

OVERALL SCORE

15 /100

Poor

Executive summary

This site is difficult for AI tools to read right now. Key strengths include sitemap and content signals, while homepage access and crawler policy need attention. Recommended next step: remove AI crawler Disallow: / rules or replace them with narrower path-level restrictions for private content only.

Recommended next step

1. Remove AI crawler Disallow: / rules or replace them with narrower path-level restrictions for private content only.
2. Remove AI crawler Disallow: / rules or add narrower Allow/Disallow rules if AI crawlers should be able to discover public content.
3. Publish /llms.txt as text or markdown with more than 200 characters, markdown headings, and at least one absolute URL.

Signal breakdown

Crawlability

Fail0/20

Remove AI crawler Disallow: / rules or replace them with narrower path-level restrictions for private content only.

Robots.txt

Fail0/15

Remove AI crawler Disallow: / rules or add narrower Allow/Disallow rules if AI crawlers should be able to discover public content.

llms.txt

Fail0/15

Publish /llms.txt as text or markdown with more than 200 characters, markdown headings, and at least one absolute URL.

Sitemap	Pass	10/10
The sitemap.xml file is valid and 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	Fail	0/10
Shorten the meta description to 160 characters or fewer. Add exactly one h1 element that describes the page topic. Add missing semantic elements: main, article, nav, footer. Add visible body copy after script and style removal so the		
Structured data	Fail	0/10
Add JSON-LD structured data with Organization or WebSite schema so AI systems can identify the site owner or website entity.		
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

Fix Structured data

HTML

```
<script type="application/ld+json">
{
  "@context": "https://schema.org",
  "@graph": [
    {
      "@type": "Organization",
      "@id": "https://ruleworld.co/#organization",
      "name": "RuleWorld - own a country on the world map",
      "description": "Buy a country and put your logo, name and link on the world map. Anyone can take it from you for more than you paid - and half of what they pay is split between everyone who owned it before.",
      "url": "https://ruleworld.co/",
      "logo": "https://ruleworld.co/icon-180.png"
    },
    {
      "@type": "WebSite",
      "@id": "https://ruleworld.co/#website",
      "name": "RuleWorld - own a country on the world map",
      "description": "Buy a country and put your logo, name and link on the world map. Anyone can take it from you for more than you paid - and half of what they pay is split between everyone who owned it before.",
      "url": "https://ruleworld.co/",
      "publisher": {
        "@id": "https://ruleworld.co/#organization"
      },
      "inLanguage": "en"
    }
  ]
}
```

Continued in the full scan report...

Full report

<https://llmscan.dev/scan/buKKPssQ-4ak1YuPW9-Mg>