

OVERALL SCORE

75 /100

AI-Readable

Executive summary

This site is mostly readable by AI tools, with a few improvements still recommended. Key strengths include homepage access and AI guidance file, while plain-text page access and crawler policy need attention. Recommended next step: add content negotiation for Accept: text/markdown on the homepage and return a markdown representation with Content-Type: text/markdown. Keep the HTML response for regular browser requests.

Recommended next step

- 1. Add content negotiation for Accept: text/markdown on the homepage and return a markdown representation with Content-Type: text/markdown. Keep the HTML response for regular browser requests.
- 2. Fix robots.txt syntax issues so each rule uses Field: value format, directives appear under a User-agent, and Sitemap entries use absolute URLs.
- 3. Add a meta description between 50 and 160 characters. Add missing semantic elements: article.

Signal breakdown

Crawlability The homepage resolves, connects, returns HTTP 200 OK, exposes a canonical URL, and is not blocked by robots.txt.	Pass	20/20
Robots.txt Fix robots.txt syntax issues so each rule uses Field: value format, directives appear under a User-agent, and Sitemap entries use absolute URLs.	Warn	7.5/15
llms.txt The llms.txt file was found and includes the expected text, length, heading, and URL signals.	Pass	15/15

Sitemap	Pass	10/10
The sitemap index is valid and the checked child sitemap contains URL entries.		
Markdown support	Fail	0/15
Add content negotiation for Accept: text/markdown on the homepage and return a markdown representation with Content-Type: text/markdown. Keep the HTML response for regular browser requests.		
Semantic HTML	Warn	7.1/10
Add a meta description between 50 and 160 characters. Add missing semantic elements: article.		
Structured data	Pass	10/10
Valid JSON-LD structured data was found with core Organization or WebSite schema.org types.		
Content signals	Pass	5/5
Consider adding Content-Signal HTTP header, AI-specific head meta tags, robots noai/noimageai directive so AI systems can consistently discover content usage preferences across robots.txt, HTTP headers, and HTML metadata.		

Suggested fixes

```
llms.txt
MARKDOWN
# Record Store 01

> Record Store 01

This llms.txt file summarizes the public, canonical resources that AI
assistants and crawlers should use to understand this site.

## Site Overview

- Canonical URL: https://recordstore01.myshopify.com/
- Site type: organization
- Recommended summary: Record Store 01

## Core URLs

- [Homepage](https://recordstore01.myshopify.com/): Primary public entry
point and canonical site overview.
- [Catalog](https://recordstore01.myshopify.com/collections/all): Important
public page discovered from the homepage navigation.
- [Collection](https://recordstore01.myshopify.com/collections): Important
public page discovered from the homepage navigation.
- [Glasses](https://recordstore01.myshopify.com/collections/glasses):
Important public page discovered from the homepage navigation.
-
[Blog](https://recordstore01.myshopify.com/blogs/discover-the-joy-of-vinyl-r
ecords-554): Continued in the full scan report...
```

robots.txt additions

TXT

```
# robots.txt additions
# Copy these blocks into the existing robots.txt file. Keep current rules
unless a note calls out a conflicting Disallow.

# Sitemap discovery
# Add this top-level directive so search and AI crawlers can discover the
canonical sitemap.
Sitemap: https://recordstore01.myshopify.com/sitemap.xml

# AI crawler access
# Add explicit Allow rules for AI crawlers so public content permissions are
clear.
User-agent: GPTBot
Allow: /

User-agent: ChatGPT-User
Allow: /

User-agent: ClaudeBot
Allow: /

User-agent: Claude-Web
Allow: /

User-agent: PerplexityBot
Allow: /
Continued in the full scan report...
```

schema.json

JSON

```
{
"@context": "https://schema.org",
"@graph": [
{
"@type": "Organization",
"@id": "https://recordstore01.myshopify.com/#organization",
"name": "Record Store 01",
"description": "Record Store 01",
"url": "https://recordstore01.myshopify.com/"
},
{
"@type": "WebSite",
"@id": "https://recordstore01.myshopify.com/#website",
"name": "Record Store 01",
"description": "Record Store 01",
"url": "https://recordstore01.myshopify.com/",
"publisher": {
"@id": "https://recordstore01.myshopify.com/#organization"
},
"inLanguage": "en",
"potentialAction": {
"@type": "SearchAction",
"target": {
"@type": "EntryPoint",
"urlTemplate":
"https://recordstore01.myshopify.com/search?q=%7Bsearch_term_string%7D"
}
}
}
]
}
Continued in the full scan report...
```

Content-Signal

TXT

Content-Signal recommendations

Use these directives to make AI-use preferences explicit for compliant crawlers and AI systems. They are advisory signals, so keep them aligned with robots.txt, terms, and access controls.

Recommended values

- ai-train=no: AI model training, fine-tuning, and dataset creation.
- search=yes: AI search indexing, snippets, and discovery.
- ai-input=yes: AI answer grounding, retrieval, and generated-response context.

HTTP response header

```
'''http
Content-Signal: ai-train=no, search=yes, ai-input=yes
Content-Usage: train-ai=n
'''
```

Best for site-wide or route-specific policies because the signal travels with every response, including pages that AI systems fetch directly.

HTML meta tag alternatives

Add these inside the document '`<head>`' when server headers are not

continued in the full scan report...

Full report

<https://llmscan.dev/scan/18F2MXvYbhSF1ABFxy18U>