

Should You Buy That Agricultural AI?

Eight Questions Before You Sign

Jackson Mambozoukuni

Auctus Agri, East London, South Africa · ORCID 0009-0009-9343-1910 · jackson@auctusagri.com
Auctus Agri Working Paper AA 2026 005 · 23 August 2026 · CC BY 4.0 · DOI: 10.5281/zenodo.22070535

Abstract

Artificial intelligence is marketed to farms as broadly transformative; however, adoption rates are uneven, and many promising pilot projects fail to become profitable. We argue that the deciding factor is rarely the algorithm but the fit between the tool and the conditions of its use. Four conditions predict that fit, and all four can be checked before money is committed. We turn them into an eight-question test, state its scoring rule, probe its sensitivity, apply it to four common vendor propositions, and release it as a free, offline instrument for farm and agribusiness managers.

Keywords: artificial intelligence; technology adoption; decision making; agribusiness management; extension

JEL codes: Q16, Q12, O33, D25

Note on the instrument ecosystem. This article is the middle layer of a three-layer appraisal ecosystem maintained at auctusagri.com. A five-question quick screen decides whether a vendor proposition is coherent enough to deserve the scored conversation this article describes; the eight questions produce a provisional routing verdict; and propositions carrying material investment exposure or irreversible error proceed to the AI-D2V appraisal framework, which underwrites the investment decision itself.

Here is the short version. Before you buy an artificial intelligence product for your farm or agribusiness, ask what decisions it changes, how often you make those decisions, what happens when it is wrong, and whose farms its data came from. If the vendor cannot answer those four questions clearly, the software's sophistication will not save the investment. Fit, not intelligence, determines whether agricultural artificial intelligence pays off.

By artificial intelligence this article means software that learns patterns from data rather than following fixed rules. The claim above is not a dismissal of that technology. It is already doing real work in agriculture: spotting disease before a scout would, guiding sprayers to treat weeds and skip crop, and forecasting yields with useful accuracy. Innovation in the sector has concentrated heavily in systems that replicate perception and judgment rather than muscle, which is precisely why so much of it lands on decisions rather than tasks (Onel et al. 2025). The difficulty is that these successes and the failures sitting beside them look nearly identical in a sales demonstration.

Most Agricultural AI Works. So, Why Do So Many Projects Fail?

The adoption picture is sharply uneven, and it has been for longer than the current excitement. Guidance and auto-steer systems now cover well over half of the acreage of major United States row crops, while information-intensive tools such as variable rate application have stalled far below expectations for two decades (McFadden, Njuki, and Griffin 2023; Lowenberg-DeBoer and Erickson 2019). The same gradient runs through livestock: large dairies adopt advanced technology at several times the rate of small ones (Gillespie, Njuki, and Terán 2024). Measured use of artificial intelligence across agriculture as a whole remains low relative to other industries (Thomasson et al. 2025),

and a federal assessment of precision agriculture identified high up-front costs and uncertain payoffs as the top reasons (U.S. Government Accountability Office 2024).

The usual explanations are all real. Rural connectivity is patchy. Capital is scarce. Farmers reasonably worry about who ends up owning their data, and surveys find that most cannot say what their technology agreements permit and do not trust providers not to share their data with third parties (Wiseman et al. 2019). Managers distrust software that produces an answer without explaining itself, and uncertain returns are hard to finance (Thomasson et al. 2025; Dara, Hazrati Fard, and Kaur 2022).

There is also a timing problem. Venture investment in agrifood technology peaked at \$51.7 billion in 2021 on the promise of transformation, then fell below \$16 billion by 2024 (AgFunder 2025). The sector has entered a colder, more demanding phase in which buyers want evidence of return rather than evidence of capability. Vendors who thrived by demonstrating what their software could do are now asked what it earned. Many cannot say.

Agriculture is not the outlier here; it is the rule. Industry surveys of enterprise artificial intelligence, none of them refereed, keep reporting the same shape: the large majority of corporate pilots produce no measurable profit-and-loss impact, not because the models are weak but because the tools are never embedded in real workflows (MIT Project NANDA 2025). The precise failure rate is contested and does not matter here; the mechanism does, and so does the least-advertised finding, which is the most useful one for a farm office: purchased tools from specialised vendors succeeded roughly twice as often as systems organisations tried to build themselves. The pilot trap, a cycle in which a technology is repeatedly shown to work in demonstrations but never becomes an ordinary, profitable part of the operation, is a general disease of buying intelligence without buying fit.

