



The MES AI Visibility Index

Which manufacturing execution vendors AI answer engines actually name, and why

112 prompts · 8 buyer personas · 3 answer engines · 336 responses

AEO Wrangler B2B Software Index · September 2026

Executive summary

We asked ChatGPT, Gemini, and Perplexity 112 buying questions about manufacturing execution software, spanning eight distinct buyer personas in discrete manufacturing. No prompt named a vendor. The vendor list in this report is entirely an output of what the models said, never an input to what we asked.

Across 336 responses, the models named 134 distinct software products and cited 4,777 URLs from 847 domains.

Seven findings stand out.

Siemens Opcenter leads by a wide margin. Opcenter is named in 83 of 112 prompts and appears in all eight personas on all three engines. Rockwell, counting Plex and FactoryTalk together, reaches 70. No other vendor clears 50 percent coverage on more than one engine. Under every weighting scheme tested, Opcenter ranks first by a margin no methodological choice threatens.

The three engines are not measuring the same market. Critical Manufacturing is named in 52.7 percent of ChatGPT prompts and 17.9 percent of Perplexity prompts. iFactory, a product most people in this category have never heard of, appears in 12.5 percent of Perplexity prompts and zero percent of ChatGPT prompts. Treating the engines as one number hides the most actionable information in the dataset.

The three engines draw on structurally different evidence. Counting both vendor categories, 82.0 percent of ChatGPT's citations were vendor domains, against 46.5 percent for Gemini and 35.4 percent for Perplexity. But raw composition understates the difference, because the engines cite at wildly different volumes: 8.1 URLs per prompt for ChatGPT, 3.0 for Gemini, 31.6 for Perplexity. ChatGPT reads vendors closely. Gemini barely cites anything. Perplexity reads everything, mostly third parties.

Visibility comes from two separate mechanisms, and the best position holds both. When ChatGPT names a product, 68 percent of the time it also cites that product's own site. On Perplexity the figure is 24 percent and on Gemini 9 percent, the latter driven almost entirely by Gemini's low citation volume rather than by any preference for third parties. Products can be visible because the model learned them or because content exists and gets retrieved. Siemens is the only vendor strong on both.

Content strategy demonstrably moves visibility, and it moves the retrieval engines first. iFactory earned 160 Perplexity citations to its own domain and converted that into 15 prompt appearances, with 93 percent of its mentions backed by a citation to its own site. Factory AI, Connektica, Fabrico, Epsilon3, and SYMESTIC show the same signature. None registers on ChatGPT yet. In AEO Wrangler client work we have repeatedly observed Perplexity and Gemini visibility arriving months ahead of ChatGPT visibility for the same content program. Retrieval-engine performance may therefore be an early signal of where the other engines land later.

Persona segmentation produces genuinely different leaderboards. MasterControl ranks 10th overall and 1st in medical device. iBase-t ranks 12th overall and 2nd in aerospace. Applied Materials ranks 11th overall and 2nd in semiconductor. First Resonance ranks 13th overall and 2nd in hard tech. Aegis ranks 7th overall and 1st in electronics. A single aggregate ranking buries every one of these.

Adjacent categories intrude heavily on shop floor questions. Only 62 percent of product mentions were core MES or MOM products. ERP took 17.6 percent and PLM 6.7 percent. Siemens Teamcenter, a PLM system, ranks 8th on the combined leaderboard. SAP S/4HANA ranks 7th. Buyers asking AI about shop floor execution are pointed at systems of record roughly a quarter of the time.

1. Why this index exists, and what it does not measure

Every vendor in this category now has a marketing team asking the same question: when a buyer asks an AI assistant which manufacturing execution system to buy, does our name come up? Until now the answers have been anecdotal. Someone tries a prompt, screenshots the result, and draws a conclusion from a sample of one.

This index replaces that with a fixed, disclosed, repeatable instrument.

What it measures is visibility. Whether a product is named, how often, by which engine, and to which buyer. It is a share of voice study.

What it does not measure is quality, fit, suitability, or market position. Being named frequently is not evidence a product is good. Being absent is not evidence a product is bad, and it is not evidence a product is losing.

That distinction is worth making concrete. Manufacturo was named in 2 of 112 prompts, appearing once on Gemini and once on Perplexity, with two citations to its own domain in the entire dataset. It competes directly against First Resonance, which was named in 17 prompts. Whatever is happening in the market is not what is happening in these answers. Read every low placement in this report as a statement about AI visibility and nothing else.

Readers wanting a fit-based evaluation of these tools should look to sources that do that work directly, such as Michael Finocchiaro's MES buyer's guide at Demystifying PLM, which appears in our source data as one of the more heavily cited third-party references in the category.

Editorial independence: AEO Wrangler has taken no vendor money in this category, has no paid placements in this index, and offers none. No vendor saw this report or its methodology before publication.

2. Method

Category boundary. Software that executes and governs production on the discrete shop floor: work order and routing execution, digital work instructions, quality and nonconformance, traceability and genealogy, equipment and machine data, with scheduling, inventory, and maintenance in scope where delivered natively. A product is excluded from the core category if it serves as the financial system of record. This system-of-record test is what separates Rockwell Plex, which we count as core, from SAP S/4HANA, which we count as adjacent ERP.

Scope is discrete manufacturing, including regulated discrete such as aerospace, medical device, and electronics. Process and batch manufacturing are out of scope.

