

ენერგო-პრო ჯგუფი

ხელოვნური ინტელექტის გამოყენების პოლიტიკა
2025



ხელოვნური ინტელექტის გამოყენების პოლიტიკა

ვერსია: 1.0 | ძალაში შედის: 2025 წლის 1 ოქტომბრიდან | მფლობელი: ხელოვნური ინტელექტის მართვის კომიტეტი (AIGC)

1. მიზანი და მოქმედების სფერო

1.1. ხელოვნური ინტელექტის გამოყენების ეს პოლიტიკა („პოლიტიკა“) ადგენს ხელოვნური ინტელექტის („ხელოვნური ინტელექტი“) საშუალებების უსაფრთხო, ეფექტური და პასუხისმგებლიანი გამოყენების სავალდებულო წესებს ENERGO – PRO a.s.-ში („კომპანია“) და მის შვილობილ კომპანიებში (ერთობლივად, „ჯგუფი“). ამ პოლიტიკის მიზნებისათვის, ჯგუფი მოიცავს DK Holding Investments, s.r.o.-ს, კომპანიის ერთადერთ აქციონერს, და მის ყველა პირდაპირ და არაპირდაპირ შვილობილ კომპანიას.

1.1. ჩვენი მიზანია ინოვაციებისა და პროდუქტიულობის მხარდაჭერა თანამედროვე ტექნოლოგიების დანერგვის გზით, ჩვენი მონაცემების, ინტელექტუალური საკუთრების, მომხმარებლებისა და რეპუტაციის დაცვის უმაღლესი დონის უზრუნველყოფის პარალელურად. ხელოვნური ინტელექტის დანერგვა უნდა მოხდეს არა მხოლოდ ინოვაციებით, არამედ ოპერაციული უსაფრთხოების, საზოგადოებრივი უსაფრთხოებისა და საზოგადოების მიერ ჩვენდამი გამოხატული ნდობისადმი მტკიცე ერთგულებით.

1.2. ჩვენ ვალდებულნი ვართ უზრუნველყოთ, რომ ხელოვნური ინტელექტის ნებისმიერი გამოყენება შეესაბამებოდეს ჩვენს ძირითად ეთიკურ პრინციპებს, როგორცაა მთლიანობის, პატივისცემის, გამჭვირვალობისა და ეთიკური ქცევის პრინციპები. უმნიშვნელოვანესია, რომ ადამიანის ზედამხედველობა/კონტროლი და პასუხისმგებლობა უმთავრესი დარჩეს ჩვენს ყველა ხელოვნურ ინტელექტზე დაფუძნებულ პროცესში.

1.3. აღნიშნული პოლიტიკა ვრცელდება ყველა თანამშრომელზე, გარე კონტრაქტორზე, კონსულტანტსა და მომწოდებელზე, რომლებსაც აქვთ წვდომა ჯგუფის საინფორმაციო სისტემებზე ან ამუშავებენ ჯგუფის მონაცემებს (ერთობლივად, „პერსონალი“). ის ვრცელდება ჯგუფში გამოყენებულ ყველა ხელოვნური ინტელექტის სისტემაზე, მიუხედავად იმისა, შემუშავებულია თუ არა ისინი შიდა სისტემის მიერ, შეძენილია მესამე მხარისგან თუ წვდომა მოხდა საჯარო სერვისის სახით.

2. განმარტებები

აღნიშნული პოლიტიკის მიზნებისათვის გამოიყენება შემდეგი განმარტებები:

2.1. **ხელოვნური ინტელექტის სისტემა:** ხელოვნური ინტელექტის შესახებ კანონის 3(1) მუხლით განსაზღვრული ხელოვნური ინტელექტის სისტემა (როგორც ეს ქვემოთ არის განსაზღვრული). ნებისმიერი მექანიზირებულ მეთოდზე დაფუძნებული სისტემა, რომელიც შექმნილია ავტონომიურობის სხვადასხვა დონით მუშაობისთვის და რომელმაც შეიძლება გამოავლინოს ადაპტირება მიღება/გამოყენების შემდეგ და რომელიც, აშკარა ან დაფარული მიზნებისთვის, შეყვანილი ინფორმაციიდან რასაც იღებს დაადგინოს, თუ როგორ დააგენერიროს მიღებული/გამოტანილი მონაცემები, როგორცაა პროგნოზები, შინაარსი, რეკომენდაციები ან გადაწყვეტილებები, რომლებსაც შეუძლიათ გავლენა მოახდინონ ფიზიკურ ან ვირტუალურ გარემოზე. ეს განსაზღვრება ვრცელდება მიუხედავად იმისა, არის თუ არა ხელოვნური ინტელექტის სისტემა დამოუკიდებელი აპლიკაცია თუ ინტეგრირებული ფუნქცია უფრო დიდი პროგრამული უზრუნველყოფის კომპლექტში.

