

Η Ευρωπαϊκή Ένωση αλλάζει τους κανόνες της AI από τις 2 Αυγούστου 2026

Ο πλήρης πρακτικός οδηγός του EU AI Act για AI προϊόντα, agents, blogs και έργα πελατών

Ενημερώθηκε: 2 Αυγούστου 2026

Από τις 2 Αυγούστου 2026 οι επιχειρήσεις που προσφέρουν ή χρησιμοποιούν AI στην Ευρωπαϊκή Ένωση αποκτούν νέες υποχρεώσεις διαφάνειας. Τι πρέπει να γράφει ένα chatbot, πώς σημαίνεται το συνθετικό περιεχόμενο, πότε ένα AI άρθρο χρειάζεται ετικέτα, τι ισχύει για deepfakes, source code και white-label έργα, και ποιος θεωρείται πραγματικά provider ή deployer.

02.08.2026

Έναρξη εφαρμογής

4 πυλώνες

Disclosure · marking · deepfakes · public text

30 ημέρες

Εφαρμόσιμο πλάνο συμμόρφωσης

Πρακτικός ενημερωτικός οδηγός
Δεν αποτελεί εξατομικευμένη νομική συμβουλή

Από την ομάδα του :

< **DevsClub** >

Η μεγάλη εικόνα σε 90 δευτερόλεπτα

Από τις **2 Αυγούστου 2026** το άρθρο 50 του ευρωπαϊκού AI Act εφαρμόζεται πλήρως. Για χιλιάδες μικρές και μεγάλες επιχειρήσεις, το πρώτο πρακτικό αποτέλεσμα δεν είναι ότι «απαγορεύεται η AI», αλλά ότι πρέπει να γίνεται σαφέστερο **πότε ένας άνθρωπος συνομιλεί με AI, πότε ένα αποτέλεσμα δημιουργήθηκε ή αλλοιώθηκε τεχνητά και πότε ένα δημοσιευμένο κείμενο ή ένα deepfake χρειάζεται εμφανή γνωστοποίηση**.

Οι κανόνες δεν αφορούν μόνο τους κατασκευαστές μεγάλων μοντέλων. Μπορούν να αγγίξουν μια εταιρεία που κυκλοφορεί ένα chatbot ή έναν AI agent με το δικό της brand, έναν εκδότη που χρησιμοποιεί AI στη ροή παραγωγής άρθρων, ένα agency που κατασκευάζει custom ή white-label λύσεις και έναν επαγγελματία που δημοσιεύει συνθετικές εικόνες, ήχο ή βίντεο.

Το κρίσιμο είναι να μη γίνουν τρία συνηθισμένα λάθη:

- Να θεωρήσουμε ότι την ευθύνη την έχει πάντοτε ο πάροχος του foundation model.
- Να βάλουμε μία γενική φράση στους όρους χρήσης και να θεωρήσουμε ότι καλύψαμε κάθε υποχρέωση.
- Να αντιμετωπίσουμε κάθε AI output με την ίδια ετικέτα, χωρίς να ξεχωρίζουμε provider, deployer, modality, κοινό και σκοπό δημοσίευσης.

Η σύντομη απάντηση: ένα επώνυμο AI προϊόν χρειάζεται συνήθως disclosure από την πρώτη άμεση αλληλεπίδραση με άνθρωπο και μηχανικά ανιχνεύσιμη σήμανση για τα in-scope συνθετικά outputs. Ο επαγγελματίας που δημοσιεύει deepfakes ή AI κείμενο δημοσίου ενδιαφέροντος έχει ξεχωριστές, ανθρώπινα αντιληπτές υποχρεώσεις. Το ποιος ευθύνεται δεν το αποφασίζει μόνο η σύμβαση· προκύπτει από το τι κάνει πραγματικά κάθε μέρος.

Τι αλλάζει, με μια ματιά

Περίπτωση	Ποιος εξετάζεται	Βασική υποχρέωση από 2/8/2026	Πρακτικό πρώτο βήμα
Chatbot, voice bot ή AI agent που μιλά απευθείας με ανθρώπους	Provider του τελικού AI system	Καθαρή ενημέρωση ότι ο χρήστης αλληλεπιδρά με AI, το αργότερο στην πρώτη αλληλεπίδραση	Προσθέστε first-turn ή opening disclosure σε κάθε ανθρώπινο κανάλι
Σύστημα που παράγει ή αλλοιώνει κείμενο, εικόνα, ήχο ή βίντεο	Provider, περιλαμβανομένων ορισμένων downstream providers	Machine-readable marking και πραγματική detectability των in-scope outputs	Χαρτογραφήστε modalities, exports και upstream metadata
Χρήση emotion recognition ή	Deployer	Ενημέρωση των προσώπων που εκτίθενται στο σύστημα	Ελέγξτε αν υπάρχει τέτοιο feature ή SDK στην

Περίπτωση	Ποιος εξετάζεται	Βασική υποχρέωση από 2/8/2026	Πρακτικό πρώτο βήμα
biometric categorisation			παραγωγή
Δημοσίευση deepfake	Deployer	Εμφανής ή ακουστή γνωστοποίηση ότι το υλικό είναι τεχνητά δημιουργημένο ή αλλοιωμένο	Βάλτε label στο σημείο πρώτης έκθεσης, όχι μόνο σε metadata
Δημοσίευση AI κειμένου για θέμα δημοσίου ενδιαφέροντος	Deployer/εκδότης	Disclosure, εκτός αν υπάρχει ουσιαστικός ανθρώπινος έλεγχος και πραγματική editorial responsibility	Ορίστε human sign-off πριν από τη δημοσίευση

1. Τι ακριβώς τέθηκε σε εφαρμογή στις 2 Αυγούστου 2026

Ο AI Act είναι ο [Κανονισμός \(ΕΕ\) 2024/1689](#). Το άρθρο 50 είναι το κεντρικό καθεστώς διαφάνειας για ορισμένα AI systems και AI-generated ή AI-manipulated outputs. Η εφαρμογή του ξεκινά στις 2 Αυγούστου 2026.

Λίγες ημέρες πριν από αυτή την ημερομηνία τέθηκε σε ισχύ και ο [Κανονισμός \(ΕΕ\) 2026/1744](#), γνωστός ως AI Omnibus. Αυτό έχει σημασία επειδή αρκετές παλαιότερες αναλύσεις στο διαδίκτυο μιλούν ακόμη για «πρόταση» ή χρησιμοποιούν παλιότερο χρονοδιάγραμμα. Στις 2 Αυγούστου 2026, όμως, οι νέες ημερομηνίες αποτελούν πλέον ισχύον δίκαιο.

Το άρθρο 50 είναι κυρίως **καθεστώς διαφάνειας**, όχι γενική απαγόρευση παραγωγής περιεχομένου με AI. Αυτό δεν σημαίνει ότι κάθε χρήση γίνεται νόμιμη μόλις προστεθεί μια ετικέτα. Εξακολουθούν να ισχύουν κανόνες για προσωπικά δεδομένα, πνευματική ιδιοκτησία, καταναλωτική προστασία, δυσφήμιση, απαγορευμένες πρακτικές AI και ειδικούς τομείς. Η ετικέτα είναι μία υποχρέωση, όχι «άδεια χρήσης».

