

Structured Data Project | Seer

Work by [Jake Labate](#), Technical SEO Manager Candidate.

Hi Kelly, Kim, Wil and the broader Seer Team,

This is a 1-off gift I'd like to present to the Seer team.

The Deliverable - An enhanced/upgraded structure data (schema markup) profile for a number of pages on the Seer website.

Deliverable Links:

- [Seer Interactive Scrape Framework](#)
- [Seer Interactive Schema Framework](#)
- [Final Schema \(.zip file\)](#)
 - [Example schema for Wil's team member profile](#)

The Current Schema - The existing schema layout on the 3 chosen focus page types:

- [Insights Schema](#)
- [Case Study Schema](#)
- No Schema on the Team Member Pages

How is it “enhanced/upgraded”?

- 1. Additional [schema.org](#) @types:** The 3 focus page types included either very limited or no schema at all. The new framework includes a wider range of schema @types that follow the quality standards (outlined later). This can help trigger more rich result types in Google and other search engines.
- 2. Additional [schema.org](#) properties:** In addition to more schema @types, the framework includes a range of new properties (nested inside both existing and new schema @types). This can help both trigger rich results (as search engines favor additional populated property values on each schema @type), but also show more information in any triggered rich result (providing additional content to searchers / taking up more real estate in the SERP).
- 3. Addition of the @id key on each and every object:** Every single schema @type object in the framework is assigned its own @id value. This allows

benefit #4 (Cross object linking / referencing), but also grants Google the ability to evaluate each schema object separately. This helps search engines to better identify the extent of the objects.

4. **Cross object linking / referencing:** In addition to implementing more extensive object nesting (a @type within a @type) allowing search engines to infer object relations, the @id feature also allows cross-@type, cross-page-type and cross-site referencing. When search engines attempt to understand a webpage's schema, they use the rest of the site's schema as context.

Implementation Methodology

Data Collection

Of course a prerequisite to deploying structured data is the data / values / content / info one would use in order to populate any @type / property framework. Ex: We would need to have the actual first and last name of a team member in order to populate the following "name" property in a "Person" @type:

```
{
  "@type": "Person",
  "name": {{ name }} <- this value
}
```

So, I created a series of custom web scrape frameworks (in JavaScript) in order to capture the data on the Seer website.

While an ideal future-proof solution (should the page design change) for capturing data on the site would be pulling this data directly from a CMS / database (which Seer might currently use), this method worked for my use case (collecting the data) in order to then pass the data to a framework in the construction of this project.

The framework scrapes basic URL, <h1> tag, title tag, meta description, open graph tags, twitter tags, and the current structured data for all pages, as well as additional data for specific pages off certain folders/paths (more on this below).

[Seer Interactive Scrape Framework](#)

In addition to what was scraped off the pages with this framework, various information about Seer was manually captured, such as the HQ address, social media profiles, etc.

Page Prioritization

As expected with any site, on the Seer website there is a number of 1-off URLs (URLs that are not part of a greater URL path shared with other URLs), for example:

www.seerinteractive.com/

- [insights/events-calendar/lets-talk-about-ga4](#)
- [insights/analyzing-ai-overview-data-w/-paid-conversions-finding-when-to-care](#)
- [work/analytics/google-analytics-4-migration](#)

However, most of the URLs on the site (of the total of 1839 in [the sitemap](#)), do fall into a few shared paths, where the page layouts, and content share a similar structure. The most common folder paths being:

https://www.seerinteractive.com/

- [insights/{{ slug }}](#)
- [work/case-studies/{{ slug }}](#)
- [people/team/{{ slug }}](#)

Because structured data is primarily a SERP asset (as the primary benefit comes from triggering rich results in the SERP) it's likely that a coherent structured data framework effort for any site would best attend to the pages acquiring the most organic traffic.

However, without my access to this information, to disproportionately aid in Seer's SERP appearance through structured data, I created enhanced / upgraded profiles for the pages in the most common 3 path types - the 3 types being: Insight pages, Case Study pages and Team Member pages.

Framework Development

The schema framework was developed so that for every page, the schema matches the following **quality standards**:

1. **Valid schema.org vocabulary / syntax:** All pages went through rigorous testing and were refined until there were 0 errors and 0 warnings in both the industry standard Google Rich Result Test and the Schema Validator.
2. **Indicative of page and/or page content:** The schema @types and properties (and their populated values) as well as their mapped relational hierarchy are indicative of either the placement of the page in the website (ex: BreadcrumbList), or precisely relevant to the page content (ex: a Person @type on a team member page).
3. **Accurate:** Schema was constructed based on several manual checks for both how the data is captured, and populated in the framework. The data populated in the final framework is reliably the same data that is on each page.

The schema.org vocabulary includes hundreds of schema @types, and thousands of properties. Instead of considering all of these we may populate, a focus may be placed in developing a framework for the schema @types and properties which [Google says they specifically support](#) in the creation of rich results.

While other properties may provide additional value for other search engines, as well as Google (which they just might not tell us), it makes sense to focus on these few.

However, some schema @types and properties **heavily align with relevant information** on the pages, but are not supported according to Google (such as a "Person" @type on the /people/team/ pages). As a comprehensive practice, some of these were included in the framework.

FYI: I was sure to capture all existing schema @types and properties for pages, and either identically re-used them in the new schema framework, or wrap the same values in the new upgraded framework. There was only 1 alteration to an existing schema property **value**:

- I used "Seer Interactive" instead of "New Equation, LLC DBA Seer Interactive" for the BlogPosting publisher name. I did this so that I might link the full (and now reusable) Organization @type object as the publisher value.

[Seer Interactive Schema Framework](#)

I've implemented this type of dynamic structured data framework for Hooray Agency's hotel clients. Because Hooray served only 1 type of client (hotels), the structured data framework was similar from client-to-client. I combined these frameworks into a [master hotel schema framework](#), which the agency currently uses for their premium SEO clients.

Hotels Using This: (see validation test links in the dev console - on all pages):

- [Global Ambassador Hotel](#) - (see [structured data validation](#))
- [Hotel Bennett](#) - (see [structured data validation](#))
- [Maui Kā'anapali Villas](#) - (see [structured data validation](#))