



CREATE LINKS TAB

Qualitative causal mapping involves taking passages of text, e.g. from interviews or documents, and identifying sections which make causal claims. We highlight each of these sections and specify a causal factor at each end of each link (for example Lost job or Went hungry). This means creating a new factor or reusing an existing one. Usually we create these factors inductively as we code, and revise and review and consolidate them as part of the process, as with any other kind of qualitative content analysis.

To code a causal link,

- With your mouse, highlight a piece of text within the statement which makes a causal claim. Your selection must remain within one statement and must not cross into another statement.
- Watch how that passage is copied for you into the "Quote" window. (Usually, you don't need to think about this window: you can edit the text if you really need to but it **has to remain an exact quote of one part of the text.**)
- Start to type the name of the influence factors at the **start** of the link(s) which you are going to make, in the first drop-down menu.
- If there is an existing factor which matches what you want, you can select it.
- Otherwise, you will create a new factor with the contents of what you have typed; finish what you have typed with a comma or a tab character if you want to continue to select or create another factor.
- If you want to create more than one link, you can select or create additional factors in the same box (as shown in the video below).
- When you have finished, press Enter.
- Repeat the process in the other box to specify the factors at the **end** of the link (or ends of the links).
- Press the green Save button which is now active.
- The link is created in the Map window.
- When you have finished coding one source, click the right arrow in the source navigator to code the next source.

Source Text Viewer

🌟 **What you can do here:** Read your source documents and create causal links by highlighting text. When you highlight a passage that claims or implies that one thing influences or causes

another, a popup lets you identify the cause and effect. This is where you do the core work of mapping out causal relationships from your source material, a process which we call *coding* The text viewer shows full text from the selected source. If you have selected multiple sources, it shows the text from the first selected source. You can:

- Highlight sections to identify causal claims. Highlighting opens the causal link editor.
- Examine and edit existing highlights by clicking on them.

Inside the header, there is an info  icon which toggles open/shut a panel beneath it which shows the values of the custom columns for the current source e.g. gender etc.

Navigation Controls

Navigate sources:

-  **Previous source**
-  **Next source**

Navigate highlights within the current source:

-  **First highlight in source**
-  **Previous highlight**
-  **Next highlight**
-  **Last highlight in source** (useful if you haven't finished coding the whole text yet and want to see the last highlight)

Source selection is filtered through the sources selector dropdown. When multiple sources are loaded, the first source is displayed. The next/previous buttons cycle through sources by updating the sources selector to show the next/previous source.

This is convenient because usually when coding you will want to view the [Map](#) and [Links](#) for the same source on the right.

Clicking these buttons means that if you previously had a multiple selection, you now have only one.

Dealing with long documents in the source text viewer

For documents longer than ~30-40 pages, the text viewer automatically splits content into manageable chunks for better performance. Navigation controls appear in the "Source text" header:

-  **Chunk selector** (dropdown): e.g. "Chunk 1 of 5" — jump directly to a chunk.

-  **Previous / Next chunk** (arrow buttons): move between chunks.
-  **Next chunk** (button at end of chunk): appears at the end of each chunk (except the last).

Visual Highlighting

Each section of coded text, each causal claim, is shown with a highlight.

For overlapping or identical highlights with multiple links, overlaps are shown with varying color opacity. Clicking on multiple highlights shows a link selector for each section.

- Multiple highlights shown with varying color opacity
- Click on overlapping highlights to select specific links

Link Editor screen

Opens when you highlight text or click on existing links.

Fields:

-  **Cause selector** (multi-select dropdown + free text): choose one or more causes (project-wide list, sorted by frequency).
-  **Effect selector** (multi-select dropdown + free text): choose one or more effects (same project-wide list). After you select Causes, the Effect dropdown boosts the most common effects of those causes to the top.
-  **Quote** (text area): editable evidence text; supports ellipses like *Actual quote [this text is ignored] quote continues....*
-  **Chain** (toggle): if on, saving keeps the editor open and uses the previous Effect as the next Cause.
-  **Plain coding** (toggle): if on, saves a self-loop (Cause = Effect) tagged *#plain_coding* so it counts as “theme present” rather than a causal claim. Plain codings can be hidden with the *Exclude self-loops* filter.
-  **Tags** (multi-select input): add tags like *#hypothetical* or *check*.
-  **Favourites** (3 toggle buttons): heart / exclamation / star.

Actions:

-  **Save** (button): create/update link(s).
-  **Delete** (button): remove an existing link.
-  **Cancel** (button): close without saving.

Links in Causal Map only have one cause and one effect. You can add multiple causes and/or effects to the boxes, and the system creates all combinations when saving. So if you put

unemployment and violence in the Cause box, and stress and worry in the Effect box, the system will create four links.