საწარმოს ხელოვნური ინტელექტი: დამტკიცებული ხელოვნური ინტელექტის სისტემა, რომელიც რეგულირდება ხელშეკრულებით (მათ შორის, ქვემოთ განსაზღვრული მონაცემთა დამუშავების დანართით), რომელიც უზრუნველყოფს ჩვენი მონაცემების დაცვას. ასეთი შეთანხმებები აშკარად უნდა კრძალავდეს გამყიდველისთვის ჯგუფის მონაცემების გამოყენებას საკუთარი ან მესამე მხარის მოდელების ტრენინგისთვის.

მაგალითები: Microsoft Copilot M365-ისთვის, DeepL Pro.

საჯარო ხელოვნური ინტელექტი: დამტკიცებული ხელოვნური ინტელექტის სისტემა, რომელიც, როგორც წესი, ხელმისაწვდომია საჯარო ინტერნეტში და რომელიც არ გვთავაზობს მონაცემთა დაცვის სახელშეკრულებო გარანტიებს და შეიძლება გამოიყენოს გაცემული ინფორმაცია (მოთხოვნილი ინფორმაცია) და მიღებული ინფორმაცია (პასუხი) საკუთარი ან მესამე მხარის მოდელების გასაუმჯობესებლად ან გადასამზადებლად.

მაგალითები: ChatGPT-ის, Google Gemini-ის, Google Translate-ის სტანდარტული ვერსიები.

2.2. მონაცემთა კლასიფიკაცია: ყველა მონაცემი კლასიფიცირებულია ოთხ დონეზე

საჯარო (დონე 1): ნებისმიერი ინფორმაცია, რომელიც კანონიერად/კანონთა შესაბამისად არის საჯარო დომეინში. ის მოიცავს როგორც ჯგუფის მიერ ოფიციალურად გამოქვეყნებულ ინფორმაციას, ასევე გარე, საჯარო წყაროებიდან მიღებულ ინფორმაციას.

მაგალითები: გამოქვეყნებული პრესრელიზები, აკადემიური ნაშრომები, საინფორმაციო სტატიები, საჯარო სამთავრობო მონაცემთა ბაზები, საჯაროდ ხელმისაწვდომ ვებგვერდებზე განთავსებული ინფორმაცია.

შიდა (დონე 2): არასაჯარო ინფორმაცია, რომლის არაავტორიზებული გამჟღავნება წარმოადგენს დაბალ რისკს.

მაგალითები: სამუშაო გუნდის ზოგადი შეხვედრის ოქმები (რომლებიც არ შეიცავს პერსონალურ მონაცემებს ან კომერციულ საიდუმლოებებს), შიდა სახელმძღვანელოები, არასენსიტიური პროექტის პროექტები.

კონფიდენციალური (დონე 3): მგრძნობიარე/სენსიტიური სახის მონაცემები, რომელთა გამჟღავნებამ შეიძლება გამოიწვიოს საშუალო ან მძიმე ზიანი ჯგუფისთვის, მისი პარტნიორებისთვის ან მომხმარებლებისთვის.

მაგალითები: ფინანსური გეგმები, ბიზნეს სტრატეგიები, არააუდირებული ფინანსური შედეგები, ანონიმური კონტრაქტები.

მკაცრად კონფიდენციალური (დონე 4): ყველაზე მგრძნობიარე/სენსიტიური მონაცემები, დაცული კანონით ან კონტრაქტით, სადაც არაავტორიზებულმა გამჟღავნებამ შეიძლება გამოიწვიოს სერიოზული სამართლებრივი, ფინანსური ან რეპუტაციული ზიანი.