Το ενημερωμένο χρονοδιάγραμμα

Ημερομηνία	Τι συμβαίνει	Γιατί ενδιαφέρει μια επιχείρηση
2 Φεβρουαρίου 2025	Εφαρμογή ορισμών, απαγορευμένων πρακτικών και AI literacy	Οι ομάδες που χρησιμοποιούν AI πρέπει ήδη να διαθέτουν κατάλληλη γνώση και επίγνωση κινδύνων
2 Αυγούστου 2025	Έναρξη σημαντικών κανόνων για GPAI models, governance και penalties	Οι υποχρεώσεις του model provider δεν είναι ίδιες με εκείνες του τελικού AI product
2 Αυγούστου 2026	Πλήρης εφαρμογή του Article 50	Chatbot disclosure, marking/detectability, deepfake και public-interest text disclosures ξεκινούν εδώ

Ημερομηνία	Τι συμβαίνει	Γιατί ενδιαφέρει μια επιχείρηση
2 Δεκεμβρίου 2026	Λήξη της ειδικής παράτασης του Article 50(2) για generative systems που είχαν ήδη τεθεί στην αγορά πριν από 2/8/2026	Δεν είναι γενική περίοδος χάριτος για όλο το Article 50
2 Δεκεμβρίου 2027	Εφαρμογή των κύριων high-risk κανόνων για συστήματα του Annex III	Αφορά, μεταξύ άλλων, συγκεκριμένες χρήσεις σε εργασία, εκπαίδευση, κρίσιμες υπηρεσίες και δικαιοσύνη
2 Αυγούστου 2028	Εφαρμογή high-risk κανόνων για συστήματα του Article 6(1)/Annex I	Αφορά AI που λειτουργεί ως προϊόν ή safety component σε ρυθμιζόμενα προϊόντα

Η παγίδα της «τετράμηνης παράτασης»

Η νέα μεταβατική ρύθμιση είναι στενή: αφορά **μόνο** την υποχρέωση machine-readable marking και detectability του Article 50(2), για παρόχους generative AI systems που είχαν τεθεί στην αγορά πριν από τις 2 Αυγούστου 2026. Αυτά τα συστήματα έχουν έως τις 2 Δεκεμβρίου 2026 για τη συγκεκριμένη συμμόρφωση.

Δεν μετατίθεται μέχρι τον Δεκέμβριο:

- η ενημέρωση ότι ο χρήστης συνομιλεί με AI,
- η ενημέρωση για emotion recognition ή biometric categorisation,
- η εμφανής γνωστοποίηση deepfake,
- η γνωστοποίηση για AI-generated public-interest text,
- η απαίτηση οι γνωστοποιήσεις να είναι σαφείς, διακριτές και προσβάσιμες.

Επίσης, η επίσημη καθοδήγηση διευκρινίζει ότι περιεχόμενο που δημιουργήθηκε πριν από τις 2 Αυγούστου 2026 δεν χρειάζεται αναδρομικά νέο provider mark. Αν όμως ένα παλαιότερα δημιουργημένο AI κείμενο δημοσιεύεται μετά την έναρξη εφαρμογής για να ενημερώσει το κοινό πάνω σε θέμα δημοσίου ενδιαφέροντος, η υποχρέωση του deployer εξετάζεται κατά τον χρόνο της δημοσίευσης.

2. Provider, deployer και downstream provider: ο ρόλος προηγείται της υποχρέωσης

Πριν εξετάσουμε τι label χρειάζεται, πρέπει να βρούμε **ποιος έχει ποιον νομικό ρόλο**.

- **Provider** είναι, σε γενικές γραμμές, το πρόσωπο ή η εταιρεία που αναπτύσσει —ή αναθέτει να αναπτυχθεί— ένα AI system και το διαθέτει στην αγορά ή το θέτει σε λειτουργία με το δικό του όνομα ή σήμα, δωρεάν ή επί πληρωμή.
- **Deployer** είναι αυτός που χρησιμοποιεί επαγγελματικά ένα AI system υπό τη δική του εξουσία. Η καθαρά προσωπική, μη επαγγελματική χρήση εξαιρείται από τον συγκεκριμένο ορισμό.

- **Downstream provider** είναι ο provider ενός AI system που ενσωματώνει AI model τρίτου ή δικό του.
- **Importer** και **distributor** είναι ξεχωριστοί ρόλοι της αλυσίδας διάθεσης, οι οποίοι μπορεί επίσης να είναι κρίσιμοι σε προϊόντα τρίτων χωρών.

Οι ρόλοι δεν είναι αμοιβαία αποκλειόμενοι. Η ίδια εταιρεία μπορεί να είναι provider ενός branded agent, deployer ενός εξωτερικού εργαλείου δημιουργίας εικόνων και customer ενός άλλου SaaS.

Γιατί το API του μεγάλου μοντέλου δεν «παίρνει μαζί του» όλη την ευθύνη

Ας υποθέσουμε ότι μια μικρή εταιρεία δημιουργεί ένα AI προϊόν με δικό της interface, δικές της ροές, δικό της brand και δικούς της πελάτες, αλλά χρησιμοποιεί API τρίτου μοντέλου. Ο προμηθευτής του μοντέλου παραμένει provider του μοντέλου και έχει τις δικές του υποχρεώσεις. Η μικρή εταιρεία, όμως, μπορεί να είναι provider του **τελικού AI system** που διαθέτει στους χρήστες.

Η διάκριση αυτή είναι θεμελιώδης:

Το «δεν εκπαιδεύω δικό μου μοντέλο» δεν σημαίνει «δεν είμαι provider κανενός AI system».

Το ίδιο ισχύει όταν ο πελάτης εισάγει δικό του API key, όταν χρησιμοποιείται open-source model ή όταν ο vendor αποκαλεί το προϊόν «wrapper». Ο νομικός χαρακτηρισμός κοιτάζει τα πραγματικά γεγονότα: ανάπτυξη, ανάθεση, branding, διάθεση, intended purpose και έλεγχο της χρήσης.

Πίνακας ρόλων για συνηθισμένα έργα

Σενάριο	Πιθανός ρόλος	Τι πρέπει να ελεγχθεί
Εταιρεία διαθέτει chatbot/agent με δικό της brand και τρίτο model API	Downstream provider του τελικού system	Ποια features ελέγχει, ποια outputs παράγει, πώς ενημερώνονται οι χρήστες
Εταιρεία αγοράζει τρίτο SaaS για εσωτερική παραγωγή περιεχομένου	Deployer	Σκοπός χρήσης, κοινό, ανθρώπινος έλεγχος, τύπος δημοσίευσης
Agency πουλά δικό του branded AI προϊόν σε πελάτη	Το agency μπορεί να είναι provider· ο πελάτης deployer	Αν ο πελάτης απλώς χρησιμοποιεί ή επαναδιαθέτει/αλλάζει το σύστημα
Πελάτης αναθέτει white-label σύστημα και το διαθέτει με το brand του	Ο πελάτης μπορεί να είναι provider και deployer	Ποιος ανέθεσε την ανάπτυξη, με ποιο όνομα τίθεται σε λειτουργία, ποιος αποφασίζει τον σκοπό
Developer παρέχει reusable engine/component κάτω από white-label λύση	Ενδέχεται να παραμένει provider άλλου system ή component	Τι ακριβώς διαθέτει ο developer και με ποιο brand

Σενάριο	Πιθανός ρόλος	Τι πρέπει να ελεγχθεί
EU reseller ξένου branded AI προϊόντος	Πιθανός importer ή distributor	Ποιος κάνει την πρώτη διάθεση στην ΕΕ και ποιος είναι ο αρχικός provider

Η σύμβαση βοηθά, αλλά δεν ξαναγράφει τον νόμο

Σε custom και white-label έργα, η σύμβαση πρέπει να περιγράφει με ακρίβεια ποιος υλοποιεί disclosure, marking, detection, human review, χειρισμό παραπόνων και συνεργασία με τις αρχές. Δεν μπορεί όμως να μετατρέψει έναν πραγματικό provider σε «απλό τεχνικό προμηθευτή» μόνο επειδή έτσι γράφει μια παράγραφος.