Consolidation rule. Where one vendor sells more than one product inside the core category, those products are counted as one entry. This affects two vendors. Rockwell Plex and Rockwell FactoryTalk, including PharmaSuite and MedDeviceSuite, are reported as Rockwell. Dassault DELMIA Apriso and DELMIAworks are reported as Dassault DELMIA. Sub-brands and legacy names already roll up, so Camstar and SIMATIC IT count as Siemens Opcenter. Products in different categories are not consolidated, which is why Siemens Teamcenter and SAP S/4HANA remain separate adjacent entries rather than merging into their vendors' core products. Per-product figures for both consolidated vendors appear in Appendix D.

Personas. Eight, chosen because they carry materially different vendor sets rather than because they are convenient market segments.

Code	Persona	Buyer
HT	Hard tech and new space	Founder or VP Manufacturing, 50 to 300 employees
AD	Aerospace and defense supplier	Program or manufacturing engineering lead
MD	Medical device	Quality or operations lead
EL	Electronics and EMS	Operations or engineering lead at a contract manufacturer
AT	Automotive tier supplier	Plant or corporate manufacturing lead
PM	Precision machining and industrial equipment	Owner or operations manager, mid market
SC	Semiconductor and advanced packaging	Fab or OSAT operations leadership
EN	Enterprise multi-plant standardization	Corporate IT or OT architect

Prompts. Fourteen per persona, 112 in total, in three categories.

High Buyer Intent (48 prompts) phrases an active purchase decision qualified by size or situation. Example: which manufacturing execution software should a hardware startup buy when moving from prototype to first production line.

Category With Intent (48 prompts) phrases a functional or standards requirement that only a product can satisfy. Example: which manufacturing platforms generate electronic device history records compliant with 21 CFR Part 11.

Open Question (16 prompts) phrases the problem as how or what approach rather than which product. Example: how should a hardware startup manage production traceability as it scales. These are counted alongside the other two categories, because most of them named products.

No prompt contains a vendor, product, or brand name. Every prompt in the first two categories is constructed so that a product is the only acceptable answer, because questions phrased around a problem rather than a product return methodology and name nobody.

Engines and dates. ChatGPT (GPT-5.6 Sol, High). Gemini (3.6 Flash Extended). Perplexity (Sonar 2). All 336 responses collected in a single window in September 2026. Every prompt was run in a fresh incognito or temporary session on all three engines, so no answer could be influenced by a prior prompt, by account history, or by personalization. Each response was captured in full along with its complete source list.

Extraction. Products were identified using a curated lexicon with alias matching, built by reading all 336 responses rather than from a predefined vendor list. Sources were extracted as raw URLs, reduced to domains, and classified

by hand for the top tier and by rule for the remainder.

3. The weighting problem

The three engines produce wildly different volumes of evidence. Perplexity returned 31.6 URLs per prompt. ChatGPT returned 8.1. Gemini returned 3.0. Perplexity alone accounts for 74 percent of all citations in the dataset.

That asymmetry does not affect the mention counts, because each engine answered the same 112 prompts and named a similar number of products per prompt: 6.7 for ChatGPT, 6.5 for Gemini, 5.8 for Perplexity. Mention data is naturally balanced. Source data is not, and any source metric that sums raw URLs across engines is really a Perplexity metric wearing a disguise. Every source figure in this report is therefore expressed as a share of that engine's own citations.

For the mention ranking, three composite methods are available.

Union coverage counts the distinct prompts where a product was named on any engine. It is how most published leaderboards work, and it treats a product named by one engine in 15 prompts identically to a product named by all three in 15 prompts.

Equal channel weight computes each product's coverage rate within each engine, then averages the three. Every engine contributes one third regardless of how much it says.

Usage weight applies a 56 / 28 / 16 split for ChatGPT, Gemini, and Perplexity, reflecting estimated buyer usage of each assistant.

This index leads with equal channel weight and publishes usage weight beside it in every table. Equal weighting makes no claim about market share, so the ranking does not depend on the usage split being correct. Usage weighting answers the separate question of how much commercial consequence that visibility carries today. Where the two disagree, both figures are shown.

Union coverage appears as a reference column but does not drive the ranking. Every table reports per-engine rates alongside the composite, so readers who prefer a different weighting can recompute the ranking themselves.

Effect on the ranking. The choice barely matters at the top and matters considerably in the middle. Siemens ranks first under all three schemes. Second place is genuinely contested: Rockwell leads on equal weight at 40.5, while Critical Manufacturing leads on usage weight at 43.1 against Rockwell's 42.9, a gap too small to call. Below that the choice moves things significantly. 42Q ranks 11th under union coverage and 15th under equal weight, because its visibility is almost entirely on ChatGPT. iFactory ranks 15th and 17th respectively under those two, and 26th under usage weight, for the mirror-image reason.

4. Core MES and MOM ranking

The primary leaderboard. Vendors meeting the category boundary in section 2, ranked by equal channel weight. Coverage figures are the percentage of the 112 prompts in which the vendor was named on that engine.