None of this should surprise an economist. Nearly seventy years ago, the classic study of hybrid corn showed that a plainly superior technology diffused across the American Midwest at a speed set almost entirely by its profitability in each district, not by its ingenuity (Griliches 1957). Farmers adopted where and when it paid, and not before. The artificial intelligence industry is currently rediscovering that result at considerable expense. However, as a diagnosis, the catalogue of barriers helps the individual manager very little. A manager sitting across a table from a vendor does not need a taxonomy of sectoral obstacles. They need to decide this week whether this product is worth the money for this operation. The list of barriers does not tell them that. Neither, unfortunately, does the demonstration, which is engineered to be persuasive rather than diagnostic.

The Question That Matters Is Fit, Not Sophistication

Consider what a vendor actually tells you. The statement “Our model detects leaf disease with 98% accuracy” is about the software in isolation. It is probably true. It is also nearly useless on its own, because it says nothing about whether that accuracy will turn into money on your farm.

To turn accuracy into money, a chain of events has to be completed. The software has to change a decision you actually make. Changing that decision has to improve the outcome. Moreover, the improvement has to be worth more than the cost of producing it. Break any link, and the software can be flawless while the investment returns nothing. A disease detector on a farm that already sprays preventively on a fixed schedule makes no decision. A yield forecast that arrives after the marketing decision has been made improves no outcome. A model that costs more to feed and maintain than the losses it prevents fails the third test.

This is the practical meaning of a simple point: a technology produces an effect only when the conditions that let it work are actually present. The useful consequence is that those conditions are largely visible in advance. A manager can move the evaluation upstream of the cheque, from “is this impressive?” (which demonstrations are engineered to answer) to “will this work here?” (which they are not).

Four Conditions That Predict Whether AI Pays

Four conditions do most of the work in separating the investments that pay from those that stall. Figure 1 summarises how they combine.

Frequency. Machine learning repays its fixed setup cost through repetition. A model informing a decision made many times a season, such as whether to spray a given block, spreads its cost across many uses and compounds small per-decision gains. It is no accident that the clearest commercial success in agricultural artificial intelligence to date, camera-guided targeted spraying, sits on a decision made millions of times per pass: one manufacturer reports average herbicide savings of 59% across more than a million acres (John Deere 2024). Stated properly, the economics is that the expected value across the decision opportunities a model informs must recover its fixed and recurrent costs. Frequency is the commonest way that happens, because many small gains compound; it is not the only way. A once-a-year cultivar, insurance or marketing decision can carry enough value in a single use to pay for the model that improves it. What a low-frequency proposition loses is room for error: it must clear the same cost bar from far fewer chances, so the stakes of the decision have to do the work repetition would otherwise do.

Reversibility. No model is right every time, so what matters is the cost of its mistakes. Where a wrong call is cheap to correct, automation is attractive and moderate accuracy suffices. Where a wrong call destroys a crop, denies a farmer credit, or puts a machine among people, the burden of proof rises steeply, and a human must stay in the loop, a principle that runs through responsible-use guidance for the sector (Dara, Hazrati Fard, and Kaur 2022). Reversibility, not accuracy alone, should set the bar a system must clear.

Representative data. A model is only as transferable as the data behind it. A disease classifier built on large, well-instrumented farms in one region frequently underperforms on different soils, cultivars, and management. This is the most under-examined condition, precisely because a demonstration on the vendor's data reveals nothing about performance on yours. Ask whose farms the training data came from, and ask where yours will go: data anxieties are not paranoia but a documented, rational response to agreements most farmers have never been able to read on their own terms (Wiseman et al. 2019).

Total cost and ownership. The purchase price is a fraction of the total cost, which includes integration, connectivity, sensor upkeep, data cleaning, staff training, and periodic retraining as conditions drift. Machine learning systems are unusual among technologies in that their maintenance burden is systematically underestimated at purchase; engineers call the accumulating upkeep hidden technical debt, meaning a bill that grows quietly until someone has to pay it (Sculley et al. 2015). Many propositions that look profitable at sticker price lose money once these are counted. Compounding this, the most common failure of a deployed model is not technical but organisational: nobody owns keeping it running after the pilot, so it quietly decays. Ask who runs it in year two, and name them.