Ένα καλό συμβατικό παράρτημα AI θα πρέπει να καλύπτει τουλάχιστον:

- το mapping system, model, provider, deployer, importer και supplier,
- ποιος ελέγχει το intended purpose και τις ουσιώδεις αλλαγές,
- ποιος εμφανίζει το disclosure της πρώτης αλληλεπίδρασης,
- ποιος εξασφαλίζει και δοκιμάζει machine-readable marking και detection,
- ποιος διατηρεί upstream metadata κατά export, transcoding ή επεξεργασία,
- ποιος βάζει visible labels για deepfakes ή public-interest text,
- τι evidence και τεχνική τεκμηρίωση παρέχει κάθε vendor,
- πώς αντιμετωπίζονται παράπονα, incidents και αιτήματα αρχών,
- τι συμβαίνει όταν αλλάζει model, modality ή product flow.

3. Οι τέσσερις μεγάλες υποχρεώσεις του Article 50

Το άρθρο 50 δεν επιβάλλει ένα ενιαίο «AI label». Χωρίζει τις υποχρεώσεις ανά actor και χρήση.

Διάταξη	Υπόχρεος	Τι απαιτεί	Σημαντική εξαίρεση ή όριο
Article 50(1)	Provider	Σύστημα που προορίζεται για άμεση αλληλεπίδραση με φυσικά πρόσωπα να τους ενημερώνει ότι αλληλεπιδρούν με AI	Όταν αυτό είναι αντικειμενικά προφανές στο συγκεκριμένο context· η εξαίρεση ερμηνεύεται στενά
Article 50(2)	Provider	In-scope synthetic ή manipulated text, audio, image και video να φέρει machine-readable mark και να είναι ανιχνεύσιμο	Standard editing, μη ουσιώδης αλλαγή και ορισμένες κατηγορίες outputs που εξειδικεύουν οι Guidelines
Article 50(3)	Deployer	Ενημέρωση των προσώπων που εκτίθενται σε emotion recognition ή biometric categorisation	Ειδική εξαίρεση για ορισμένες νόμιμες χρήσεις επιβολής του νόμου
Article 50(4)	Deployer	Εμφανής γνωστοποίηση για deepfakes και για συγκεκριμένο	Για το κείμενο, μπορεί να ισχύει εξαίρεση με

Διάταξη	Υπόχρεος	Τι απαιτεί	Σημαντική εξαίρεση ή όριο
		AI-generated public-interest text	ουσιαστικό human review/editorial control και editorial responsibility

Όλες οι γνωστοποιήσεις των παραγράφων 1 έως 4 πρέπει, σύμφωνα με το Article 50(5), να είναι **σαφείς, διακριτές, προσβάσιμες και διαθέσιμες το αργότερο κατά την πρώτη αλληλεπίδραση ή έκθεση**.

Από εδώ προκύπτει μια χρήσιμη διάκριση:

- Το machine-readable mark του provider απευθύνεται κυρίως σε συστήματα και εργαλεία ανίχνευσης.
- Το visible ή audible disclosure του deployer απευθύνεται στον άνθρωπο που βλέπει, ακούει ή διαβάζει το αποτέλεσμα.

Σε αρκετές περιπτώσεις χρειάζονται **και τα δύο**. Το κρυφό metadata δεν αντικαθιστά την εμφανή ετικέτα ενός deepfake, και μια ορατή λεζάντα δεν αντικαθιστά αυτομάτως το machine-readable marking του συστήματος παραγωγής.

4. Περίπτωση A: branded chatbot ή AI agent

Ένα AI προϊόν μπορεί να συνομιλεί, να απαντά σε πελάτες, να γράφει και να στέλνει μηνύματα, να κλείνει ραντεβού, να πραγματοποιεί φωνητικές κλήσεις ή να ενεργεί σε εξωτερικά συστήματα. Για τη συμμόρφωση, πρέπει να ξεχωρίσουμε την **άμεση ανθρώπινη αλληλεπίδραση** από τις καθαρά backend ενέργειες.

Πότε ενεργοποιείται το disclosure της πρώτης αλληλεπίδρασης

Η επίσημη καθοδήγηση χρησιμοποιεί τέσσερα σωρευτικά στοιχεία:

1. Υπάρχει AI system.
2. Το σύστημα έχει σχεδιαστεί για πραγματική αμφίδρομη ανταλλαγή.
3. Το AI επικοινωνεί απευθείας, χωρίς άνθρωπο που μεσολαβεί και παρουσιάζει την απάντηση ως δική του.
4. Ο συνομιλητής είναι φυσικό πρόσωπο.

Αυτό περιλαμβάνει κλασικό web chat και voice bot, αλλά μπορεί να περιλαμβάνει και agent που ξεκινά επικοινωνία με τρίτο πρόσωπο για booking, διαπραγμάτευση ή αλληλογραφία. Σε αυτή την περίπτωση, οι Guidelines προτείνουν να γίνεται σαφές τόσο ότι πρόκειται για AI όσο και για λογαριασμό ποιου ενεργεί.

Δεν υπάρχουν αυτομάτως στο Article 50(1):

- ένα backend API request,
- επικοινωνία agent-to-agent χωρίς ανθρώπινη έκθεση,
- machine-to-machine data exchange,
- μια απλή μονόδρομη αυτοματοποιημένη ειδοποίηση χωρίς πραγματική συνομιλία,
- εσωτερική επεξεργασία που παραδίδεται σε άνθρωπο-χειριστή, ο οποίος αποφασίζει πώς θα απαντήσει.

Η διάκριση έχει σημασία. Η φράση «κάθε email που συντάσσει agent πρέπει να γράφει υποχρεωτικά ότι είναι AI» είναι υπερβολική. Όταν όμως ο agent συμμετέχει σε πραγματική άμεση αλληλεπίδραση, το disclosure πρέπει να αποτελεί μέρος της αρχιτεκτονικής του καναλιού και όχι χειροκίνητη επιλογή του χρήστη.

Τι θεωρείται καλό disclosure

Ο Κανονισμός δεν επιβάλλει μία συγκεκριμένη πρόταση. Το μήνυμα πρέπει να είναι σαφές για έναν φυσιολογικό χρήστη, όχι κρυμμένο σε όρους χρήσης ή σε τεχνικό metadata.

Ενδεικτικές διατυπώσεις:

- «Συνομιλείτε με σύστημα τεχνητής νοημοσύνης.»
- «Είμαι AI βοηθός και απαντώ για λογαριασμό της [εταιρείας/ομάδας].»
- Σε voice channel: «Η κλήση εξυπηρετείται από σύστημα τεχνητής νοημοσύνης.»
- Σε email thread με agent: εμφανής ένδειξη στην αρχή της πρώτης άμεσης ανταλλαγής.

Μία σαφής ενημέρωση στην έναρξη συνήθως μπορεί να αρκεί. Persistent badge ή περιοδική υπενθύμιση είναι ισχυρές UX πρακτικές, ειδικά σε μεγάλη διάρκεια, ευαίσθητο θέμα ή περιβάλλον όπου το AI μοιάζει ιδιαίτερα ανθρώπινο, αλλά δεν αποτελούν γενικό, άκαμπτο format του νόμου.