Rank	Vendor / Product	ChatGPT	Gemini	Perplexity	Equal weight	Usage weight	Union (of 112)	Personas
1	Siemens Opcenter	64.3%	64.3%	49.1%	59.2	61.9	83	8
2	Rockwell (Plex + FactoryTalk)	47.3%	37.5%	36.6%	40.5	42.9	70	8
3	Tulip	31.2%	44.6%	36.6%	37.5	35.9	66	8
4	Critical Manufacturing	52.7%	38.4%	17.9%	36.3	43.1	62	8
5	Dassault DELMIA	34.8%	21.4%	23.2%	26.5	29.2	48	8
6	SAP Digital Manufacturing	27.7%	22.3%	23.2%	24.4	25.5	47	8
7	Aegis FactoryLogix	20.5%	16.1%	9.8%	15.5	17.6	30	7
8	GE Vernova Proficy	17.9%	9.8%	8.0%	11.9	14.0	25	7
9	AVEVA MES	15.2%	6.2%	12.5%	11.3	12.2	24	6
10	MasterControl Mx	9.8%	13.4%	8.0%	10.4	10.5	18	5
11	Applied Materials SmartFactory	7.1%	12.5%	10.7%	10.1	9.2	17	5
12	iBase-t Solumina	9.8%	8.9%	5.4%	8.0	8.9	14	3
13	Eyelit	8.9%	3.6%	8.0%	6.8	7.3	15	3
13	First Resonance ION	3.6%	14.3%	2.7%	6.8	6.4	17	3
15	42Q	16.1%	0.9%	2.7%	6.5	9.7	21	5
16	Infor MES / CloudSuite	3.6%	8.9%	6.2%	6.2	5.5	16	6
17	iFactory	0.0%	0.9%	12.5%	4.5	2.2	15	7
17	iTAC MOM Suite	8.0%	2.7%	2.7%	4.5	5.7	13	2
19	IBM SiView	6.2%	2.7%	2.7%	3.9	4.7	10	1
19	Körber PAS-X	6.2%	2.7%	2.7%	3.9	4.7	10	3
21	Epicor Advanced MES	3.6%	2.7%	3.6%	3.3	3.3	7	2
21	Oracle Manufacturing / Fusion	3.6%	2.7%	3.6%	3.3	3.3	9	6
21	Sepasoft MES	8.9%	0.9%	0.0%	3.3	5.3	10	5
24	Epsilon3	0.0%	4.5%	4.5%	3.0	2.0	10	3
25	Connektica MES	0.0%	0.9%	7.1%	2.7	1.4	8	3
25	Factory AI	0.0%	0.0%	8.0%	2.7	1.3	9	5
25	PICO MES	4.5%	2.7%	0.9%	2.7	3.4	9	4
25	WorkCell / ShopVue	0.9%	4.5%	2.7%	2.7	2.2	8	5
29	Parsec TrakSYS	0.0%	2.7%	4.5%	2.4	1.5	7	4

Below this line, 24 further core products were named in fewer than seven prompts and are listed in Appendix A. At that volume the difference between two appearances and one is noise rather than a competitive gap, so they are reported as an unordered set.

Reading the table. The structure is one dominant vendor, a three-way contest for second, a stable middle, and a long tail.

Opcenter is named in three quarters of all prompts and is the only entry to exceed 49 percent coverage on every engine.

Second place does not resolve cleanly. Rockwell wins on equal weight through consistency, ranking in the top four on all three engines. Critical Manufacturing wins on usage weight through a commanding ChatGPT position, where it is second at 52.7 percent, but collapses to eighth on Perplexity at 17.9 percent. Tulip sits between them with the most balanced profile of the three, ranking second on Gemini and Perplexity while only fifth on ChatGPT. Three vendors, three different shapes of visibility, separated by less than four points.

5. Extended ranking including adjacent products

The same measure applied to every product named, regardless of category. This table answers a different question: when a buyer asks an answer engine about shop floor software, what are they actually shown?

Rank	Product	Category	ChatGPT	Gemini	Perplexity	Equal weight	Union
1	Siemens Opcenter	Core	64.3%	64.3%	49.1%	59.2	83
2	Rockwell (Plex + FactoryTalk)	Core	47.3%	37.5%	36.6%	40.5	70
3	Tulip	Core	31.2%	44.6%	36.6%	37.5	66
4	Critical Manufacturing	Core	52.7%	38.4%	17.9%	36.3	62
5	Dassault DELMIA	Core	34.8%	21.4%	23.2%	26.5	48
6	SAP Digital Manufacturing	Core	27.7%	22.3%	23.2%	24.4	47
7	SAP S/4HANA	ERP	25.9%	25.0%	22.3%	24.4	46
8	Siemens Teamcenter	PLM	34.8%	19.6%	8.0%	20.8	48
9	Aegis FactoryLogix	Core	20.5%	16.1%	9.8%	15.5	30
10	GE Vernova Proficy	Core	17.9%	9.8%	8.0%	11.9	25
11	AVEVA MES	Core	15.2%	6.2%	12.5%	11.3	24
11	Epicor Kinetic	ERP	12.5%	12.5%	8.9%	11.3	22
13	PTC Windchill	PLM	4.5%	16.1%	11.6%	10.7	25
14	MasterControl Mx	Core	9.8%	13.4%	8.0%	10.4	18
15	Applied Materials SmartFactory	Core	7.1%	12.5%	10.7%	10.1	17