Put together, the pattern is compact. Agricultural artificial intelligence tends to earn its keep on decisions that are frequent, reversible, and supported by data resembling your own, at a total cost that pays back within the asset's life, with someone accountable for it.

It tends to disappoint wherever those conditions are absent, regardless of how elegant the underlying model is.

Notice what is absent from that list. Nothing about the algorithm. Nothing about accuracy percentages, model architecture, or whether the product uses the currently fashionable technique. Those things matter, but algorithmic sophistication is not the buyer's first question. Calibration, drift and failure modes do eventually become the buyer's problem, through the losses they can create, and that is exactly what the questions on reversibility and evidence exist to price. The four conditions are the buyer's problem from the start, and the buyer is the only person who can assess them.

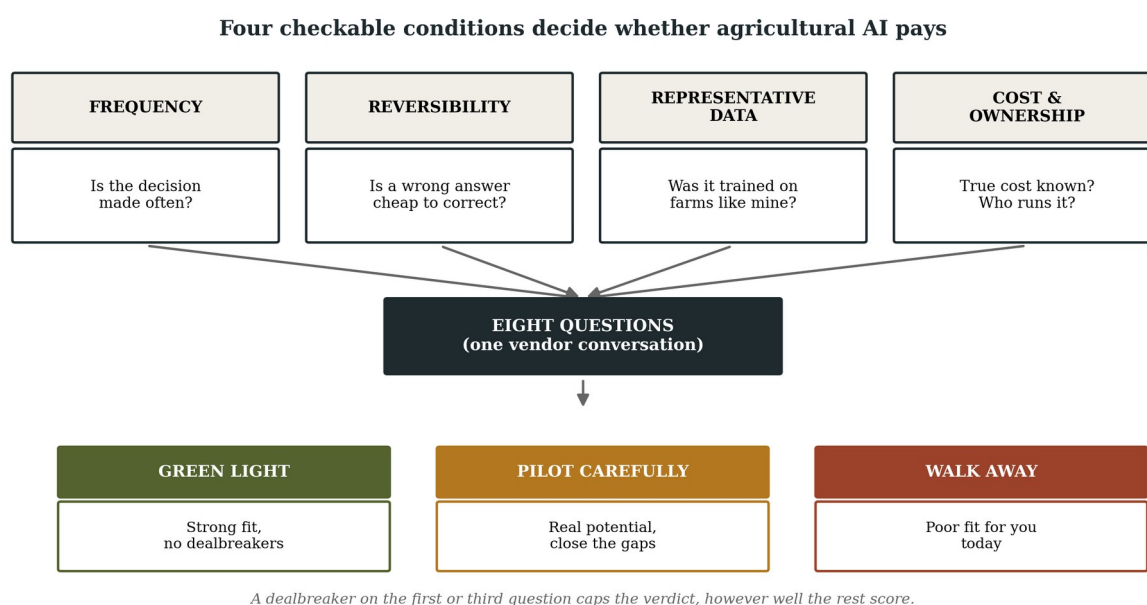


Figure 1. Four checkable conditions determine whether agricultural artificial intelligence pays off. A failure on the first or third question caps the verdict regardless of the others.

Eight Questions a Manager Can Actually Ask

To be useful at the moment of decision, those conditions must become questions answerable from a single vendor conversation. Table 1 presents them. Each is scored, and the total maps to one of three verdicts: green light, pilot carefully, or walk away for now.

#	Question	What it tests	Condition it serves	Score range
1	What decision does this product actually change?	Is a real decision affected?	The decision premise the chain starts from	-1 to 2
2	How often is that decision made?	Frequency: does the benefit compound?	Frequency	0 to 2
3	If the product is wrong, how bad is it?	Reversibility; stakes of error	Reversibility	-1 to 2
4	Whose farms was the model trained on?	Representative training data	Representative data	-1 to 2
5	What evidence of performance can they show?	Proof beyond the demonstration	Verification of the other answers	0 to 2
6	Do you know the true total cost, not the price?	Total cost of ownership	Cost and ownership: the money	-1 to 2
7	Who runs it after the pilot ends?	Operational ownership	Cost and ownership: the people	0 to 2
8	Does it need connectivity and the records you have?	Fit to the actual infrastructure	Cost and ownership: the infrastructure	-1 to 2

Table 1. The eight questions, what each tests, the condition each serves, and the scoring range. Totals run from minus five to sixteen: five of the eight questions can score minus one, and the maximum is 16.