მაგალითები: პერსონალური მონაცემები, როგორც ეს განსაზღვრულია GDPR-ით (მოცემულია ქვემოთ) (მაგ., თანამშრომლის ან მომხმარებლის ჩანაწერები), სავაჭრო საიდუმლოებები, ფასის მგრძნობიარე/სენსიტიური ინფორმაცია, ინფორმაცია, რომელიც დაცულია გაუმჟღავნებლობის შეთანხმებით (NDA), ნოუ-ჰაუ, სისტემაში შესვლის მონაცემები და პაროლები, მესაკუთრის/მფლობელის კოდი, კრიტიკული ინფრასტრუქტურის ოპერაციული

მონაცემები (მაგ., SCADA სისტემები), ქსელის დეტალური დაუცველობის შეფასებები ან ელექტროენერგიის გენერაციის ერთეულების სქემები.

2.2. **მონაცემთა დამუშავების დანართი (DPA):** იურიდიულად სავალდებულო ხელშეკრულება, რომელიც დადებულია მონაცემთა მაკონტროლებრსა და მონაცემთა დამმუშავებელს შორის GDPR-ის 28-ე მუხლის შესაბამისად. იგი არეგულირებს პერსონალური მონაცემების დამუშავებას დამმუშავებლის მიერ მაკონტროლებრის სახელით. აღნიშნული არის სავალდებულო მოთხოვნა GDPR-ის თანახმად, როდესაც მაკონტროლებელი ქირაობს დამმუშავებელს პერსონალური მონაცემების სამართავად.

2.3. **Deepfake:** Any AI-generated or AI-manipulated video, audio or image content that resembles existing persons, objects, places, subjects or events and could falsely appear to a reasonable person as authentic or truthful.

დიდი სიყალბე: ნებისმიერი ხელოვნური ინტელექტის მიერ გენერირებული ან ხელოვნური ინტელექტის მიერ მანიპულირებული ვიდეო, აუდიო ან გამოსახულების კონტენტი, რომელიც წააგავს არსებულ პირებს, ობიექტებს, ადგილებს, სუბიექტებს ან მოვლენებს და შეიძლება ადამიანისთვის იყოს ცრუდ აღქმული, როგორც ავთენტური ან ნამდვილი/მართალი.

2.4. **ჰალუცინაცია:** ფენომენი, რომლის დროსაც ხელოვნური ინტელექტის მოდელი წარმოქმნის აბსურდულ, ფაქტობრივ მცდარ ან მთლიანად შეთხზულ შედეგებს, მიუხედავად იმისა, რომ ისინი წარმოდგენილია თავდაჯერებულად და თავისუფლად.

2.5. **მინიშნება:** ნებისმიერი მიწოდებული ინფორმაცია, როგორიცაა კითხვა, ინსტრუქცია ან მონაცემთა ერთობლიობა, რომელსაც მომხმარებელი აწვდის ხელოვნური ინტელექტის სისტემას ინფორმაციის მისაღებად.

3. როლები, პასუხისმგებლობები და მმართველობა

ხელოვნური ინტელექტის მმართველობის კომიტეტი (AIGC): ერთობლივი გუნდი, რომელიც შედგება უმაღლესი ხელმძღვანელობის, იურიდიული, ადამიანური რესურსების და სხვა ძირითადი დეპარტამენტების წარმომადგენლებისგან. ის ამტკიცებს ხელოვნური ინტელექტის სტრატეგიას, ამ პოლიტიკას და პასუხისმგებელია რეესტრის წარმოებაზე.

ხელოვნური ინტელექტის ნავიგატორი: ხელოვნურ ინტელექტთან დაკავშირებულ საკითხებზე ხელმძღვანელობისა და მხარდაჭერის ძირითადი საკონტაქტო პირი. ეს როლი უზრუნველყოფს პროაქტიულ რჩევებს, უჭერს მხარს AIGC-ს და ეხმარება პერსონალს ხელოვნურ ინტელექტთან დაკავშირებულ შეკითხვებში.