Rank	Product	Category	ChatGPT	Gemini	Perplexity	Equal weight	Union
16	iBase-t Solumina	Core	9.8%	8.9%	5.4%	8.0	14
17	Oracle NetSuite	ERP	5.4%	8.9%	8.9%	7.7	18
18	Katana	ERP	2.7%	9.8%	9.8%	7.4	18
19	Microsoft Dynamics 365	ERP	2.7%	4.5%	13.4%	6.8	18
19	Eyelit	Core	8.9%	3.6%	8.0%	6.8	15
19	First Resonance ION	Core	3.6%	14.3%	2.7%	6.8	17
19	Fulcrum	ERP	10.7%	8.9%	0.9%	6.8	15
23	42Q	Core	16.1%	0.9%	2.7%	6.5	21
24	Infor MES / CloudSuite	Core	3.6%	8.9%	6.2%	6.2	16

The adjacency problem is larger than expected. Core MES and MOM products account for only 62 percent of all product mentions. ERP takes 17.6 percent, PLM 6.7 percent, SMT and equipment vendor software 4.2 percent, and the remainder splits across SCADA and platform tools, semiconductor yield and metrology, machine monitoring, QMS, work instructions, labeling, and scheduling.

Two adjacent products crack the top eight. SAP S/4HANA ties SAP Digital Manufacturing at 24.4, which means the engines are as likely to point a shop floor buyer at SAP's ERP as at SAP's MES. Siemens Teamcenter at 20.8 outranks every specialist MES vendor except the top six.

For vendors, the strategic reading is that the competitive set in AI answers is wider than the competitive set in a normal deal cycle. A specialist MES vendor competes for visibility against the ERP and PLM incumbents that already own the buyer's mental model, and the engines reproduce that mental model faithfully.

6. Per-engine leaderboards

The three engines disagree more than they agree, and compositing hides that structure.

6.1 ChatGPT (GPT-5.6 Sol, High)

Rank	Vendor	Coverage	Own-domain URLs
1	Siemens Opcenter	64.3%	107
2	Critical Manufacturing	52.7%	78
3	Rockwell	47.3%	50
4	Dassault DELMIA	34.8%	51
5	Tulip	31.2%	39
6	SAP Digital Manufacturing	27.7%	22

Rank	Vendor	Coverage	Own-domain URLs
7	Aegis FactoryLogix	20.5%	44
8	GE Vernova Proficy	17.9%	13
9	42Q	16.1%	17
10	AVEVA MES	15.2%	15

ChatGPT is the most vendor-grounded of the three engines. 82.0 percent of the URLs it cited were vendor domains and 8.4 percent were standards and regulatory bodies such as the FDA, IPC, ISA, and SAE. It cited no statistics content farms at all and almost no trade press. When ChatGPT names a product, 68.3 percent of the time it also cites that product's own website in the same answer.

This makes ChatGPT the hardest engine to influence with third-party content and the most responsive to a vendor's own site. It is also the engine where incumbency shows most clearly, since the leaders here are the vendors with the largest and best-structured documentation estates.

6.2 Gemini (3.6 Flash Extended)

Rank	Vendor	Coverage	Own-domain URLs
1	Siemens Opcenter	64.3%	6
2	Tulip	44.6%	14
3	Critical Manufacturing	38.4%	8
4	Rockwell	37.5%	0
5	SAP S/4HANA	25.0%	0
6	SAP Digital Manufacturing	22.3%	0
7	Dassault DELMIA	21.4%	1
8	Siemens Teamcenter	19.6%	0
9	Aegis FactoryLogix	16.1%	3
9	PTC Windchill	16.1%	0
11	First Resonance ION	14.3%	12

Gemini is the hardest engine to read. It returned only 3.0 URLs per prompt, roughly a tenth of Perplexity's volume, while naming about as many products per answer as the other two engines. Of the little it does cite, 46.5 percent is vendor domains, a higher share than Perplexity's 35.4 percent. So Gemini is not indifferent to vendor sources. It simply shows very little evidence for anything it says.

The consequence is that Gemini's citation support rate of 9.4 percent, reported in section 8, should be read as a volume artifact rather than as a preference for third-party sources. Gemini appears to be answering primarily from training rather than retrieval, which makes its rankings the closest available proxy for what these models have absorbed rather than what they can currently find.

A second consequence is that individual domains occupy improbably large shares of Gemini's citation pool. Fabrico accounts for 5.41 percent of all Gemini citations in this study from 18 URLs. Read those percentages as thin rather than dominant.

First Resonance at 14.3 percent is the standout. It is the only newer entrant to place top eleven on the engine that appears least influenced by current content.

6.3 Perplexity (Sonar 2)

Rank	Vendor	Coverage	Own-domain URLs
1	Siemens Opcenter	49.1%	27
2	Rockwell	36.6%	21
2	Tulip	36.6%	34
4	SAP Digital Manufacturing	23.2%	12
4	Dassault DELMIA	23.2%	8
6	SAP S/4HANA	22.3%	0
7	Critical Manufacturing	17.9%	7
8	Microsoft Dynamics 365	13.4%	33
9	iFactory	12.5%	160
9	AVEVA MES	12.5%	0

Perplexity is the retrieval engine, and the most open to influence. It returned 31.6 URLs per prompt, nearly four times ChatGPT and ten times Gemini, and only 35.4 percent of those citations were vendor domains. The rest came from consultancy and SI blogs at 10.3 percent, statistics content farms at 12.2 percent, media at 5.2 percent, analysts at 4.9 percent, and a long tail of small and unverified sites at 22.8 percent.