Τι συνήθως δεν αρκεί

- Μια γενική φράση «ο ιστότοπος χρησιμοποιεί AI» στο footer.
- Αναφορά μόνο στους όρους χρήσης.
- Η λέξη «assistant» χωρίς να δηλώνεται ότι είναι AI.
- Τεχνική φράση «powered by LLM» που δεν είναι σαφής στο ευρύ κοινό.
- Κρυφό metadata που ο άνθρωπος δεν βλέπει.

Η εξαίρεση «είναι προφανές» είναι στενή

Σε ένα εργαλείο coding που χρησιμοποιείται από εκπαιδευμένους developers, σε περιβάλλον όπου η φύση της λειτουργίας είναι αδιαμφισβήτητη, μπορεί να είναι προφανές ότι ο χρήστης αλληλεπιδρά με AI. Σε δημόσιο helpdesk με φυσική γλώσσα και ανθρώπινο τόνο, η ίδια υπόθεση είναι πολύ πιο αδύναμη.

Η ασφαλής σχεδιαστική απόφαση είναι απλή: όταν το disclosure κοστίζει μία καθαρή γραμμή στην αρχή, δεν αξίζει να στηριχθεί το προϊόν σε μια αμφισβητήσιμη εξαίρεση.

Checklist για conversational AI

- Καταγραφή όλων των καναλιών: web, mobile, email, voice, messaging.
- First-interaction disclosure ανά κανάλι και γλώσσα.
- Δήλωση του principal όταν agent επικοινωνεί με τρίτο για λογαριασμό κάποιου.
- Accessibility review για visual και audible notices.
- Έλεγχος ότι το disclosure δεν χάνεται σε embedded, white-label ή API-driven έκδοση.
- Regression test μετά από αλλαγή UI, provider ή conversational flow.
- Ξεχωριστή εξέταση των outputs για Article 50(2), πέρα από το chatbot notice.

5. Machine-readable marking: το δυσκολότερο τεχνικό κομμάτι

Το Article 50(2) απαιτεί από providers συστημάτων που παράγουν synthetic audio, image, video ή text να εξασφαλίζουν ότι τα in-scope outputs:

- φέρουν machine-readable marking,
- μπορούν να ανιχνευθούν ως τεχνητά δημιουργημένα ή αλλοιωμένα,
- χρησιμοποιούν λύσεις αποτελεσματικές, διαλειτουργικές, ανθεκτικές και αξιόπιστες,
- στον βαθμό που αυτό είναι τεχνικά εφικτό, με βάση τη modality, το κόστος και την κατάσταση της τεχνικής.

Ο νόμος είναι τεχνολογικά ουδέτερος. Δεν επιβάλλει ένα συγκεκριμένο JSON schema, ένα μοναδικό watermark, ένα public detector API ή πλήρη αλυσίδα provenance για κάθε output.

Marking και detection δεν είναι το ίδιο

Έννοια	Τι σημαίνει πρακτικά	Παραδείγματα τεχνικών
Machine-readable marking	Το output μεταφέρει σήμα που μπορεί να διαβαστεί από μηχανή	Signed metadata, cryptographic credentials, embedded watermark, container metadata
Detectability	Υπάρχει λειτουργικός τρόπος να διαπιστωθεί ότι το περιεχόμενο είναι AI-generated/manipulated	Detector που επαληθεύει credential/watermark, registry lookup, fingerprint matching
Human-facing disclosure	Ο άνθρωπος ενημερώνεται εμφανώς ή ακουστά	Caption, badge, on-screen notice, spoken announcement

Ένα απλό database log που δεν συνδέεται αξιόπιστα με το αρχείο μπορεί να βοηθά στο audit, αλλά δεν ισοδυναμεί αυτομάτως με machine-readable mark πάνω ή μέσα στο output. Αντίστροφα, metadata που αφαιρείται στο πρώτο export ή upload είναι ανεπαρκές ως μοναδική λύση.

Μπορώ να βασιστώ στον upstream provider;

Ναι, ένας downstream provider μπορεί να χρησιμοποιήσει συμμορφούμενη λύση του upstream model/system provider ή τρίτου. Αυτό είναι συχνά η πιο λογική αρχιτεκτονική. Πρέπει όμως να επιβεβαιώσει ότι το σήμα:

- υπάρχει στις modalities και στα endpoints που χρησιμοποιεί,
- επιβιώνει στις δικές του μετατροπές, exports και pipelines,
- είναι πραγματικά ανιχνεύσιμο,
- δεν αφαιρείται από resize, compression, transcoding ή re-encoding χωρίς εναλλακτική,
- καλύπτει το τελικό branded system και όχι μόνο ένα demo του upstream vendor.

Η ευθύνη για το τελικό system-level αποτέλεσμα δεν εξαφανίζεται επειδή ο vendor υποσχέθηκε «AI metadata». Ο provider πρέπει να μπορεί να δείξει τι δοκίμασε και τι συμβαίνει όταν το upstream mark λείπει ή αποτυγχάνει.

Τι δεν απαιτεί ρητά το Article 50

Ένα εσωτερικό provenance record με output ID, model version, timestamp, content hash, prompt reference, operator και test result είναι εξαιρετική πρακτική. Δεν αποτελεί όμως αυτούσιο statutory checklist. Το ίδιο ισχύει για ένα public verification endpoint ή για συγκεκριμένη υλοποίηση adapter/fallback.

Η σωστή διατύπωση είναι:

Χρειάζεται επαρκές, δοκιμασμένο αποτέλεσμα marking και detectability. Η αρχιτεκτονική μπορεί να είναι upstream, system-level, third-party ή συνδυαστική, αρκεί να καλύπτει τις απαιτήσεις και να μπορεί να τεκμηριωθεί.

Ο εθελοντικός Code of Practice

Ο [Code of Practice on Transparency of AI-generated Content](#) είναι προαιρετικός, όχι δεύτερος νόμος. Προσφέρει, όμως, αναγνωρισμένο πλαίσιο τεχνικών και οργανωτικών μέτρων για όσους τον υπογράφουν.

Για εικόνα, ήχο και βίντεο προωθεί συνδυασμό τεχνικών, επειδή κανένα μεμονωμένο μέσο δεν είναι άτρωτο. Για ελεύθερο κείμενο εξετάζει text watermarking σε μεγαλύτερα outputs. Περιλαμβάνει επίσης δοκιμές, monitoring, διατήρηση upstream signals, διαδικασίες συμμόρφωσης και πρόσβαση σε αποτελέσματα detection με τρόπο κατανοητό.

Μια εταιρεία που δεν υπογράφει τον Code δεν παρανομεί γι' αυτό και μόνο. Πρέπει, όμως, να μπορεί να αποδείξει ότι συμμορφώνεται με το Article 50 μέσω άλλων επαρκών μέσων.

6. Τι γίνεται με source code, SQL, JSON, YAML και agent outputs

Οι τελικές Guidelines της Επιτροπής ξεκαθαρίζουν μια κρίσιμη πρακτική λεπτομέρεια: ορισμένα αυστηρά τεχνικά outputs βρίσκονται εκτός του πεδίου του Article 50(2). Σε αυτά περιλαμβάνονται source code σε γλώσσες programming, scripting, markup, query ή configuration, καθώς και τεχνικά στοιχεία που αποτελούν αναπόσπαστο μέρος του κώδικα.

Οι Guidelines αναφέρουν, μεταξύ άλλων:

- source code και ενσωματωμένα comments,
- SDKs, SQL και infrastructure-as-code,
- YAML και JSON configuration schemas,
- scripts, machine-readable specifications, APIs και software libraries,
- πολύ σύντομες ακολουθίες, single words, UI labels, alt text και icon-scale graphics,
- αποκλειστικά machine-to-machine outputs χωρίς ανθρώπινη έκθεση,
- agent-to-agent messages, web requests και intermediate agent processing που δεν προορίζονται να γίνουν αντιληπτά από άνθρωπο.