The mapping to the four conditions is deliberate but not one-to-one, and the table's fourth column states it. Question 1 tests the premise the whole chain starts from, that a decision exists to improve. Questions 2, 3 and 4 test the first three conditions directly. Question 5 is verification: the evidence that the other answers are more than the vendor's say-so. Questions 6, 7 and 8 unpack the fourth condition into its three components, because total cost of ownership fails through money, through people and through infrastructure, and a manager needs to check all three.

Using the test takes one conversation: ask the eight questions, score each answer, apply the two caps, and read the verdict. The scoring rule is simple enough to state in full. Each answer earns 2, 1, 0, or -1 points, for a maximum of 16 points. A proposition earns a green light if it scores 11 or more and no answer scores minus one. It earns a pilot carefully if it scores six or more and neither the first question (what decision changes) nor the third (how bad a wrong answer is) scores minus one. Everything else is a walk away.

The scoring is deliberately asymmetric, and three distinct mechanisms produce that asymmetry; it is worth stating them formally because they do different jobs. First, two **hard caps**: a minus one on question 1 (no identifiable decision) or question 3 (irreversible errors) makes the verdict walk away regardless of the total. Second, a **green-light veto**: a minus one on any question blocks a green light but not a pilot, because a dealbreaker anywhere is exactly what a pilot exists to resolve, while a green light claims there is none. Third, the ordinary **numerical thresholds** of eleven and six, which do the undramatic work of ranking everything the caps and the veto let through. The cap on question 3 is deliberately unconditional rather than an interaction with the evidence score: strong evidence lowers an error rate, but it does not undo an irreversible error, so no quantity of evidence lifts the cap. A strong average cannot rescue a proposition that fails a fundamental test; one broken link nullifies the rest of the chain. That structure carries most of the argument, and the appendix shows it: moving the thresholds changes almost nothing, while removing the caps and the veto changes the verdicts.

Two design choices follow. First, the test scores fit, not merit. A low score does not mean the product is bad; it means the product is poorly matched to this operation, now. The same tool can be a walk-away for one farm and a green light for its neighbour. Second, the instrument is free, requires no registration, and runs entirely on the user's device without collecting or transmitting any data, because anxiety about data ownership is itself among the documented barriers to adoption (Wiseman et al. 2019; Thomasson et al. 2025).

Others have worked on this problem. The most developed effort, recently validated with German farmers, evaluates digital technologies against seven criteria, including technological maturity, legal conditions, and social acceptance (Michels and Isenberg 2026). It is careful work, and its authors are candid about what it deliberately does not do: it reveals where constraints arise, but it sets no decision cutoffs, and the farmers who tested it asked for economic criteria it does not yet include. This article supplies both. It is narrower, addressing artificial intelligence rather than digital technology in general. It is unapologetically economic, built on how often a decision recurs, what an error costs, and what the thing costs to run. Moreover, it ends in a verdict rather than a profile, because a manager standing in front of a vendor has to say yes or no.

Testing the Filter on Four Common Propositions

An instrument that returns the same answer for everything is worthless. Table 2 applies the eight questions to four vendor propositions of a kind familiar to most managers,

described from publicly available product categories rather than from any individual company. The four are illustrative rather than empirical, but they show the test discriminating, and doing so for different reasons.

One caution before the table. On the definition this article uses, two of the four propositions may not be artificial intelligence at all: a satellite dashboard that computes standard vegetation indices learns nothing, and some credit scorecards are fixed formulas wearing the label. They are scored anyway, because they are sold as AI, and because the test is deliberately label-indifferent. Fit determines the return; what the brochure calls the product does not.