Perplexity is also the only engine that cites LinkedIn in volume: 115 URLs across 32 prompts and all eight personas. No other engine cited LinkedIn once.

The competitive implication is direct. On Perplexity, a vendor competes not only through its own site but through every third-party page that lists it. That is why the challenger cohort in section 8 exists almost exclusively here.

7. Persona leaderboards

Fourteen prompts per persona, ranked by equal channel weight within the persona. Vendors whose persona rank differs from their overall rank by four or more places are flagged in bold, since those gaps represent visibility a single aggregate ranking would erase.

Hard tech and new space

Rank	Vendor	Score	Overall rank
1	Tulip	76.2	3
2	First Resonance ION	42.9	13
3	Siemens Opcenter	35.7	1
4	Rockwell	23.8	2
5	Critical Manufacturing	19.0	4
6	Aegis FactoryLogix	16.7	7

The only persona where Opcenter does not rank first or second. Tulip's 76.2 is the highest single-persona score any vendor achieves anywhere in this study. First Resonance appears in 12 of 14 hard tech prompts, and 12 of its 17 total appearances come from this persona alone. This is the clearest case in the dataset of a challenger owning a segment in AI answers.

Aerospace and defense supplier

Rank	Vendor	Score	Overall rank
1	Siemens Opcenter	54.8	1
2	iBase-t Solumina	52.4	12
3	Dassault DELMIA	42.9	5
4	Rockwell	38.1	2
5	SAP Digital Manufacturing	23.8	6
6	Critical Manufacturing	19.0	4

iBase-t appears in 11 of 14 aerospace prompts and in only three other prompts across the entire study. It is the most concentrated position of any vendor here, and within its segment it is nearly level with Siemens.

Medical device

	Rank	Vendor	Score	Overall rank
	1	MasterControl Mx	69.0	10
	2	Siemens Opcenter	61.9	1
	3	Tulip	54.8	3
	4	Rockwell	45.2	2
	5	Critical Manufacturing	38.1	4
	6	Körber PAS-X	19.0	19

The only persona where a vendor other than Siemens or Tulip takes first place outright. MasterControl appears in 12 of 14 medical device prompts and in six prompts elsewhere. Rockwell's strong showing here comes via PharmaSuite and MedDeviceSuite rather than Plex, which the consolidation makes visible.

Electronics and EMS

	Rank	Vendor	Score	Overall rank
	1	Aegis FactoryLogix	71.4	7
	2	Siemens Opcenter	69.0	1
	3	Critical Manufacturing	64.3	4
	4	Tulip	31.0	3
	5	Rockwell	26.2	2
	6	iTAC MOM Suite	23.8	17

Aegis appears in 12 of 14 electronics prompts and takes first place ahead of Siemens. This persona also draws the most adjacent competition of any: SMT equipment vendor software from Cogiscan, Panasonic, Fuji, Yamaha, Mycronic, ASMPT, and Koh Young appears throughout, competing directly against MES for the same answer slot.

Automotive tier supplier

Rank	Vendor	Score	Overall rank
1	Siemens Opcenter	71.4	1
2	Rockwell	66.7	2
3	SAP Digital Manufacturing	50.0	6
3	Dassault DELMIA	50.0	5
5	Tulip	33.3	3
6	Critical Manufacturing	28.6	4

The most conventional persona in the study and the one that most closely mirrors the overall ranking. No vendor moves more than three places. Automotive is where the established enterprise vendors are strongest and where no challenger has any foothold at all.

Precision machining and industrial equipment

Rank	Vendor	Score	Overall rank
1	Tulip	42.9	3
2	Rockwell	33.3	2

Rank	Vendor	Score	Overall rank
3	Siemens Opcenter	26.2	1
4	Epicor Advanced MES	23.8	21
5	Infor MES / CloudSuite	14.3	16
5	Dassault DELMIA	14.3	5

Opcenter's weakest persona by a wide margin, at 26.2 against an overall 59.2. This is also the persona where adjacency is most severe: Epicor Kinetic, ProShop, Fulcrum, Global Shop Solutions, JobBOSS, Odoo, Katana, and MRPeasy all appear repeatedly, and machine monitoring vendors including MachineMetrics, Datanomix, Amper, and Guidewheel compete for the same answers. A mid market machine shop asking an AI which shop floor system to buy is more likely to be shown an ERP than an MES.

Semiconductor and advanced packaging

Rank	Vendor	Score	Overall rank
1	Siemens Opcenter	81.0	1
2	Applied Materials SmartFactory	71.4	11
3	Critical Manufacturing	64.3	4
4	Eyelit	40.5	13
5	IBM SiView	31.0	19
6	SAP Digital Manufacturing	14.3	6

The most distinctive persona in the study, and the one that best validates the persona design. Opcenter's 81.0 is its highest score anywhere. Applied Materials appears in 13 of 14 semiconductor prompts and almost nowhere else. IBM SiView appears in 10 prompts, all of them semiconductor, making it the single most segment-locked product in the dataset. Eyelit takes 11 of its 15 appearances here.

Enterprise multi-plant standardization

Rank	Vendor	Score	Overall rank
1	Rockwell	78.6	2
2	Siemens Opcenter	73.8	1
3	Dassault DELMIA	66.7	5
3	SAP Digital Manufacturing	66.7	6
5	AVEVA MES	57.1	9
6	Critical Manufacturing	50.0	4

The only persona where Opcenter is beaten outright, and the persona with the highest scores overall, meaning answers here concentrate on a small set of names. AVEVA takes 13 of its 24 total appearances in this one persona. Unified namespace and platform vendors including HighByte, HiveMQ, Litmus, Solace, EMQX, and Cognite appear almost exclusively here.