Αυτό είναι σημαντικό για coding agents και αυτοματισμούς. Δεν σημαίνει, όμως, ότι κάθε αρχείο με κατάληξη `.json`, κάθε README ή κάθε τεχνικό κείμενο εξαιρείται μαγικά.

Output	Πιθανή αντιμετώπιση	Γιατί
Source code, SQL query, YAML config	Συνήθως εκτός marking κατά τις Guidelines	Αυστηρά τεχνική μορφή που λειτουργεί ως code/configuration
JSON payload μεταξύ δύο services, χωρίς ανθρώπινη έκθεση	Συνήθως εκτός	Αποκλειστικά M2M και αυτόματη επεξεργασία
JSON report που εμφανίζεται αυτούσιο σε analyst	Χρειάζεται νέα αξιολόγηση	Υπάρχει ανθρώπινη έκθεση και ενδέχεται να λειτουργεί ως περιεχόμενο
README ή architecture report σε φυσική γλώσσα	Δεν είναι αυτομάτως source code	Είναι human-facing explanatory text
AI coding assistant που συνομιλεί με developer	Μπορεί να απαιτεί disclosure 50(1)	Η συνομιλία είναι διαφορετικό output από τον παραγόμενο κώδικα
Agent intermediate reasoning ή browser action που δεν βλέπει άνθρωπος	Εκτός της συγκεκριμένης ανθρώπινης έκθεσης	Δεν προορίζεται να παρουσιαστεί ως αντιληπτό περιεχόμενο

Υπάρχει επίσης στενή αντιμετώπιση για συγκεκριμένα B2B/industrial outputs, όταν είναι σωρευτικά αυστηρά τεχνικά, προσβάσιμα μόνο σε περιορισμένο προκαθορισμένο επαγγελματικό κοινό, εσωτερικά και προστατευμένα από εξωτερική διάδοση. Δεν είναι γενική εξαίρεση για κάθε B2B SaaS.

7. Περίπτωση B: blog, ενημερωτικό site και AI-assisted άρθρα

Ένα blog ή ενημερωτικό site που χρησιμοποιεί AI για έρευνα, draft, περίληψη, μετάφραση ή παραγωγή κειμένου είναι συνήθως **deployer** του εργαλείου. Η υποχρέωση visible disclosure για κείμενο ενεργοποιείται όταν συντρέχουν σωρευτικά τρία στοιχεία:

5. Το κείμενο έχει δημιουργηθεί ή αλλοιωθεί με AI.
6. Δημοσιεύεται με σκοπό να ενημερώσει το κοινό.
7. Αφορά θέμα δημοσίου ενδιαφέροντος.

Η έννοια του δημοσίου ενδιαφέροντος είναι ευρεία. Μπορεί να περιλαμβάνει πολιτική, δημόσια διοίκηση, δικαιοσύνη, θεμελιώδη δικαιώματα, ασφάλεια, δημόσια υγεία, περιβάλλον, προστασία καταναλωτή και οικονομικές, χρηματοπιστωτικές, επιστημονικές, τεχνολογικές ή πολιτιστικές εξελίξεις που τροφοδοτούν τον δημόσιο διάλογο.

Δεν είναι κάθε tutorial ή product description αυτομάτως public-interest text. Ένα τεχνικό how-to για μια βιβλιοθήκη μπορεί να μην έχει τον ίδιο χαρακτήρα με άρθρο για την ασφάλεια ενός AI συστήματος που χρησιμοποιείται σε νοσοκομεία. Η ταξινόμηση γίνεται με βάση θέμα, σκοπό, κοινό και context.

Η σημαντική εξαίρεση της ανθρώπινης επιμέλειας

Δεν απαιτείται το συγκεκριμένο AI label στο public-interest text όταν υπάρχουν σωρευτικά:

- **ουσιαστικός ανθρώπινος έλεγχος ή editorial control**, και
- φυσικό ή νομικό πρόσωπο που φέρει **editorial responsibility** για τη δημοσίευση.

Ουσιαστικός έλεγχος δεν είναι ένα γρήγορο skim, ορθογραφική διόρθωση, αυτόματη επαλήθευση από δεύτερο AI ή τυπικό πάτημα «publish». Ο reviewer πρέπει να έχει σχετική γνώση και επαγγελματική κρίση, να ελέγχει τουλάχιστον τα πραγματικά στοιχεία και να έχει εξουσία να αλλάξει ή να απορρίψει το κείμενο.

Κρίσιμη λεπτομέρεια: αν μετά το ανθρώπινο sign-off το AI κάνει ουσιώδη αλλαγή στο κείμενο, η βάση της εξαίρεσης μπορεί να χαθεί. Το τελικό ανθρώπινο review πρέπει να γίνεται **μετά** την τελευταία ουσιαστική AI παρέμβαση.

Ένα πρακτικό editorial workflow

8. **Classification:** Είναι το κείμενο AI-generated/manipulated; Δημοσιεύεται στο κοινό; Είναι public-interest;
9. **Source check:** Άνοιγμα των πρωτογενών πηγών, έλεγχος ημερομηνιών, όρων και κρίσιμων αριθμών.
10. **Substantive edit:** Δομή, λογική, ακρίβεια, ισορροπία, παραδείγματα και αποφυγή παραπλανητικών συμπερασμάτων.

11. **Legal/subject review όπου χρειάζεται:** Ιδίως για υγεία, χρηματοοικονομικά, νομικά, ασφάλεια ή ρυθμιζόμενους κλάδους.
12. **Final human sign-off:** Από πρόσωπο με πραγματική εξουσία έγκρισης, αλλαγής ή απόρριψης.
13. **Publication gate:** Καμία ουσιώδης AI αλλαγή μετά το sign-off χωρίς νέο review.
14. **Evidence:** Προαιρετική καταγραφή reviewer, ημερομηνίας, πηγών και έκδοσης ως καλή πρακτική.

Το Article 50 δεν απαιτεί ρητά per-article audit log με hash. Ένα τέτοιο αρχείο είναι, ωστόσο, χρήσιμο για εσωτερική συνέπεια, έλεγχο ποιότητας και απόδειξη ότι το editorial process ήταν πραγματικό.

Χρειάζεται κάθε AI-assisted άρθρο ετικέτα;

Όχι. Ένα άρθρο που πέρασε ουσιαστικό human review και έχει υπεύθυνο εκδότη μπορεί να εμπίπτει στην εξαίρεση. Επίσης, ένα καθαρά ιδιωτικό report προς συγκεκριμένο πελάτη ή μια προσωπική AI απάντηση συνήθως δεν είναι «δημοσίευση προς το κοινό» για το συγκεκριμένο σκέλος.

Προσοχή: αυτό δεν βγάζει το output από ολόκληρο το Article 50. Το provider-level marking του συστήματος παραγωγής ή το chatbot disclosure μπορεί να εξακολουθούν να ισχύουν.