Vendor proposition	Q1	Q2	Q3	Q4	Q5	Q6	Q7	Q8	Total	Indicative cost (relative)	Verdict
Camera-guided selective sprayer	2	2	1	0	1	0	2	2	10	High: hardware retrofit plus subscription	Pilot carefully
Satellite imagery dashboard subscription	-1	1	2	0	0	2	0	2	6	Lowest of the four: small annual subscription	Walk away
Credit-scoring tool for input finance	2	2	-1	-1	1	0	0	1	4	Moderate: fee per assessment	Walk away
Autonomous weeder, specialty crops	2	2	1	0	2	2	2	2	13	Highest of the four: major capital purchase	Green light

Table 2. Illustrative worked examples, not empirical results. The eight-question test was applied to four archetypal vendor propositions constructed from publicly described product categories. No claim of predictive accuracy is made. Two propositions fail for opposite reasons; the most expensive passes; the cheapest is refused.

The camera-guided selective sprayer scores well on frequency and on ownership, but the evidence is the vendor's own, and the training data comes from other regions. It earns a pilot carefully: promising, with a local trial required before scale.

The satellite imagery dashboard subscription looks harmless. It is cheap, it works offline, and the imagery is genuinely useful. It nevertheless walks away, because the vendor cannot name the decision it changes. Managers subscribe, look at the maps for a season, and change nothing. The cost is small; the return is zero.

The credit-scoring product for input finance fails differently. It touches a frequent and genuinely valuable decision, but denying a farmer credit is irreversible, and the model was trained on other farmers in other places. High stakes and unrepresentative training data are the combination the test is designed to catch.

The autonomous weeder for specialty crops earns a green light: frequent, recoverable errors, independent evidence, known running cost, and a named operator. Notably, it is also the most expensive of the four. Price is not the variable that matters.

The pattern in Table 2 is the point. Two propositions fail, for opposite reasons. One passes despite being costly. The cheapest of the four is the one to refuse.

Using the Test in Extension and Advisory Work

The eight questions were written for a manager, but they may be most valuable in the hands of the people managers turn to for advice.

Extension officers and farm advisors are routinely asked some version of “Should I buy this?” They are rarely equipped to answer, as evaluating a machine learning product appears to require technical expertise that most advisors do not claim to have. This article argues that it does not. Every one of the eight questions can be answered without knowing what a neural network is. An advisor who asks what decision the product changes, who trained it, and who will run it in year two is conducting a more rigorous evaluation than a technical review of the algorithm would. The point generalises: the value of these systems in agriculture lies largely in augmenting managerial judgment rather than replacing it, which makes the judgment itself the thing worth disciplining (Hankins, Kim, and Villacis 2026). Adoption research reaches the same conclusion from the other direction: farmers prioritise profitability over trust, and both are properties of fit rather than of the algorithm (Yeo and Keske 2024).

The questions also work as a teaching device. Walking a group of producers through a single vendor proposition, scoring it together and debating the answers, surfaces the economics of technology adoption more effectively than a lecture on it. Disagreement about the scores is productive because it is usually about the farm, not the software.

Lenders and insurers evaluating whether to finance a technology purchase face the same problem from the other side, and the same questions apply. A proposition that changes no decision will not generate the cash flow to service the loan taken out to buy it.

Finally, the questions are a useful discipline for vendors. A serious vendor welcomes them and has answers ready. A weak one deflects toward accuracy statistics. That asymmetry is itself diagnostic, and costs nothing to observe.

What This Test Cannot Do

An instrument built to resist overclaiming should not overclaim. This is a structured heuristic, not a predictive model. Its weights are reasoned rather than calibrated against outcome data, and thoughtful practitioners will weight them differently. It scores fit; it does not verify a vendor's technical claims, which still requires scrutiny and, where the stakes are high, an independent trial.

It also assumes a decision maker with the discretion to buy or decline. Smallholders operating through cooperatives face collective and institutional questions that this test does not reach, and adapting it to that setting is worthwhile future work. Finally, it is a snapshot. Connectivity improves, data accumulates, prices fall. A walk away today may be a green light in two seasons, which is why the instrument is built to be rerun.

These limits follow from the argument itself. A decision aid is a technology, too, and its value depends on the conditions under which it is used.

So What?

The debate about artificial intelligence in agriculture is conducted mostly between people selling it and people dismissing it. Both positions ask the wrong question: whether the technology works. It does, sometimes, under conditions that can be named. The right

question is narrower and far more useful: will it work here, on this decision, at this price, with this data, run by this person?