Summary of persona effects. Five vendors rank in a persona top two while sitting outside the overall top nine: First Resonance, iBase-t, MasterControl, Aegis, and Applied Materials. For those vendors the aggregate leaderboard understates their position in the only segment where they compete. For Siemens, the persona view reveals the one genuine weakness in an otherwise total dominance, which is the mid market machine shop.

8. Source analysis: what is actually driving visibility

This is where the mention data and the citation data have to be read together, because coverage alone cannot distinguish a product the models remember from a product the models can find.

8.1 The citation support rate

For every product mention, we checked whether that same engine cited that product's own domain in the same answer.

Engine	Mentions	Backed by own-domain citation	Support rate	URLs per prompt
ChatGPT	751	513	68.3%	8.1
Perplexity	651	156	24.0%	31.6
Gemini	724	68	9.4%	3.0

The Gemini figure requires the caveat from section 6.2. Its low support rate reflects how little it cites, not a preference for third parties. Product-level support rates below are therefore computed across all three engines combined, where the ChatGPT and Perplexity signal dominates.

By product, among those named in seven or more prompts:

Product	Mentions	Support rate	Reading
Connektica MES	9	100%	Retrieval-driven
iFactory	15	93%	Retrieval-driven
Epsilon3	10	70%	Retrieval-driven
42Q	22	68%	Retrieval-driven
Factory AI	9	67%	Retrieval-driven
PICO MES	9	67%	Retrieval-driven

Product	Mentions	Support rate	Reading
WorkCell / ShopVue	9	67%	Retrieval-driven
Eyelit	23	65%	Retrieval-driven
First Resonance ION	23	57%	Mixed, leaning retrieval
Sepasoft MES	11	55%	Mixed
IBM SiView	13	54%	Mixed
Critical Manufacturing	122	51%	Mixed
iBase-t Solumina	27	48%	Mixed
iTAC MOM Suite	15	47%	Mixed
Aegis FactoryLogix	52	46%	Mixed
Dassault DELMIA Apriso	87	45%	Mixed
Infor MES / CloudSuite	21	43%	Mixed
Siemens Opcenter	199	42%	Mixed
Tulip	126	37%	Leaning recall
MasterControl Mx	35	37%	Leaning recall
AVEVA MES	38	37%	Leaning recall
SAP Digital Manufacturing	82	34%	Leaning recall
GE Vernova Proficy	40	32%	Leaning recall
Rockwell Plex	114	31%	Leaning recall
Körber PAS-X	13	31%	Leaning recall
Applied Materials SmartFactory	34	29%	Recall
Rockwell FactoryTalk	51	2%	Almost pure recall
DELMIAworks	14	0%	Pure recall
Epicor Advanced MES	11	0%	Pure recall

Two mechanisms, and they are not mutually exclusive. Products at the top of this list are being found. Products at the bottom are being remembered. Rockwell FactoryTalk is named in 34 prompts with essentially no supporting citation anywhere, which means that visibility currently rests on how thoroughly the name is embedded in training data.

Trained-in visibility is not a finished position that content cannot improve. Siemens sits at 42 percent support with 107 own-domain citations on ChatGPT alone, meaning it is both remembered and findable, and that combination explains the size of its lead. Nothing in this data suggests Siemens has reached a ceiling. A vendor with strong recall that also publishes well should expect to extend its lead, because the two mechanisms compound rather than substitute.

The corresponding risk falls on vendors whose position is recall-only. Training corpora refresh on a cycle nobody outside the labs controls, retrieval indexes update continuously, and a name that is well remembered today has no defense if a competitor systematically out-publishes it in the retrieval layer. Incumbency in this category is a lead, not a moat.

8.2 The content-driven challengers, and why they matter more than their rank suggests

Products whose Perplexity coverage substantially exceeds their ChatGPT coverage, ordered by that gap.

Product	ChatGPT	Perplexity	Gap	Own-domain URLs on Perplexity
iFactory	0.0%	12.5%	+12.5	160
Microsoft Dynamics 365	2.7%	13.4%	+10.7	33
Factory AI	0.0%	8.0%	+8.0	54
Katana	2.7%	9.8%	+7.1	0
Connektica MES	0.0%	7.1%	+7.1	30
PTC Windchill	4.5%	11.6%	+7.1	0
Tulip	31.2%	36.6%	+5.4	34
Accevo QMS	0.0%	5.4%	+5.4	18
Epsilon3	0.0%	4.5%	+4.5	26
Parsec TrakSYS	0.0%	4.5%	+4.5	9

iFactory is the clearest case. It ranks 17th on the core leaderboard and 9th on Perplexity specifically, it is cited 160 times to its own domain, and 93 percent of its mentions come with a citation to its own site. It appears in seven of the eight personas, which is broader persona reach than iBase-t, MasterControl, or First Resonance achieve. On ChatGPT it does not exist.

A vendor looking only at ChatGPT would conclude its AEO program had failed. The sequencing explains why that reading is premature.

