions that manipulated the two factors we expected to matter for sourcing: reading instructions (*sourcing* or *comprehension*) and discrediting information applied to one author (early, late, or never). In the sourcing instructions given to three of the groups, learners were told that they would be given two texts that would present different perspectives on the events and later they would have to indicate which author (if any) mentioned a particular event or detail. In comprehension instructions, given to only one group, nothing was said about sources; rather participants were told to learn the story and to be prepared to be tested on particular events and details. Finally, we also manipulated when, or if, participants were given information discrediting one of the authors (i.e., they were told that the first author was not very credible and should be believed only if the information was given by both authors). Two of the three sourcing instruction groups and the comprehension instruction group were given discrediting information, which always applied to the first text after it was read. The discrediting information was provided early (discredit the first text immediately after reading it), late (discredit the first text after both texts had been read but before writing the essay), and never (no discrediting information). Thus, the four groups were: sourcing instruction with discrediting early, sourcing instruction with discrediting late, sourcing instruction with discrediting never, comprehension instruction with discrediting late.

The procedure can be summarized as follows. All students read the first text; at that point, one group (sourcing-early discredit) was told that the first text was not trustworthy. This condition, we thought, should discourage readers from integrating information across the two texts. Next, for all groups, the first text was removed and the second text was read. After reading the second text, students were given a 1-minute distractor task to reduce text recency effects. Then they were given a list of propositions

(sentences that paraphrased information that was in one or both texts or information that actually was not in the texts) and asked which author or authors, if any, mentioned the information. The test consisted of 24 items from each of these types: first author only, second author only, both authors, neither author (i.e., foils—half that contradicted facts mentioned and half that were true but unstated in the presented texts). Following this multiple-choice recognition test, two groups (sourcing-late and comprehension-late) were given discrediting information. Finally, all students were asked to write a short essay summarizing the events in Panama by including only information that was trustworthy. The final two tasks were the Nelson-Denny test of reading ability and a modified version of the verbal reasoning section of the Law School Admission Test. These were given to gain information about individual differences in our student sample that could be correlated with our performance measures.

Predictions

According to the separate representation model, one would expect readers to be very accurate at identifying an author, given that the content was represented. This would require only a search of the two individual situation models to find the information. Although these sequential searches may require high effort, they should lead to accurate identification. Because there is no integration, this model would predict that unique items would be better identified than shared items. We predict this because two mentions within a text would make the representation for that proposition stronger in the individual representation, thus more likely to be recalled. Without integration, there would be no strengthening of overlapping propositions during the second reading. Because shared items are weaker within a text, and there is no increased salience due to reinstatement of that proposition during the second reading, we predict that any advantage comes from more accurately sourcing unique items over shared items. Additionally, the unique items may have an advantage over shared items if the search is terminated after finding the proposition in the first author's situation model.

In contrast, the mush model predicts that participants will not be very accurate in determining who said what. Although, in this model, readers should do well at recognizing that a proposition was contained in what they read (content information), they should have trouble indicating which author mentioned it (source information) because propositions are not marked for source. Essentially, they should be at chance.

The tag-all model, like the separate representation model, should result in extremely accurate performance because everything is tagged. Although there may be better content recognition for core events, there should not be better source recognition for core events. Core and noncore events should be equally well recognized because both are tagged. Shared items should be sourced better than unique items because the model is integrated. While reading the second text, any propositions already in memory will be marked for the new source. This may increase the likelihood that the reader will at that time mark the proposition for the original text. Thus, both important and unimportant items will both be tagged at this time.

The documents' model should lead to accurate performance primarily on the core event items, because they are tagged. Noncore event items are more likely to not be tagged during reading. As with the tag-all model, the joint sources of shared items should be identified better than the single source of unique items because the reader is building an integrated model. During reading of the second text, information that is recognized as already represented becomes tagged, in effect strengthening source information even for noncore events. Readers can tag shared items that were not previously tagged (e.g., noncore events), thus leaving a predicted difference in the accuracy of sourcing shared items. This leads to the prediction of a specific interaction between centrality and uniqueness: Core events should be sourced better than noncore events, but only for unique items.

We turn now to consider selectively the results that are most important for distinguishing among these four models, emphasizing how they differ in representing source information.

Recognition Results

Content Recognition. Students were able to adequately distinguish presented items from not presented items, suggesting that they encoded and remembered content information from the texts. For items actually from the text, they had a hit rate of 82%; for foil items (neither author) they had a correct rejection rate of 70%, resulting in an overall average d' of 1.64.