Παραδείγματα εκδοτικής ταξινόμησης

Περίπτωση	Πιθανή αξιολόγηση	Ενέργεια
AI draft για άρθρο νομοθετικής αλλαγής, ουσιαστικά ελεγμένο και εγκεκριμένο από υπεύθυνο editor	Μπορεί να ισχύει η εξαίρεση	Τεκμηριωμένο human review και editorial responsibility
Πλήρως αυτοματοποιημένο news feed με AI summaries, χωρίς ουσιαστικό ανθρώπινο έλεγχο	Ισχυρή πιθανότητα disclosure	Label στο δημοσιευμένο κείμενο και provider-marking έλεγχος
Προσωπική απάντηση chatbot σε έναν αναγνώστη	Όχι public publication για 50(4) text	Ελέγξτε 50(1) interactive disclosure και 50(2) marking
Ιδιωτική έκθεση προς έναν πελάτη	Συνήθως όχι publication to the public	Μην θεωρήσετε ότι εξαιρείται αυτομάτως από όλα τα άλλα σκέλη
Product description ή διαφήμιση	Συνήθως όχι public-interest text	Επανεξέταση αν περιέχει health, safety ή sustainability claims
Τεχνολογικό άρθρο για AI που επηρεάζει εργασία, δικαιώματα ή ασφάλεια	Πιθανό matter of public interest	Human editorial workflow ή disclosure

8. Εικόνες, ήχος και βίντεο: πότε έχουμε deepfake

Δεν απαιτεί το Article 50(4) visible label για κάθε AI εικόνα. Η συγκεκριμένη υποχρέωση του deployer αφορά εικόνα, ήχο ή βίντεο που αποτελεί **deepfake**.

Η αξιολόγηση δεν σταματά στο ερώτημα «δείχνει γνωστό πραγματικό πρόσωπο;». Οι επίσημες Guidelines εξετάζουν αν:

15. το υλικό δημιουργήθηκε ή αλλοιώθηκε με AI,
16. παρουσιάζει επαρκή ομοιότητα με πρόσωπο, αντικείμενο, τόπο, οντότητα ή γεγονός που υπάρχει, μπορεί εύλογα να υπάρχει ή θα μπορούσε να έχει υπάρξει,
17. μπορεί να εμφανιστεί ψευδώς ως αυθεντικό ή αληθινό στο συγκεκριμένο context και κοινό.

Η πρόθεση εξαπάτησης δεν είναι απαραίτητο στοιχείο. Το context, η λεζάντα, το κανάλι, το ύφος και οι προσδοκίες του κοινού μπορούν να αλλάξουν την αξιολόγηση.

Η φωτορεαλιστική «stock» εικόνα δεν είναι αυτόματα ασφαλής

Μια εμφανώς αφηρημένη, εικονογραφημένη ή φανταστική σκηνή που κανείς δεν θα εκλάβει ως πραγματική καταγραφή συνήθως δεν πληροί τον ορισμό. Αντίθετα, μια φωτορεαλιστική συνθετική εικόνα που παρουσιάζει μια εύλογα πραγματική ιατρική πράξη, ένα παιδί σε θεραπεία, έναν χώρο εργασίας ή ένα περιστατικό σαν να φωτογραφήθηκε πραγματικά μπορεί να χρειάζεται case-by-case αξιολόγηση και να αποτελεί deepfake, ακόμη και αν δεν απεικονίζει επώνυμο πρόσωπο.

Παράδειγμα	Πιθανή κατάταξη	Ασφαλής πρακτική
Εμφανώς cartoon robot σε φανταστικό τοπίο	Συνήθως όχι deepfake	Δεν προκύπτει γενικό 50(4) label, αλλά διατηρήστε provider marks
Φωτορεαλιστικός «γιατρός» σε ανύπαρκτη κλινική που συνοδεύει πραγματική υπηρεσία	Μπορεί να εμφανιστεί ψευδώς αυθεντικό	Contextual review και, όπου πληρούται ο ορισμός, visible AI disclosure
Συνθετικό βίντεο γνωστού πολιτικού που λέει κάτι που δεν είπε	Κλασικό deepfake	Εμφανής/ακουστή γνωστοποίηση από την πρώτη έκθεση
Σατιρικό AI βίντεο με δημόσιο πρόσωπο	Παραμένει deepfake, αλλά ισχύει η ειδική καλλιτεχνική αντιμετώπιση	Disclosure με τρόπο που δεν εμποδίζει αδικαιολόγητα την απόλαυση του έργου
AI voice clone σε διαφημιστικό	Πιθανό deepfake και επιπλέον ζητήματα συναίνεσης/δικαιωμάτων	Audible/visible disclosure και χωριστός νομικός έλεγχος
Απλό color correction ή noise reduction χωρίς ουσιαστική αλλαγή	Μπορεί να εμπίπτει στο standard editing	Τεκμηριώστε ότι δεν άλλαξε ουσιαστικά το νόημα ή το περιεχόμενο

Για deepfakes, το metadata του provider δεν αρκεί μόνο του. Ο deployer χρειάζεται disclosure που ο άνθρωπος μπορεί να αντιληφθεί το αργότερο στην πρώτη έκθεση. Σε καλλιτεχνικά, δημιουργικά, σατιρικά ή μυθοπλαστικά έργα, η γνωστοποίηση μπορεί να είναι λιγότερο παρεμβατική, αλλά δεν εξαφανίζεται.

9. Standard editing: πότε μια AI παρέμβαση είναι μικρή και πότε ουσιώδης

Το Article 50(2) προβλέπει εξαίρεση όταν το AI εκτελεί λειτουργία υποβοήθησης για standard editing ή όταν δεν αλλοιώνει ουσιωδώς το input ή τη σημασία του.

Η επίσημη καθοδήγηση αντιμετωπίζει ως τυπικά μικρές παρεμβάσεις:

- ορθογραφική και γραμματική διόρθωση,
- μικρή γλωσσική εξομάλυνση,
- μετάφραση που δεν αλλάζει το νόημα,
- format conversion, compression ή noise reduction,
- περιορισμένο crop, resize ή color adjustment χωρίς ουσιώδη αλλαγή.

Αντίθετα, περίληψη, εκτεταμένη παράφραση, αλλαγή ύφους ή νοήματος, προσθήκη/αφαίρεση προσώπων και αντικειμένων ή κατασκευή νέας σκηνής δεν πρέπει να βαφτίζονται «standard editing» μόνο επειδή ξεκίνησαν από ανθρώπινο input.

Ένας χρήσιμος έλεγχος είναι:

Αν ο μέσος παραλήπτης θα σχημάτιζε διαφορετική εντύπωση για το περιεχόμενο, το γεγονός, το ύφος ή τον δημιουργό μετά την AI παρέμβαση, η αλλαγή πιθανότατα δεν είναι απλή τεχνική διόρθωση.

10. Περίπτωση Γ: custom και white-label AI έργα πελατών

Τα custom projects είναι η περιοχή όπου οι απλές γενικεύσεις αποτυγχάνουν συχνότερα. Μπορεί να υπάρχει developer, agency, integrator, upstream model provider, reseller και τελικός πελάτης — με περισσότερους από έναν ρόλους για κάθε μέρος.

Τέσσερα ενδεικτικά σενάρια

Σενάριο 1 — Το agency διαθέτει δικό του branded system

Το agency αναπτύσσει επαναχρησιμοποιήσιμη πλατφόρμα, τη διαθέτει με το δικό του όνομα και ο πελάτης τη χρησιμοποιεί. Το agency μπορεί να είναι provider του system και ο πελάτης deployer.

Σενάριο 2 — Ο πελάτης αναθέτει και λανσάρει με το δικό του brand

Ο πελάτης μπορεί να είναι provider επειδή ανέθεσε την ανάπτυξη και θέτει το system σε λειτουργία με το όνομά του. Αν το χρησιμοποιεί ο ίδιος, μπορεί να είναι ταυτόχρονα deployer. Ο developer μπορεί να είναι τεχνικός supplier, αλλά αυτό δεν είναι αυτόματο.