In AEO Wrangler client engagements we have repeatedly seen the same order of arrival: Perplexity first, Gemini second, ChatGPT last, with months separating the first from the last for the same body of published content. The mechanism is straightforward. Perplexity retrieves live, so newly published content can appear within days. Gemini retrieves less and leans more on training. ChatGPT, on the evidence in section 6.1, is grounded in vendor sites but appears to weight established and heavily linked domains far more conservatively, and its training corpus lags publication by a year or more.

This makes retrieval-engine visibility a leading indicator rather than a consolation prize. It is not a guarantee. We have no controlled evidence that Perplexity visibility converts to ChatGPT visibility on any reliable timetable, and some of it may never convert. But the practical guidance is unambiguous: a content program judged solely on ChatGPT results will be abandoned exactly when it is working.

Factory AI, Connektica, Epsilon3, Fabrico with 66 own-domain citations, and SYMESTIC with 59 show the same signature at smaller scale. Each has built real visibility on the retrieval engines from a standing start, in a category dominated by vendors thirty times their size.

8.3 The incumbent signature

The mirror image, products whose ChatGPT coverage substantially exceeds their Perplexity coverage.

Product	ChatGPT	Perplexity	Gap
Critical Manufacturing	52.7%	17.9%	-34.8
Siemens Teamcenter	34.8%	8.0%	-26.8
Siemens Opcenter	64.3%	49.1%	-15.2
42Q	16.1%	2.7%	-13.4
Dassault DELMIA Apriso	34.8%	21.4%	-13.4
Rockwell Plex	42.0%	25.9%	-16.1
Aegis FactoryLogix	20.5%	9.8%	-10.7
GE Vernova Proficy	17.9%	8.0%	-9.9

Critical Manufacturing's gap is the largest in the study in either direction. It is the second most visible vendor on ChatGPT and only the seventh most visible on Perplexity. The pattern suggests deep documentation and long-standing web presence that the training corpus has absorbed, paired with weaker performance in live retrieval against third-party sources.

Read alongside section 8.2, the two tables describe opposite exposures. The challengers are winning the layer that updates fastest and losing the layer that carries the most buyers. The incumbents are doing the reverse. Neither position is stable, and the vendor that fixes its weaker side first will be the one that consolidates.

8.4 Who owns the citation layer

Top cited domains, with the channel concentration flag from the source dataset.

Domain	Type	Prompts	Total URLs	Concentration
siemens.com	Vendor	58	75	All three, ChatGPT-skewed
criticalmanufacturing.com	Vendor	51	80	All three, ChatGPT-skewed
ifactoryapp.com	Vendor	47	166	No ChatGPT
linkedin.com	Social	32	115	Perplexity only
3ds.com	Vendor	31	48	All three, ChatGPT-skewed
blogs.sw.siemens.com	Vendor	29	47	No Gemini
gitnux.org	Content farm	28	99	Perplexity only
tulip.co	Vendor	27	70	All three, balanced

Domain	Type	Prompts	Total URLs	Concentration
plex.rockwellautomation.com	Vendor	26	44	No Gemini
reliamag.com	Media	25	65	No ChatGPT
f7i.ai (Factory AI)	Vendor	25	54	Perplexity only
fabrico.io	Vendor	24	66	No ChatGPT
demystifyingplm.com	Analyst	21	49	No ChatGPT
caddissystems.com	Consultancy	20	49	No ChatGPT
guideflow.com	Content farm	20	43	No ChatGPT

Only 21 of the 847 domains in this study are cited by all three engines. 536 are cited by Perplexity alone.

Three observations that matter for anyone trying to influence these systems.

Traditional analyst coverage is a small factor. Gartner accounts for 21 URLs across 12 prompts, 0.55 percent of ChatGPT citations and 0.37 percent of Perplexity citations. Demystifying PLM, an independent MES buyer's guide, was cited in 21 prompts across all eight personas, more than double Gartner's volume, and outperforms Gartner on both Gemini and Perplexity. Specialist independent content is competing successfully with institutional analyst brands inside these systems.

Statistics content farms are a meaningful citation channel on Perplexity and nowhere else. gitnux, wifitalents, zipdo, worldmetrics, guideflow, sysgenpro and similar account for 12.2 percent of all Perplexity citations, more than double what analysts receive there. ChatGPT cites none of them.

LinkedIn is a real retrieval surface, but only on Perplexity, where it appears in 32 prompts across all eight personas. Given that most B2B manufacturing software marketing already runs through LinkedIn, this is the single most underexploited finding in the source data.

9. What this means

For Siemens. The visibility contest is won by a margin no reasonable methodology overturns, and it is won on both mechanisms at once. The exposure is narrow and specific: mid market machine shops, where Opcenter scores 26.2 against its overall 59.2 and loses to Tulip. The strategic risk here is behavioral rather than competitive. At 42 percent citation support there is substantial room to publish into, and a lead built partly on recall will decay if it is treated as permanent.

For the second-place contest. Rockwell, Tulip, and Critical Manufacturing are separated by less than four points and by completely different profiles. Rockwell wins on consistency across engines. Critical Manufacturing wins on the engine most buyers use and is badly exposed on the engine that updates fastest. Tulip is the most balanced and the only one of the three with a persona it owns outright.

For challengers with strong segment positions. First Resonance in hard tech, iBase-t in aerospace, MasterControl in medical device, Aegis in electronics, and Applied Materials in semiconductor each own a segment while sitting outside the overall top nine. None has meaningful visibility outside its home persona. Whether that is a moat or a ceiling is the question each of them should be asking.