Source Recognition. Given that students were able to distinguish presented from nonpresented content, how accurately were they able to recognize the author of the content? To address this question, we looked at responses that identified one or the other or both sources. Overall students

were able to identify the actual author at a better than chance rate, 56% of responses compared to a 33% chance level. This suggests that a model that assumes full integration and no representation of sources (the mush model) cannot be correct. Students were able to identify the source of content they recognized, not perfectly, but rather well relative to the chance level implied by the mush model.

Table 13.1 gives the average conditional percent correct recognition for centrality (core vs. noncore) and uniqueness (unique vs. shared). There was a significant effect of centrality with core events identified more accurately than noncore events. Only the documents' model predicts that more important items are more likely to be marked. There was also a significant effect of uniqueness, with shared items identified more accurately than unique items (63% > 48%). Recall that for every item, the number of mentions was not confounded with whether information was shared or unique. Thus, this result seems to reflect a genuine advantage for information that is shared between the texts over information that is unique to one text. This result argues against the separate representation model, which appears to predict either an advantage for unique information or no difference, depending on what assumptions are made about text procedures. The uniqueness effect, however, does support the documents' model and the tag-all model. Both models predict that during the integration process evoked by the second text, information encountered in the second text that was also contained in the first text will be tagged for both sources. Finally, there was a significant interaction of centrality by uniqueness, specifically, a smaller effect of centrality for shared than unique items. This is exclusively consistent with the documents' model, which predicts that core events are more likely to be tagged, as found with the unique items. For shared information, the noncore events should be noticed during the integration processes and should be tagged at that point. Thus, shared items should not have an appreciable benefit for centrality, as we found.

In summary, only the documents' model is consistent with all the recognition results. Models that allow only integration are ruled out by the

TABLE 13.1

Mean Conditional Percent Correct Recognizing the Author Identity
by Centrality and Uniqueness

	Centrality	Uniqueness	
		Unique	Shared
Core Events		54	64
Noncore Events		42	61

high level of source accuracy. Models that involve sourcing based on full tagging or on separate representations are ruled out by the selective sourcing of core information.

Trustworthy Essays

The recognition data show that students built a representation that included markings for the sources of some of the information. Whether students can use their source-marked representation in a meaningful way, such as writing an essay they consider trustworthy, is another question. By requesting students to write essays that were trustworthy, we are asking them to use information that they believed to be correct. By telling three groups of students that one of their sources was not an expert after all, but a politician who wrote to influence public opinion, we were suggesting that they perhaps should not use information from this discredited source in their essay. However, because the discrediting information always came after students had read the source to be discredited, the options for separating the two sources were limited. For complete success, students would have needed to tag all the information from the second text and then used that exclusively in writing their essay. That option was available to the early discredit group, who received discrediting information prior to reading the second text; however the late discredit group had already read the second text prior to receiving discrediting information. Their only chance to write a selective essay, that is, drawing only from the second and not from the first text, was to have already tagged the sources during reading.

In analyzing the student essays, Britt et al. (1998) made judgments of the source of each proposition contained in a student essay. For our purposes, the most important data are those based on conditional categorizations that took account of the students' judgment of that proposition in the recognition test. For each proposition in the student essay, Britt et al. asked how the student sourced it in the recognition test. For example, for a proposition written by both authors, the proposition was classified as unique if the student indicated on the recognition test that only one of the authors had mentioned it. Then each proposition in the essay could be classified according to the subject's judgment of its source: the discredited text (Clark), the nondiscredited text (LaCosta), or both. The basic result, shown in Table 13.2, is that students who were given discrediting information about Clark wrote essays that contained more information drawn from LaCosta, the trustworthy source, than did the group who received no discrediting infor-

TABLE 13.2

Mean Number of Essay Statements, Classified by Recognition Results, for Each Condition and Text

Condition	Text	
	Clark (Discredited)	LaCosta (Credited)
Sourcing instruction with discrediting early	.35	1.20
Sourcing instruction with discrediting late	.45	1.05
Sourcing instruction with discrediting never	.75	.60
Comprehension instruction with discrediting late	.45	1.00

mation. It did not matter whether the discrediting information came early or late or whether comprehension or sourcing instructions were provided.

INTEGRATION AND SEPARATION: A SUMMARY OF PROGRESS

We have argued that in learning from multiple texts a reader can achieve two goals that at first suggest incompatibility—developing an integrated representation of situations while keeping the source of the information separate. This is accomplished by the reader's acquisition of a documents' model, with the selective tagging of source information that allows integrated situations and partially separated sources.