About the factor label dropdown menus

By creating links, you also create the names of your factors.

In Causal Map, a factor is its label. Once you create a label, there is nothing else to add.

Factor names which contain semicolons ; get special treatment as they separate the different parts of 📌 Hierarchical factors .

After beginning to create links between factors, already-coded factors will appear in the dropdown menus in the to and from factor boxes. For added convenience. The most frequently coded factors will appear at the top of this list

#doubtful? #hypothetical? Adding link tags

Link tags

Link tags are available as a special kind of memo when coding a link: you can use them to provide any kind of additional information.

There is no need to actually use a hash # at the start of a link tag, though you can if you want. Just use any unique single word which is easy to search and filter on, like #nutrition or nutrition# or nutrition—.

As usual in Causal Map, you can apply one or more tags, and you can either select existing tags or create new ones on the fly.

Later, you can filter the map (see 🌟 Transforms Filters: Include or exclude tags) to show only links containing or beginning or ending with specific hashtags (or parts of hashtags), and also for links which *do not* contain specific hashtags or parts of hashtags.

You can also use tags to narrow down your searches in 🔗 The Manage Links tab.

You can display tags on your map.

Conceptually, there are two kinds of tag.

Ordinary link tags

You can use any tag which does not begin with a ? to record any other information about the link, e.g.:

- respondent doesn't like this connection
- respondent feels good about the outcome
- for you, the analyst, e.g.
 - respondent is answering a different question
 - to tag links you want to come back and review.

Weak tags

Weak tags are a special kind of tag. They are *caveats*. If you use weak tags, you should make sure that by default your maps do not include any link with a weak tag.

This is just a convention, it makes no difference to the Causal Map app.

They begin with **?** and are used to mark any link which you are not sure is always valid across the global context for the whole global map, for example:

- **the causal connection** is only valid for a specific context, e.g.
 - the respondent says this is only true for their village, not for other villages e.g. **?village X**
 - a link is only projected for the future e.g. **?future**
- you are unsure about **the claim about the causal connection**
 - a link is only a hypothesis e.g. **?hypothetical**
 - you as the analyst are not confident in the claim e.g. **?doubtful**
 - the source themselves are not sure e.g. **?source seems unsure**
 - to add other qualifying information e.g. **?probably hearsay**
 - to mark the fact that a connection is **weak or non-existent**, e.g.
 - Respondent makes a substantive claim that X does *not* influence Y, e.g. **?zero influence**
 - Respondent makes a substantive claim that X only insignificantly influences Y, e.g. **?weak**

AI Coding

Requires AI subscription

- **Model dropdown** - Select AI model
- **Prompt box** - Enter coding instructions
- **Add source prompt** - Toggle, default ON
- **Response displays** - View AI responses and full JSON

Additional controls hidden behind gear icon (experimental):

- **Temperature slider** - Control randomness (default 0)

Iterative Processing:

Normally you can put all your tasks to the AI into one iteration. However sometimes you might need to add another task which is best done once the entire original table has been created, for example:

```
Give me a links table where each row is a causal claim, with columns for Cause, Effect and Quote
```

```
====
```

```
Now look at the table and return the same table but with an extra column `Importance` with an integer with which you rank the links according to importance, with 1 being most important.
```

If your prompt contains lines with "====" on their own, each section before and after the line is treated as a separate iteration. . First iteration is normal; subsequent ones include the full prior conversation history (all previous User prompts and AI responses) to build on earlier results. Only the results of the last iteration are added to the links table; all iterations are logged in the responses panel.

Source prompt: it is just to describe the context/background info about each source. Not necessary e.g. where all the sources are from the same context which can be described in the main prompt. But important where some differ, e.g. mid-term reports or whatever.

If Add Source Prompt is ON, then show a text area above #text-viewer-content with usual greenish Save button to edit the corresponding source prompt for the current source.

Workflow:

- Select one (or more) sources to process using the sources dropdown
- Select "Skip coded sources" if you don't want to recode sources which have already been coded
- Toggle "Add source prompt" to append the new Source Prompt field before the beginning of the main prompt
- Click "Process Sources" button
- Confirmation dialog shows model, word count, and warnings
- AI processes sources in batches
- Results are also logged to the responses table on the right of the screen

[Tips on using the prompt history](#)

- Timeouts: per-iteration budget scales by model and iteration count (cap 540s total): Flash 120s/iteration; Pro 270s/iteration.
- Concurrency: Radio group labeled "Concurrency" (1–5) next to Region in AI settings. Default 1. Increase if you want faster processing but may risk 429/timeouts.
- Logging & Responses: each chunk inserts a pending row in [ai_logs](#) (status pending) and updates to success/error on completion; Responses tab auto-refreshes as rows update.