ხელოვნური ინტელექტის სისტემის მფლობელი: თუ ხელოვნური ინტელექტის მართვის კომიტეტი (AIGC) სხვაგვარად არ გადაწყვეტს, AINavigator იმოქმედებს, როგორც თითოეული დამტკიცებული AIS სისტემის მფლობელი თითოეული დამტკიცებული AIS სისტემისთვის. AIS სისტემის მფლობელი პასუხისმგებელია:

1. უზრუნველყოს რომ ხელოვნური ინტელექტის სისტემა გამოყენებული იქნას მხოლოდ მისი დამტკიცებული მიზნისთვის და მისი ავტორიზებული პარამეტრების ფარგლებში;

2. გამოიყენება როგორც აუდიტებისა და საქმიანობის შესრულების შედეგების მიმოხილვებისთვის ძირითადი ბიზნეს კონტაქტი;
3. მომხმარებლის წვდომის მართვა IT-თან კოორდინაციით; და
4. ნებისმიერი მნიშვნელოვანი შესრულების საკითხის ან ინციდენტის შესახებ AI Navigator-ისთვის შეტყობინება AIGC-ში გადასაცემად.

პერსონალი: პერსონალის თითოეული წევრი პასუხისმგებელია ამ პოლიტიკის გაგებასა და შეჯერებაზე და შესაბამისობაში მოყვანაზე.

4. ხელოვნური ინტელექტის გამოყენების ძირითადი პრინციპები

ჯგუფში ხელოვნური ინტელექტის ყველა სახის გამოყენება უნდა შეესაბამებოდეს შემდეგ ძირითად პრინციპებს:

- 4.1. **ადამიანის ჩართულობა პროცესში (სავალდებულო ადამიანური ზედამხედველობა):** ხელოვნური ინტელექტი დამხმარე ინსტრუმენტია და არა ავტონომიური გადაწყვეტილების მიმღები. ხელოვნური ინტელექტის სისტემას შეუძლია დაეხმაროს და მხარი დაუჭიროს ადამიანურ შესაძლებლობებს, მაგრამ ის არ ცვლის ადამიანის მსჯელობის უნარს და ანგარიშვალდებულებას. ხელოვნური ინტელექტის მიერ გენერირებული ყველა შედეგი, განსაკუთრებით ის, რასაც პოტენციური სამართლებრივი ან ფინანსური შედეგები მოჰყვება, გამოყენებამდე მნიშვნელოვნად უნდა იქნას განხილული, დამოწმებული და დამტკიცებული კვალიფიციური და პასუხისმგებელი პირის მიერ. შემოწმების დონე უნდა იყოს გადაწყვეტილების რისკის პროპორციული (მაგ., ანგარიშში ერთი ფაქტის დადასტურება ნაკლებ შემოწმებას/დაკვირვებას მოითხოვს, ვიდრე ხელოვნური ინტელექტის მიერ გენერირებული ქსელის ტექნიკური მომსახურების გრაფიკის დამტკიცება).
- 4.2. **ანგარიშვალდებულება:** ხელოვნური ინტელექტის მიერ გენერირებულ შედეგებზე დაფუძნებული ნებისმიერი გადაწყვეტილების ან ქმედების საბოლოო პასუხისმგებლობა ეკისრება მის მომხმარებელს და არა თავად ხელოვნური ინტელექტის სისტემას. პერსონალმა უნდა გაითვალისწინოს, რომ ხელოვნური ინტელექტის სისტემამ შეიძლება შექმნას „ჰალუცინაციები“ და მათ ყოველთვის უნდა გადაამოწმონ საბოლოო ინფორმაცია დამოუკიდებელი და სანდო წყაროდან.
- 4.3. **გამჭვირვალობა:** იმ შემთხვევაში, თუ ხელოვნური ინტელექტის სისტემა შექმნილია ფიზიკურ პირებთან პირდაპირი ურთიერთქმედებისთვის (მაგალითად, ჩატბოტის მეშვეობით), ამ პირებს უნდა ეცნობოთ, რომ ისინი კონტაქტობენ ხელოვნური ინტელექტის სისტემასთან. პერსონალმა, რომელიც ქმნის ან მანიპულირებს გამოსახულების, აუდიოს ან ვიდეო შინაარსს, რომელიც ავტენტურად არის მიჩნეული (მარკეტინგული მასალებისთვის ხელოვნური ინტელექტის მიერ გენერირებული ილუსტრაციებიდან დაწყებული Deepfakes-ით დამთავრებული), ნათლად და თვალსაჩინოდ უნდა გაამჟღავნოს, რომ შინაარსი ხელოვნურად არის გენერირებული ან მანიპულირებული, რაც უფრო დეტალურად არის აღწერილი 6.3(iii) ნაწილში.
- 4.4. **არა-დისკრიმინაცია და სამართლიანობა:** ხელოვნური ინტელექტის სისტემები უნდა იქნას გამოყენებული ისე, რომ თავიდან იქნას აცილებული უსამართლო მიკერძოება ან დისკრიმინაცია ინდივიდების ან ჯგუფების მიმართ დაცული მახასიათებლების (მაგ., ასაკი, სქესი, ეთნიკური კუთვნილება) საფუძველზე.