In the present development of the documents' model, two components interact to affect the overall situational coherence a reader might establish from reading multiple texts. The intertext model, which includes document nodes to represent source information as well as links among sources and between sources and their content, proves especially useful in cases where several possible causal connections exist between the events described by different authors. When readers are confronted with texts that only partially overlap in their situational connections, they cannot compute a coherent representation on the basis of causal connectedness alone. When two authors propose incompatible causes for a single event, the intertext model allows these incompatible situations to be integrated and separated at the same time by source tagging each substory (or event sequence). Source tagging explains how readers maintain overall coherence in a mental model when causal connections are incoherent. The source-content connections override the event-event causal connections.

In general support for the framework of the documents' model is the study by Britt et al. (1998), which found that students were frequently able to identify which source had contained a given proposition. Information shared between the texts was identified more often than information that was unique to one text, supporting the assumption that information is integrated during reading. Supporting the selectivity assumption, and the corollary assumptions that information important to the narrative structure is privilege for selectivity, was the observation that students identified core events more than noncore events. The documents' model predicts both integration and clear but limited source tagging. Finally, this study found that the effect of centrality (core events) was stronger for information that was unique to one text than for information that was shared. We suggest that whereas core events are marked during the first reading, the unmarked events not central to the story are activated during the reading of the second text and marked during a corroboration process. Finally, the study also found that students were able to use this source-tagged structure when writing an essay. They could selectively use information from a source that had not been discredited and avoid information from a source that had been discredited.

The finding that sourcing is controlled by how central information is to a story establishes one element of the documents' model. Determining additional influences on the creation of source-content links requires further research. One clear influence should come from cross-referencing. An author who cites another source and links that source to the situation immediately provides a source link for the student's document model. This, in effect, gives for free what must be earned in learning situations such as the one we explored, where each source ignores all other sources. Other likely influences center on the task. Whereas our study allowed what was important to emerge from the texts themselves, importance can be established explicitly in the task. In the Panama Canal case, students might be given documents on Panama and asked to read them so that they can evaluate what President Roosevelt thought about the situation in Panama. This should strongly encourage tagging Roosevelt information that is relevant to the task. Nor must instructions be so explicit to be effective. Rouet et al. (1996), in requiring students to read documents in order to come to an informed opinion on a controversy, would have implicitly encouraged students to read from the perspective of the controversy, thus making controversial events more important to source.

The results from the Britt et al. (1998) experiment, in conjunction with prior results, show that readers can represent source information and its

connection to content. Clearly, this ability is one that can be cultivated, and indeed it is not a routine ability even for college students. Both reading ability and verbal reasoning ability influenced the accuracy of this source tagging in Britt et al. If students are expected to learn from more than a single textbook excerpt, it helps if they are skilled at reading arguments and considering sources. The problem is that many students are not skilled in this way, and seem to have little experience prior to advanced college courses in such text study.

We are addressing this problem in a project that teaches high school students to attend to document sources using a computer-based learning environment. Some research has suggested that high school students do little spontaneous sourcing and corroboration (Wineburg, 1991), and students' implicit document knowledge is expressed only under some conditions (Britt, Rouet, & Perfetti, 1996; Rouet et al., 1996). The computer environment we have developed, called the Sourcer's Apprentice, presents students with a historical controversy and a set of excerpts from several different types of documents. The students identify different source characteristics (e.g., who the author is, when the document was written, etc.), answer questions about the content of the excerpts, and then write a short argumentative essay about the controversy. Additionally, Sourcer's Apprentice provides explicit tutoring about source characteristics and corroboration strategies and it provides directed practice exercises. We believe that intelligent use of texts must include paying attention to source information during reading and that training in the use and evaluation of sources will help students deal more effectively in multiple-text situations.

CONCLUSION

In a sense, we have come full circle in studying history text learning. We began by assuming and then verifying that learning from multiple history texts was, for the nonexpert, a question of learning a story. And story learning was about learning temporal-causal event structures. When we turned to developing a more elaborate account of the kinds of information readers could acquire, the idea of sources linked to situations was a natural extension. Now, in examining what students actually learn about sources when they read may take us back to the narrative itself. The students in the Britt et al. (1998) study appeared to selectively source documents in accord with the documents' representation model (Perfetti et al., 1999). Core information appears to be privileged in sourcing, just as it is in narrative

history learning (Perfetti et al., 1995). Students appear most likely to mark sources for elements that are central to the story.

For purposes of history learning—at both the narrative and the source levels—we have renewed reason to appreciate the need for a rigorous accounting of causal structures of the kind developed by Trabasso and colleagues. In terms of the documents' model, whether we examine the situations model component or the intertext model component, we find evidence that learners take account of what is "important" in a particular way: Events that are core to the temporal-causal structure of the story are learned and if different sources create differences in this structure, that fact also will be noticed. Additional questions about what these circumstances are and how they control the learning situation will presumably take us far from temporal-causal analysis, but at first look, there is reason to believe that the narrative structure not only provides the form of what the reader learns about the situation, but also provides the platform from which to erect links to sources.

