

LLM Scan

PUBLIC AI VISIBILITY REPORT

melis.website

Scanned Aug 30, 2026, 17:30 UTC

OVERALL SCORE

35 /100 Poor

Executive summary

This site is difficult for AI tools to read right now. Key strengths include sitemap and page structure, 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 Fail 0/20

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

Robots.txt Fail 0/15

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

llms.txt Fail 0/15

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	Pass	10/10
The homepage has a valid title, meta description, heading structure, semantic landmarks, substantial body copy, and descriptive anchor text.		
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

Melis Guclu | Software Engineer

> Software engineer working across AI/LLM systems, deep learning, full-stack web and mobile development, cloud infrastructure, and native iOS/macOS applications.

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://melis.website/
- Site type: person
- Recommended summary: Software engineer working across AI/LLM systems, deep learning, full-stack web and mobile development, cloud infrastructure, and native iOS/macOS applications.

Core URLs

- [Homepage](https://melis.website/): Primary public entry point and canonical site overview.
- github.com/melisguclu: Important public page discovered from the homepage navigation.
- linkedin.com/in/melis-guclu: Important public page discovered from the homepage navigation.
- [Covotiv GmbH](https://www.covotiv.com/): Important public page discovered

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.

AI crawler access

Add explicit Allow rules for blocked AI crawlers; remove or narrow conflicting Disallow rules if your crawler target requires precedence.

User-agent: GPTBot

Allow: /

User-agent: ClaudeBot

Allow: /

User-agent: Google-Extended

Allow: /

Crawl-rate guidance

Add Crawl-delay under a User-agent block when the site needs gentler crawler pacing.

User-agent: *

Crawl-delay: 10

schema.json

```
JSON
{
  "@context": "https://schema.org",
  "@graph": [
    {
      "@type": "Organization",
      "@id": "https://melis.website/#organization",
      "name": "Melis Guclu | Software Engineer",
      "description": "Software engineer working across AI/LLM systems, deep learning, full-stack web and mobile development, cloud infrastructure, and native iOS/macOS applications.",
      "url": "https://melis.website/",
      "logo": "https://melis.website/apple-touch-icon.png",
      "sameAs": [
        "https://github.com/melisguclu",
        "https://linkedin.com/in/melis-guclu",
        "https://github.com/melisguclu/dockyard",
        "https://github.com/melisguclu/Othello-AI"
      ]
    },
    {
      "@type": "WebSite",
      "@id": "https://melis.website/#website",
      "name": "Melis Guclu | Software Engineer",
      "description": "Software engineer working across AI/LLM systems, deep learning, full-stack web and mobile development, cloud infrastructure, and native iOS/macOS applications."
    }
  ]
}
```

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
```

Full report

<https://llmscan.dev/scan/HEmsLuzHSaHB5Bn4o6jGg>