- 4.5. **უსაფრთხოება და კონფიდენციალურობა გეგმის შესაბამისად:** მონაცემთა დაცვა და უსაფრთხოება ფუნდამენტურია. ყველა ხელოვნური ინტელექტის პროექტი, შესყიდვიდან დაწყებული, განლაგებითა და საბოლოო დემონტაჟით დამთავრებული, თავიდანვე უნდა ითვალისწინებდეს კონფიდენციალურობისა და უსაფრთხოების პრინციპებს.

5. მარეგულირებელი კანონმდებლობა და პოლიტიკები

5.1. სამართლებრივი და მარეგულირებელი ჩარჩო

- 5.2. ჯგუფის მიერ ხელოვნური ინტელექტის სისტემების შემუშავება, შესყიდვა და გამოყენება რეგულირდება მოქმედი კანონმდებლობისა და საერთაშორისო სტანდარტების ჩარჩოთი. ყველა პერსონალი ვალდებულია დაიცვას შემდეგი დებულებების უახლესი ვერსიები, შეზღუდვის გარეშე:

ევროპის პარლამენტისა და საბჭოს 2024 წლის 13 ივნისის **რეგულაცია (EU) 2024/1689**, რომელიც ადგენს ხელოვნური ინტელექტის ჰარმონიზებულ წესებს (**„ხელოვნური ინტელექტის შესახებ კანონი“**).

ევროპარლამენტისა და საბჭოს 2023 წლის 13 დეკემბრის **რეგულაცია (EU) 2023/2854** მონაცემებზე სამართლიანი წვდომისა და მათი გამოყენების ჰარმონიზებული წესების შესახებ (**„მონაცემთა შესახებ აქტი“**).

ევროპარლამენტისა და საბჭოს 2022 წლის 14 დეკემბრის **დირექტივა (EU) 2022/2555** კავშირის მასშტაბით კიბერუსაფრთხოების მაღალი საერთო დონის უზრუნველყოფის ზომების შესახებ (**„NIS 2 დირექტივა“**).

ევროპის პარლამენტისა და საბჭოს 2019 წლის 17 აპრილის **დირექტივა (EU) 2019/790** საავტორო უფლებებისა და მომიჯნავე უფლებების შესახებ ციფრულ ერთიან ბაზარზე (**„საავტორო უფლებების დირექტივა“**).

ევროპარლამენტისა და საბჭოს 2016 წლის 27 აპრილის **რეგულაცია (EU) 2016/679** პერსონალური მონაცემების დამუშავებასთან დაკავშირებით ფიზიკური პირების დაცვისა და ასეთი მონაცემების თავისუფალი გადაადგილების შესახებ (**„GDPR“**).

ISO/IEC 42001:2023 – ინფორმაციული ტექნოლოგიები — ხელოვნური ინტელექტი — მართვის სისტემა.

ყველა მოქმედი ეროვნული კანონმდებლობა იურისდიქციებში სადაც ჯგუფი ოპერირებს მათ შორის, მაგრამ არა მხოლოდ, შრომის სამართლის, ინტელექტუალური საკუთრებისა და სამოქალაქო პასუხისმგებლობის მარეგულირებელი აქტები.