REFERENCES

- Britt, M. A., Rouet, J.-F., & Perfetti, C. A. (1996). Using hypertext to study and reason about historical evidence. In J.-F. Rouet, J. J. Levenon, A. P. Dillon, & R. J. Spiro (Eds.), *Hypertext and cognition* (pp. 43–72). Mahwah, NJ: Lawrence Erlbaum Associates.
- Britt, M. A., Sandak, R. L., Perfetti, C. A., & Rouet, J.-F. (July, 1998). *Content integration and source separation in learning from multiple texts*. Paper presented at the Annual meeting of the Society for Text and Discourse. Madison, WI.
- Goldman, S. R., Varma, S., & Côté, N. (1996). Extending capacity-constrained construction integration: Toward "smarter" and flexible models of text comprehension. In B. K. Britton & A. C. Graesser (Eds.), *Models of understanding text* (pp. 73–113). Mahwah, NJ: Lawrence Erlbaum Associates.
- Kintsch, W., & Welsch, D. M. (1991). The construction-integration model: A framework for studying memory for text. In W. E. Hockley & S. Lewandowsky (Eds.), *Relating theory and data: Essays on human memory* (pp. 367–385). Hillsdale, NJ: Lawrence Erlbaum Associates.
- Kushner, J. M. (1996). *The role of author information in reading*. Unpublished doctoral dissertation, University of Pittsburgh, Pittsburgh, PA.
- Mandler, J. M., & Johnson, N. S. (1977). Remembrance of things parsed: Story structure and recall. *Cognitive Psychology*, 9, 111–151.
- Omanon, R. C. (1982). The relation between centrality and story category variation. *Journal of Verbal Learning and Verbal Behavior*, 21, 326–337.
- Perfetti, C. A., Britt, M. A., & Georgi, M. C. (1995). *Text-based learning and reasoning: Studies in history*. Hillsdale, NJ: Lawrence Erlbaum Associates.
- Perfetti, C. A., Rouet, J.-F., & Britt, M. A. (1999). Toward a theory of documents representation. In H. van Oostendorp & S. R. Goldman (Eds.), *The construction of mental representations during reading*. Mahwah, NJ: Lawrence Erlbaum Associates.
- Rouet, J.-F., Britt, M. A., Mason, R. A., & Perfetti, C. A. (1996). Using multiple sources of evidence to reason about history. *Journal of Educational Psychology*, 88, 478–493.
- Rumelhart, D. E. (1975). Notes on a schema for stories. In D. G. Bobrow & A. Collins (Eds.), *Representation and understanding: Studies in cognitive science* (pp. 211–236). New York: Academic Press.
- Stein, N. L., & Glenn, C. G. (1979). An analysis of story comprehension in elementary school children. In R. O. Freedle (Ed.), *New directions in discourse processing* (pp. 53–120). Hillsdale, NJ: Lawrence Erlbaum Associates.
- Thomdyke, P. W. (1977). Cognitive structures in comprehension and memory of narrative discourse. *Cognitive Psychology*, 9, 77–110.
- Trabasso, T., & Sperry, L. (1985). Causal relatedness and importance of story events. *Journal of Memory and Language*, 24, 595–611.
- Trabasso, T., & van den Broek, P. (1985). Causal thinking and the representation of narrative events. *Journal of Memory and Language*, 24, 612–630.
- van den Broek, P., & Trabasso, T. (1986). Causal network versus goal hierarchies in summarizing text. *Discourse Processes*, 9, 1–15.
- van den Broek, P., Ridsen, K., Fletcher, C. R., & Thurlow, R. (1996). A "landscape" view of reading: Fluctuating patterns of activation and the construction of a stable memory representation. In B. K. Britton & A. C. Graesser (Eds.), *Models of understanding text* (pp. 165–187). Mahwah, NJ: Lawrence Erlbaum Associates.
- Voss, J. F., & Wiley, J. (1997). Developing understanding while writing essays in history. *International Journal of Educational Research*, 27, 255–265.
- Wiley, J., & Voss, J. F. (1996). The effects of playing historian on learning in history. *Applied Cognitive Psychology*, 10, 63–72.
- Wineburg, S. S. (1991). Historical problem solving: A study of the cognitive processes used in the evaluation of documentary and pictorial evidence. *Journal of Educational Psychology*, 83, 73–87.