That question can be answered before the money moves. Answering it well is neither optimism nor cynicism about artificial intelligence. It is ordinary management discipline, applied to a technology that has been unusually successful at avoiding it, and it is the same discipline the economics of adoption has recommended since hybrid corn: pay for what pays, here. If a vendor cannot tell you what decision their product changes, there is not yet one worth paying for. Everything else in this article is an elaboration of that sentence.

Appendix. Score Construction and Sensitivity

This appendix explains why the thresholds are where they are and tests whether the verdicts depend on that choice. Readers who wish to inspect or alter the rule can do so in the companion instrument, whose source is open.

Why these thresholds? Eleven of sixteen is the point at which a proposition must have answered most questions at the top of the range, which is what a green light should require. Six of sixteen is roughly the point below which a proposition fails more questions than it passes. Neither number is derived from outcome data, and neither is claimed to be optimal. They are conventions, and the analysis below shows that the verdicts depend little on them.

The two hard caps are not conventions. A product that changes no identifiable decision cannot generate value by any route, so no combination of other strengths should let it pass. A product whose errors are irreversible exposes the operation to losses it cannot undo, and the cap is unconditional on the evidence score for the reason given in the main text: evidence lowers an error rate, it does not undo an irreversible error. These are the two places where averaging would be misleading, so it is disabled there. The green-light veto is a third, softer mechanism: any minus one blocks green because a green light asserts that no dealbreaker exists, but it leaves the pilot route open, since resolving a single identified dealbreaker is what a pilot is for.

Sensitivity of the four worked examples. Table A1 recomputes the four propositions of Table 2 under five alternative rules: stricter thresholds, looser thresholds, no hard cap at all, a hard cap extended to the data question, and equal weighting, in which the minus-one option is removed so that every item scores from 0 to 2.

Alternative rule	Sprayer	Dashboard	Credit tool	Weeder	Verdicts changed
Baseline (11 / 6; caps on Q1, Q3)	Pilot	Walk	Walk	Green	baseline
Stricter thresholds (12 / 7)	Pilot	Walk	Walk	Green	0 of 4
Looser thresholds (10 / 5)	Green	Walk	Walk	Green	1 of 4
No hard cap (score only)	Pilot	Pilot	Walk	Green	1 of 4
Hard cap extended to Q4 (data)	Pilot	Walk	Walk	Green	0 of 4
Equal weights, no negative scores	Pilot	Pilot	Pilot	Green	2 of 4

Table A1. Verdicts under alternative scoring rules. The verdicts are stable to movement within the thresholds and unstable to the removal of the hard cap and negative scores.

The pattern is informative, and it is not the one a defender of the instrument would have chosen. Tightening the thresholds changes nothing. Extending the hard cap to the data

question changes nothing. Loosening the thresholds changes a single verdict, moving the sprayer from pilot carefully to green light, a difference of emphasis rather than direction.

What changes the verdicts is removing the two features that carry the argument. Remove the hard cap and score the propositions based solely on their totals: the satellite dashboard, which changes no decision, is upgraded from walk away to pilot carefully. Remove the negative scores as well, so that a dealbreaker merely earns zero rather than a penalty, and both walk-away propositions become pilots. The arithmetic, in other words, is doing very little work. The asymmetry is doing almost all of it.

This is the appropriate sensitivity result for an instrument of this kind. It should be robust to the arbitrary parts, which the thresholds are, and sensitive to the substantive parts, which the hard caps and penalties are. It also identifies the claim a critic should attack: not the numbers, but the judgment that a product that changes no decision, or fails irreversibly, should be refused, whatever else is true of it.

These are worked examples rather than validation. Establishing predictive value would require applying the test prospectively to real purchases and observing outcomes; establishing reliability would require an inter-rater study. Both are worth doing. Neither is done here, and no claim to either is made.

Companion Instruments

The eight-question test described here is available as a free, self-contained instrument that runs offline in any web browser, and as a spreadsheet that scores up to three competing propositions side by side, both released under an open licence. A five-question quick screen that triages propositions before they earn this conversation accompanies them. Current links and DOIs for all companion instruments are maintained at auctusagri.com.