5.3. სხვა შიდა პოლიტიკებთან ურთიერთობა

- 5.4. აღნიშნული პოლიტიკა ავსებს და უნდა იქნას განმარტებული ჯგუფის სხვა სავალდებულო პოლიტიკებთან ერთად. პერსონალი ასევე ვალდებულია უზრუნველყოს, რომ ხელოვნური ინტელექტის სისტემების გამოყენება შეესაბამებოდეს ყველა სხვა შესაბამის შიდა პოლიტიკას, მათ შორის, მაგრამ არა მხოლოდ:

- ქცევის კოდექსი
- ადამიანის უფლებათა პოლიტიკა

- ადამიანური რესურსების პოლიტიკა
- უსაფრთხოების პოლიტიკა
- მონაცემთა დაცვის პოლიტიკა (შიდა და გარე)
- კონტრაქტორისა და ქვეკონტრაქტორის მართვის გეგმა

5.5. აღნიშნულ პოლიტიკასა და სხვა შიდა პოლიტიკას შორის ნებისმიერი კონფლიქტის შემთხვევაში, გამოიყენება უფრო შემზღუდავი ნორმა. საკითხი AIGC-მ დაუყოვნებლივ უნდა გადასცეს AI Navigator-ს შემდგომი განმარტებისთვის.

6. მისაღები და აკრძალული გამოყენება

6.1. ოქროს წესი: მონაცემების შესაბამის ხელოვნური ინტელექტის სისტემასთან შესაბამისობაში მოყვანა.

ყველაზე მნიშვნელოვანი პრინციპი არის იმის უზრუნველყოფა, რომ მონაცემები დამუშავდეს მხოლოდ შესაბამისი დონის უსაფრთხოების მქონე ხელოვნური ინტელექტის სისტემის მიერ. ამ პრინციპის შეუსრულებლობა წარმოადგენს ამ პოლიტიკის სერიოზულ დარღვევას.

კლასიფიცირების მონაცემები	საჯარო ხელოვნური ინტელექტი	საწარმოს ხელოვნური ინტელექტი
დონე 1 : საჯარო	დაშვებული	დაშვებული
დონე 2: შიდა	დაშვებული	დაშვებული
დონე 3 - კონფიდენციალური	დაშვებული შეზღუდვებით	დაშვებული
დონე 4 : მკაცრად კონფიდენციალური	აკრძალული	მხოლოდ AIGC დასტურით, სავალდებულო DPIA და IT Security Risk Assessment ის საგანი

6.2. დამტკიცებული ხელოვნური ინტელექტის სისტემები

ზოგადი წესი

პერსონალს უფლება აქვს გამოიყენოს მხოლოდ ის ხელოვნური ინტელექტის სისტემები, რომლებიც ჩამოთვლილია ოფიციალურ „დამტკიცებული ხელოვნური ინტელექტის სისტემების რეესტრში“ (შემდგომში „რეესტრი“). ნებისმიერი ხელოვნური ინტელექტის სისტემის გამოყენება, რომელიც არ არის ჩამოთვლილი რეესტრში, მკაცრად აკრძალულია.

ამ პოლიტიკის ძალაში შესვლის თარიღისთვის დამტკიცებული ხელოვნური ინტელექტის სისტემების სია თან ერთვის დანართ A-ს სახით. რეესტრის ავტორიტეტულ, მოქმედ ვერსიას აწარმოებს და რეგულარულად აახლებს AIGC და ხელმისაწვდომია ჯგუფის M365 SharePoint-ზე.

მინიმალური რისკის მქონე ხელოვნური ინტელექტის სისტემების გამოყენება

პირი, რომელიც გამოიყენებს ამ გამონაკლისს, ვალდებულია დაუყოვნებლივ აცნობოს AI Navigator-ს ასეთი AI სისტემის გამოყენების შესახებ ელექტრონული ფოსტით. ნებისმიერი ასეთი სისტემის გამოყენება ყოველთვის უნდა შეესაბამებოდეს აღნიშნულ პოლიტიკას და ჯგუფის ყველა სხვა პოლიტიკას სრულად, მათ შორის, მაგრამ არა მხოლოდ, მონაცემთა დაცვის (GDPR), ინფორმაციული უსაფრთხოებისა და IT უსაფრთხოების შესახებ პოლიტიკას.