Σενάριο 3 — White-label front end πάνω σε reusable engine

Ο πελάτης εμφανίζεται ως provider του τελικού προϊόντος, αλλά ο developer συνεχίζει να διαθέτει ξεχωριστό underlying system ή component. Οι υποχρεώσεις πρέπει να χαρτογραφηθούν ανά σύστημα και όχι με μία ετικέτα για όλη τη σχέση.

Σενάριο 4 — Αλλαγή σκοπού ή ουσιώδους τροποποίηση

Σε high-risk περιβάλλοντα, rebranding, substantial modification ή αλλαγή intended purpose μπορεί να μεταφέρει ρόλο provider σε άλλο μέρος. Αυτό απαιτεί ειδικό έλεγχο του Article 25 και δεν λύνεται με ένα γενικό white-label template.

Responsibility matrix για κάθε έργο

Υποχρέωση/ενέργεια	Product owner	Developer/integrator	Upstream vendor	Deployer/editor
Χαρακτηρισμός ρόλων και intended purpose	Accountable με βάση τα facts	Παρέχει τεχνικά facts	Παρέχει model/system scope	Περιγράφει πραγματική χρήση
First-interaction disclosure	Ο provider εξασφαλίζει την ύπαρξή του	Υλοποιεί σε UI/voice/email flows	Παρέχει capabilities όπου υπάρχουν	Δεν το αφαιρεί ή παρακάμπτει
Machine-readable marking	Ο provider του τελικού system εξασφαλίζει συμμόρφωση	Ενσωματώνει και δοκιμάζει	Παρέχει marks/detection evidence	Διατηρεί marks στις δικές του ροές
Deepfake/public-text visible label	Υποστηρίζει feature/policy	Υλοποιεί publication controls	Παρέχει provenance όπου διαθέσιμο	Ο deployer αποφασίζει και εμφανίζει το label
Evidence και testing	Διατηρεί συνολικό φάκελο	Παραδίδει test results	Παρέχει documentation	Τεκμηριώνει review/publication decisions
Regulator/complaint handling	Ορίζει owner και SLA	Παρέχει τεχνική συνδρομή	Συνεργάζεται συμβατικά	Παρέχει χρήση και publication records

Ο πίνακας είναι λειτουργικό template, όχι αυτόματη νομική κατανομή. Η τελική ευθύνη ακολουθεί τους ρόλους που προκύπτουν από τον Κανονισμό.

11. Το Article 50 δεν κάνει ένα απλό chatbot «high-risk»

Η διαφάνεια του Article 50 και η κατάταξη high-risk είναι δύο διαφορετικοί άξονες. Ένα σύστημα δεν γίνεται high-risk επειδή χρησιμοποιεί ισχυρό μοντέλο, έχει agents ή παράγει κώδικα.

Η high-risk αξιολόγηση εξαρτάται κυρίως από το **intended purpose** και τις κατηγορίες των Annex I και Annex III. Ενδεικτικά, μπορεί να αφορά συγκεκριμένες χρήσεις σε:

- πρόσληψη, επιλογή, αξιολόγηση ή διαχείριση εργαζομένων,
- εισαγωγή, πρόσβαση ή αξιολόγηση στην εκπαίδευση,
- κρίσιμες υποδομές και safety components,
- πρόσβαση σε βασικές ιδιωτικές ή δημόσιες υπηρεσίες,
- ορισμένα συστήματα πιστοληπτικής αξιολόγησης ή ασφάλισης,
- law enforcement, migration, border control, justice και δημοκρατικές διαδικασίες,
- ορισμένη διαλογή επειγόντων περιστατικών υγείας.

Ένας γενικός coding assistant, ένα blog drafting tool ή ένα customer-service bot συνήθως δεν είναι high-risk μόνο λόγω της τεχνολογίας του. Αν όμως το ίδιο σύστημα επαναχρησιμοποιηθεί για screening υποψηφίων ή για απόφαση πρόσβασης σε ουσιώδη υπηρεσία, η αξιολόγηση αλλάζει.

Μετά τον Κανονισμό 2026/1744, οι κύριοι high-risk κανόνες για Annex III εφαρμόζονται από τις 2 Δεκεμβρίου 2027 και για Annex I product systems από τις 2 Αυγούστου 2028. Αυτές οι ημερομηνίες δεν αναβάλλουν το Article 50.

12. Enforcement, πρόστιμα και τι ισχύει στην Ελλάδα

Η παράβαση του Article 50 μπορεί να οδηγήσει σε διοικητικό πρόστιμο έως **15 εκατομμύρια ευρώ ή 3% του συνολικού παγκόσμιου ετήσιου κύκλου εργασιών του προηγούμενου οικονομικού έτους**. Για επιχειρήσεις εκτός της ειδικής μεταχείρισης μικρότερων οντοτήτων εφαρμόζεται το υψηλότερο ανώτατο όριο, ενώ για SMEs και SMCs το χαμηλότερο.

Αυτά είναι ανώτατα όρια, όχι αυτόματο τιμολόγιο. Η αρχή εξετάζει παράγοντες όπως βαρύτητα, διάρκεια, πρόθεση ή αμέλεια, μέγεθος, συνεργασία, προηγούμενες παραβάσεις και μέτρα αποκατάστασης.

Η εποπτεία μοιράζεται μεταξύ εθνικών αρχών και, για συγκεκριμένα συστήματα ή πλατφόρμες που εμπίπτουν στη δική του αρμοδιότητα, του AI Office. Στην Ελλάδα, ο **N. 5321/2026** όρισε την Αρχή Προστασίας Δεδομένων Προσωπικού Χαρακτήρα ως αρμόδια αρχή εποπτείας της αγοράς για το Article 50 και ως ενιαίο σημείο επαφής/καταγγελιών. Η ελληνική εφαρμογή παρουσιάζεται στην [επίσημη ανακοίνωση της Ειδικής Γραμματείας ΤΝ και Διακυβέρνησης Δεδομένων](#) και στο [ΦΕΚ Α' 114/20.07.2026](#).

Η πρακτική σημασία είναι ότι η συμμόρφωση δεν πρέπει να μείνει σε marketing copy. Χρειάζεται owner, τεχνική υλοποίηση, test evidence και διαδικασία διόρθωσης.

13. Ένα εφαρμόσιμο πλάνο συμμόρφωσης 30 ημερών

Εβδομάδα 1 — Inventory και ρόλοι

18. Καταγράψτε κάθε AI feature, model, API, modality και output channel.
19. Σημειώστε ποιος αναπτύσσει, ποιος αναθέτει, ποιος βάζει το brand και ποιος χρησιμοποιεί.
20. Διαχωρίστε provider, deployer, importer, distributor και supplier ανά σύστημα.
21. Καταγράψτε ποια outputs βλέπει άνθρωπος και ποια μένουν αποκλειστικά M2M.
22. Εντοπίστε συστήματα που είχαν ήδη τεθεί στην αγορά πριν από 2/8/2026.

Παραδοτέο: AI system register με owner, intended purpose, roles, modalities και ημερομηνία διάθεσης.

Εβδομάδα 2 — UX disclosures και publication policy

23. Προσθέστε disclosure στην πρώτη αλληλεπίδραση για chat, voice και εξωτερικά agent contacts.
24. Ελέγξτε ότι το μήνυμα είναι σαφές, ορατό/ακουστό και προσβάσιμο.
25. Δημιουργήστε deepfake decision gate για image/audio/video.
26. Δημιουργήστε public-interest text classification και human-review gate.
27. Ορίστε editorial responsible person ή entity.

Παραδοτέο: approved disclosure library και editorial/deepfake policy.