References

- AgFunder. 2025. Global AgriFoodTech Investment Report 2025. San Francisco, CA: AgFunder. <https://agfunder.com/research/agfunder-global-agrifoodtech-investment-report-2025/>
- Dara, R., S.M. Hazrati Fard, and J. Kaur. 2022. "Recommendations for Ethical and Responsible Use of Artificial Intelligence in Digital Agriculture." *Frontiers in Artificial Intelligence* 5:884192. <https://doi.org/10.3389/frai.2022.884192>
- Gillespie, J., E. Njuki, and A. Terán. 2024. Structure, Costs, and Technology Used on U.S. Dairy Farms. Report ERR-334. Washington, DC: U.S. Department of Agriculture, Economic Research Service. <https://doi.org/10.32747/2024.8534118.ers>
- Griliches, Z. 1957. "Hybrid Corn: An Exploration in the Economics of Technological Change." *Econometrica* 25(4):501–522. <https://doi.org/10.2307/1905380>
- Hankins, W., J. Kim, and A.H. Villacis. 2026. "Automation or Augmentation? AI and the Future of American Farming." *Choices* 41(2). <https://www.choicesmagazine.org/choices-magazine/submitted-articles/automation-or-augmentation-ai-and-the-future-of-american-farming>
- John Deere. 2024. "See & Spray Delivers 59% Herbicide Savings on More Than One Million Acres." News release. Moline, IL: Deere & Company. <https://www.deere.com/en-us/john-deere-news/see-spray-59-percent-herbicide-savings>

- Lowenberg-DeBoer, J., and B. Erickson. 2019. "Setting the Record Straight on Precision Agriculture Adoption." *Agronomy Journal* 111(4):1552-1569.
<https://doi.org/10.2134/agronj2018.12.0779>
- McFadden, J., E. Njuki, and T. Griffin. 2023. Precision Agriculture in the Digital Era: Recent Adoption on U.S. Farms. *Economic Information Bulletin* 248. Washington, DC: U.S. Department of Agriculture, Economic Research Service.
<https://doi.org/10.22004/ag.econ.333550>
- Michels, M., and L. Isenberg. 2026. "Developing a Farmer-Centred Framework for Assessing Adoption Readiness of Digital Agricultural Technologies: Validation with German Farmers." *Agricultural Systems* 237:104825.
<https://doi.org/10.1016/j.agsy.2026.104825>
- MIT Project NANDA. 2025. *The GenAI Divide: State of AI in Business 2025*. Cambridge, MA: Massachusetts Institute of Technology, Media Lab.
- Onel, G., F. Brito, J. Gars, and C. Mullally. 2025. "Sensing, Thinking, Doing: AI's Growing Role on the Farm—and What It Means for Farm Work." *Choices* 40(3).
<https://doi.org/10.22004/ag.econ.362690>
- Sculley, D., G. Holt, D. Golovin, E. Davydov, T. Phillips, D. Ebner, V. Chaudhary, M. Young, J.-F. Crespo, and D. Dennison. 2015. "Hidden Technical Debt in Machine Learning Systems." In *Advances in Neural Information Processing Systems* 28, pp. 2503-2511.
- Thomasson, J.A., Y. Ampatzidis, M. Bhandari, R.A. Ferreyra, T. Gentimis, E. McReynolds, S.C. Murray, M.B. Peterson, C.M. Rodriguez Lopez, R.L. Strong, L.O. Tedeschi, J. Vitale, and X. Ye. 2025. *AI in Agriculture: Opportunities, Challenges, and Recommendations*. Ames, IA: Council for Agricultural Science and Technology.
<https://doi.org/10.62300/IAAG042514>
- U.S. Government Accountability Office. 2024. *Precision Agriculture: Benefits and Challenges for Technology Adoption and Use*. GAO-24-105962. Washington, DC: U.S. Government Accountability Office. <https://www.gao.gov/products/gao-24-105962>
- Wiseman, L., J. Sanderson, A. Zhang, and E. Jakku. 2019. "Farmers and Their Data: An Examination of Farmers' Reluctance to Share Their Data through the Lens of the Laws Impacting Smart Farming." *NJAS - Wageningen Journal of Life Sciences* 90-91:100301. <https://doi.org/10.1016/j.njas.2019.04.007>
- Yeo, M.L., and C.M. Keske. 2024. "From Profitability to Trust: Factors Shaping Digital Agriculture Adoption." *Frontiers in Sustainable Food Systems* 8:1456991.
<https://doi.org/10.3389/fsufs.2024.1456991>