საჯარო ხელოვნური ინტელექტის სისტემების გამოყენება მე-3 დონის (კონფიდენციალური) მონაცემებისთვის

საჯარო ხელოვნური ინტელექტის სისტემების გამოყენებით მე-3 დონის (კონფიდენციალური) მონაცემების დამუშავება მკაცრად აკრძალულია, გარდა იმ შემთხვევებისა, როდესაც ყველა შემდეგი პირობა შესაბამისობაშია მოყვანილი:

(i) გამოყენებული საჯარო ხელოვნური ინტელექტის სისტემა უნდა იყოს დამტკიცებული ამ მიზნით რეესტრში და გამოყენებული იქ მითითებული პირობების შესაბამისად. ამ პოლიტიკის ბოლო გადახედვის მდგომარეობით, საჭირო მეთოდებია:

- **Adobe Acrobat AI Assistant-ისთვის:** დამტკიცებულია დამატებითი კონფიდენციალურობის პარამეტრების მოთხოვნის გარეშე გამოსაყენებლად. მისი კონფიდენციალურობის მონახაზის მიხედვით სტრუქტურა უზრუნველყოფს მომხმარებლის მონაცემების იზოლირებას და მოდელის ტრენინგისთვის გამოუყენებლობას.
- **ChatGPT-ისთვის:** აუცილებლად უნდა გამოიყენოთ „დროებითი სასაუბროს“ ფუნქცია.
- **Google Gemini-ისთვის:** „Gemini Apps Activity“ პარამეტრი უნდა იყოს გამორთული.
- **NotebookLM-ისთვის:** დამტკიცებულია დამატებითი კონფიდენციალურობის პარამეტრების მოთხოვნის გარეშე გამოსაყენებლად. მისი კონფიდენციალურობის დიზაინის მიხედვით სტრუქტურა უზრუნველყოფს მომხმარებლის მონაცემების იზოლირებას და მოდელის ტრენინგისთვის გამოუყენებლობას.

საჯარო ხელოვნური ინტელექტის სისტემის გამოყენების პირობები გარანტიას იძლევა, რომ მომხმარებლის შეყვანილი (მოთხოვნები) და სისტემისგან მიღებული (პასუხები) მონაცემები არ უნდა იქნას გამოყენებული მიმწოდებლის ან მესამე მხარის ხელოვნური ინტელექტის მოდელის ტრენინგის ან გაუმჯობესების მიზნით.

სესიიდან მიღებული ყველა შემავალი და გამომავალი მონაცემი არ შეინახება მომხმარებლის მუდმივი საუბრების ისტორიაში და პროვაიდერის მიერ ავტომატურად უნდა წაიშალოს. უსაფრთხოებისა და ოპერაციული მიზნებისთვის მომხმარებლებმა უნდა დაადასტურონ პროვაიდერის მიერ შენახვის პერიოდი, რომელიც ამჟამად შეადგენს:

- ChatGPT-ის „დროებითი სასაუბროდ“ - 30 დღემდე.
- 72 საათამდე (3 დღე) Google Gemini-სთვის „აქტივობის“ გამორთვით.

საჯარო ხელოვნური ინტელექტის სისტემაში მე-3 დონის (კონფიდენციალური) მონაცემების წარდგენამდე, მომხმარებლის პასუხისმგებლობაა სწორი ფუნქციის არჩევა და მონაცემთა შენახვის შესაბამისი პერიოდის გაგება.

6.3. აკრძალული აქტივობები

ხელოვნური ინტელექტის გამოყენება შემდეგი მიზნებისთვის მკაცრად აკრძალულია:

- (i) ნებისმიერი საქმიანობა, რომელიც ეწინააღმდეგება მოქმედ კანონმდებლობას, განსაკუთრებით ხელოვნური ინტელექტის შესახებ კანონის მე-5 მუხლში განსაზღვრულ აკრძალულ პრაქტიკას, ან არღვევს ჯგუფის ქცევის კოდექსს ან სხვა მოქმედ შიდა პოლიტიკას.
- (ii) ნებისმიერი მე-4 დონის მონაცემების შეყვანა ნებისმიერ საჯარო ხელოვნური ინტელექტის სისტემაში ან ნებისმიერი მე-3 დონის მონაცემების შეყვანა ნებისმიერ საჯარო ხელოვნური ინტელექტის სისტემაში 6.2 პუნქტში მითითებული შეზღუდვების დაცვის გარეშე.