Εβδομάδα 3 — Marking, detection και preservation

28. Ζητήστε επίσημη τεκμηρίωση από κάθε upstream vendor.
29. Δοκιμάστε outputs από όλα τα endpoints και modalities.
30. Ελέγξτε αν το mark επιβιώνει σε export, compression, upload, resize και transcoding.
31. Επιβεβαιώστε ότι η detection διαδικασία παράγει κατανοητό αποτέλεσμα.
32. Σχεδιάστε system-level ή third-party fallback όπου το upstream δεν επαρκεί.

Παραδοτέο: compatibility matrix και test report ανά modality.

Εβδομάδα 4 — Συμβάσεις, evidence και incident flow

33. Ενημερώστε συμβάσεις με πραγματικό responsibility mapping.
34. Ορίστε ποιος διατηρεί metadata και ποιος απαγορεύεται να το αφαιρεί.
35. Δημιουργήστε evidence pack: screenshots, test cases, vendor docs, versions και αποφάσεις scope.

36. Ορίστε διαδικασία παραπόνων, regulator requests και corrective releases.
37. Προγραμματίστε επανέλεγχο σε κάθε αλλαγή model, provider, modality ή intended purpose.

Παραδοτέο: Article 50 compliance file και owner-approved remediation plan.

14. Decision tree: τι πρέπει να κάνω για κάθε feature ή output;

38. Ποιος είναι ο actor;

- Αναπτύσσει ή αναθέτει και διαθέτει/θέτει σε λειτουργία με δικό του brand → πιθανός provider.
- Χρησιμοποιεί επαγγελματικά υπό τη δική του authority → deployer.
- Μπορεί να είναι και τα δύο.

39. Υπάρχει άμεση αμφίδρομη επαφή AI με φυσικό πρόσωπο;

- Ναι → disclosure από την πρώτη αλληλεπίδραση, εκτός αν είναι αντικειμενικά προφανές.
- Ο agent επικοινωνεί με τρίτο → δηλώνει AI φύση και για λογαριασμό ποιου ενεργεί.
- Μόνο backend/M2M → συνήθως όχι Article 50(1).

40. Παράγεται ή αλλοιώνεται text, audio, image ή video;

- Ναι → έλεγχος Article 50(2), standard editing και Guidelines-based scope exclusions.
- In-scope → machine-readable marking και detectability από τον provider.

41. Το image/audio/video είναι deepfake;

- Ναι → visible/audible deployer disclosure από την πρώτη έκθεση.
- Εμφανώς artistic/satirical/fictional → η γνωστοποίηση παραμένει, αλλά μπορεί να είναι λιγότερο παρεμβατική.

42. Το AI κείμενο δημοσιεύεται για ενημέρωση του κοινού σε θέμα δημοσίου ενδιαφέροντος;

- Όχι → δεν ενεργοποιείται αυτό το συγκεκριμένο deployer label.
- Ναι, χωρίς ουσιαστικό human review → disclosure.
- Ναι, με ουσιαστικό review/editorial control και editorial responsibility → μπορεί να ισχύει η εξαίρεση.

43. Άλλαξε ο σκοπός, το brand ή η αρχιτεκτονική;

- Ναι → νέα αξιολόγηση ρόλων, Article 50 και πιθανής high-risk κατάταξης.

15. Οι 12 ερωτήσεις που πρέπει να μπορεί να απαντήσει σήμερα μια ομάδα προϊόντος

44. Ποιο ακριβώς AI system διαθέτουμε και ποιο model ενσωματώνουμε;

45. Για ποιο σύστημα είμαστε provider και για ποιο deployer;
46. Πού αλληλεπιδρά το AI απευθείας με φυσικά πρόσωπα;
47. Πώς και πότε βλέπει ή ακούει ο χρήστης το AI disclosure;
48. Ποια outputs είναι text, audio, image ή video και ποια είναι πραγματικά M2M;
49. Ποια marking τεχνική χρησιμοποιείται ανά modality;
50. Πώς αποδεικνύεται η detectability και ποιος την έχει δοκιμάσει;
51. Τι συμβαίνει στο mark μετά από export, compression ή upload σε πλατφόρμα;
52. Ποιος αποφασίζει αν ένα visual είναι deepfake;
53. Ποιος κάνει το τελικό substantive human review ενός δημόσιου άρθρου;
54. Ποιος φέρει editorial responsibility και πού φαίνεται αυτό δημόσια;
55. Τι αλλάζει όταν προστίθεται νέο μοντέλο, νέο κανάλι ή white-label πελάτης;

Αν η ομάδα δεν μπορεί να απαντήσει στις παραπάνω ερωτήσεις, η πρώτη ανάγκη δεν είναι άλλο ένα banner. Είναι καθαρή χαρτογράφηση του προϊόντος και της ευθύνης.

Η συμμόρφωση πρέπει να σχεδιαστεί μέσα στο προϊόν

Από τις 2 Αυγούστου 2026, η διαφάνεια της AI στην Ευρωπαϊκή Ένωση περνά από γενικές υποσχέσεις σε συγκεκριμένες product και publication υποχρεώσεις. Το chatbot πρέπει να συστήνεται ως AI όταν πρέπει. Το συνθετικό output πρέπει να είναι μηχανικά αναγνωρίσιμο όταν βρίσκεται εντός score. Το deepfake χρειάζεται ανθρώπινα αντιληπτή γνωστοποίηση. Το δημόσιο AI κείμενο χρειάζεται είτε disclosure είτε πραγματική ανθρώπινη επιμέλεια με editorial responsibility.

Για μια μικρή ομάδα, αυτό δεν απαιτεί απαραίτητα τεράστιο compliance department. Απαιτεί όμως τέσσερα πράγματα που συχνά λείπουν:

- σαφή ρόλο για κάθε σύστημα,
- disclosure ενσωματωμένο στο UX και όχι κρυμμένο σε νομικά κείμενα,
- τεχνικά δοκιμασμένο marking/detection,
- πραγματική ανθρώπινη λογοδοσία πριν από τη δημόσια δημοσίευση.

Η πιο ασφαλής αρχή είναι να αντιμετωπίζουμε τη διαφάνεια ως λειτουργία του προϊόντος, όχι ως σημείωση που προστίθεται την τελευταία ημέρα.

Επίσημες πηγές

- [Κανονισμός \(ΕΕ\) 2024/1689 — Artificial Intelligence Act](#)
- [Κανονισμός \(ΕΕ\) 2026/1744 — AI Omnibus](#)
- [European Commission FAQ: Transparency obligations under Article 50](#)
- [European Commission Guidelines για providers και deployers](#)
- [Code of Practice on Transparency of AI-generated Content](#)
- [EU icons για επισήμανση AI-generated content](#)
- [Ενημερωμένο ευρωπαϊκό χρονοδιάγραμμα εφαρμογής του AI Act](#)
- [Ελληνικό πλαίσιο εφαρμογής του AI Act — επίσημη ανακοίνωση](#)
- [N. 5321/2026 — ΦΕΚ Α' 114/20.07.2026](#)

Σημείωση ακρίβειας Το παρόν αποτυπώνει το δίκαιο και τις επίσημες πηγές που ήταν διαθέσιμες στις 2 Αυγούστου 2026. Οι *Guidelines* και ο *Code of Practice* αποτελούν επίσημη ερμηνευτική/πρακτική καθοδήγηση, αλλά δεν αντικαθιστούν το δεσμευτικό κείμενο του Κανονισμού. Για συγκεκριμένο προϊόν ή χρήση απαιτείται εξατομικευμένη νομική αξιολόγηση.