For newer entrants. The retrieval channel is open and demonstrably winnable, and early wins there should not be dismissed because ChatGPT has not moved yet. iFactory, Factory AI, Connektica, Fabrico, Epsilon3, and SYMESTIC have built real Perplexity visibility from nothing through published content. Whether that converts into ChatGPT presence is the most commercially important open question in this dataset, and answering it properly requires a repeat run twelve months from now.

For everyone. The three engines answer the same questions from different evidence. A vendor measuring its AI visibility on one engine is measuring roughly a third of the picture, and possibly the wrong third.

10. Limitations

Single collection window. All 336 responses were captured in September 2026. These systems change continuously and a rerun would produce different numbers.

Single run per prompt per engine. No repeat sampling, so response variance within an engine is not measured. Some portion of the tail is likely to be noise.

Model version specificity. Results describe GPT-5.6 Sol High, Gemini 3.6 Flash Extended, and Sonar 2. Other model tiers within the same products may behave differently.

Lexicon extraction. Products were matched by a curated alias list. A product named once under an unusual spelling could have been missed. Products appearing in fewer than five prompts should be treated as directional.

Source classification. Domain-level classification was hand-curated for the top tier and rule-based below it. 948 URLs across 473 domains remain in an unverified bucket, mostly small vendor and SEO sites. Source composition percentages carry that uncertainty.

Citation support is a proxy. A mention accompanied by a citation to the vendor's own domain suggests retrieval-driven visibility, but does not prove causation. The engine may have named the product first and cited it afterward.

The leading-indicator claim is observational. The sequencing described in section 8.2 comes from AEO Wrangler client engagements rather than from this dataset, which is a single point in time and cannot show conversion between engines.

Usage weighting is an estimate. The 56 / 28 / 16 split reflects an assumption about buyer usage and is offered as a scenario rather than a measurement. Equal weighting is the primary ranking for that reason.

No prompt volume data. This study measures how engines answer these questions. It does not measure how often real buyers ask them.

Appendix A: Core products named in fewer than seven prompts

Reported as an unordered set, since at this volume the differences are not meaningful: MPDV HYDRA, SYMESTIC, Fabrico, OpenMES, Autodesk Fusion Operations, NoMuda VisualFactory, Advantive PINpoint, MITS MES, 24Flow, L2L, Optimal Electronics Optel, DynamxMFG / TotalControlPro, Manufacturo, Fuuz, Emerson Syncade, Volvsoft, Radley MES, EZ-MES, Shoplogix, SYSTEMA, Tempo Manufacturing Cloud, ForgePlan.

Appendix B: Adjacent products by category

ERP: SAP S/4HANA, Epicor Kinetic, Oracle NetSuite, Katana, Microsoft Dynamics 365, Fulcrum, ProShop, Odoo, Infor LN / SyteLine, MRPeasy, JobBOSS / ECI, Cetec, Acumatica, IFS Cloud, Global Shop Solutions, MIE Trak Pro, QAD, Genius ERP, Fishbowl, SYSPRO, Sage X3, Deltek Costpoint, inFlow, WorkGuru, WEEC.

PLM: Siemens Teamcenter, PTC Windchill, Arena, Propel, Duro, Onshape, OpenBOM.

SMT and equipment: Cogiscan, Panasonic PanaCIM, Siemens Valor, Fuji Nexim, Yamaha YSUP, ASMPT WORKS, Mycronic MYPro, JUKI IFS-NX, Koh Young, Saki, WSW LOJISTIX.

Semiconductor yield and metrology: camLine LineWorks, INFICON FabGuard, PDF Solutions, Onto Innovation, Synopsys, Teradyne, Advantest, BISTel.

SCADA, platform, and UNS: Inductive Automation Ignition, AVEVA PI System, Litmus Edge, HiveMQ, HighByte, Azure IoT Operations, Cognite, EMQX, Solace PubSub+.

Machine monitoring: MachineMetrics, TeepTrak, Guidewheel, Amper, Datanomix, Scytec DataXchange, Predator MDC.

QMS: Accevo, SG Systems V5, Net-Inspect, 1factory, SimplerQMS, SafetyChain.

Work instructions: VKS, Dozuki, Augmentir.

Labeling and serialization: TEKLYNX, Systech, TraceLink.

Scheduling and APS: PlanetTogether, Asprova.

Appendix C: Consolidated vendors, per-product detail

Product	ChatGPT	Gemini	Perplexity	Equal weight	Union
Rockwell Plex	42.0%	33.9%	25.9%	33.9	61
Rockwell FactoryTalk	11.6%	13.4%	20.5%	15.2	34
Rockwell combined	47.3%	37.5%	36.6%	40.5	70

Product	ChatGPT	Gemini	Perplexity	Equal weight	Union
Dassault DELMIA Apriso	34.8%	21.4%	21.4%	25.9	48
DELMIAworks	2.7%	2.7%	7.1%	4.2	13
Dassault combined	34.8%	21.4%	23.2%	26.5	48

Combined figures are the union of prompts naming either product, not the sum of the two rows.

Appendix D: Prompt set and underlying data

The complete set of 112 prompts, the eight persona definitions, and the full mention and source datasets are available at no charge on request to info@aeowrangler.ai.

AEO Wrangler is an independent consultancy specializing in Answer Engine Optimization. This index is editorially independent and carries no paid placement.