(iii) აკრძალულია Deepfakes-ის და სხვა ხელოვნური ინტელექტის მიერ გენერირებული სინთეზური მედიის შექმნა ან გავრცელება მავნე, თაღლითური ან შეცდომაში შეყვანის მიზნით — ან ნებისმიერი გზით, რომელიც მიზნად ისახავს, სავარაუდოდ ნებისმიერი პირის შეცდომაში შეყვანას და დარწმუნებას რომ შინაარსი ავთენტურია.

ხელოვნური ინტელექტის მიერ გენერირებული ან მანიპულირებული სურათების, აუდიოს, ვიდეოს ან მსგავსი გამოშვებული მონაცემების შექმნა და გამოყენება სხვაგვარად დაშვებულია ლეგიტიმური ილუსტრაციული, საგანმანათლებლო, დიზაინის, სასწავლო ან სხვა საქმიანი მიზნებისთვის, იმ პირობით, რომ:

ა. შინაარსი ნათლად არის გამჟღავნებული სამიზნე აუდიტორიისთვის, როგორც ხელოვნური ინტელექტის მიერ გენერირებული/მანიპულირებული (მაგალითად, აქტივზე განთავსებული იარლიყით, წარწერით, watermark, მეტამონაცემების შეტყობინებით, სლაიდის footnote ან სხვა თანმხლები ინფორმაციით, რომელიც ჩანს გამოყენების მომენტში); და

ბ. შინაარსი არ არის წარმოდგენილი, როგორც რეალური სამყაროს პირების, საგნების, ადგილების, სუბიექტების ან მოვლენების ფაქტობრივი ჩანაწერი.

- (ii) რეალისტური ადამიანის მსგავსი ავტარების შექმნა და გამოყენება თანამშრომლებთან, მომხმარებლებთან ან საზოგადოებასთან მოტყუებითი ურთიერთობის ან მათი გაყალბების მიზნით.

(v) პირებთან დაკავშირებით საბოლოო, ავტომატიზირებული გადაწყვეტილებების მიღება (მაგ., დაქირავების, შესრულების მიმოხილვის ან სამსახურიდან გათავისუფლების დროს) მნიშვნელოვანი ადამიანური ზედამხედველობისა და მიმოხილვის გარეშე.

(vii) დამტკიცებული ხელოვნური ინტელექტის სისტემის უსაფრთხოების ზომების ან გამოყენების შეზღუდვების გვერდის ავლის ნებისმიერი მცდელობა.

6.4. წახალისებული აქტივობები

მივესალმებით და წავახალისებთ დამტკიცებული ხელოვნური ინტელექტის სისტემების გამოყენებას. ქვემოთ მოცემული სია საილუსტრაციოა და არა ამომწურავი, და წარმოადგენს სასარგებლო გამოყენების მაგალითებს:

- (i) ელექტრონული ფოსტის, ანგარიშების ან პრეზენტაციების შედგენაში დახმარება.
- (ii) დოკუმენტების თარგმნა (მონაცემთა კლასიფიკაციის წესების პირობით).
- (iii) დოკუმენტების შეჯამება და მონაცემების ანალიზი ახალი თვისებების განსასაზღვრად.
- (iv) იდეების გენერირება და კონცეპტუალური პროექტების შექმნა.
- (v) განმეორებადი დავალებების ავტომატიზაცია (მაგ., ცხრილის ფორმულების გენერირება).
- (vi) საჯარო მონაცემებზე დაყრდნობით ენერგეტიკული ბაზრის ტენდენციების შესახებ ტექნიკური ანგარიშების საწყისის შედგენა, ყველა ფაქტებითა და ციფრებით რომელიც შესაბამისად იქნება დამოწმებული დარგის ექსპერტის მიერ.

ნებისმიერი პოტენციური გამოყენების შემთხვევა რომელიც არ არის განსაზღვრული, პერსონალმა უნდა გამოიჩინოს სათანადო სიფრთხილე და უზრუნველყოს, რომ შემოთავაზებული გამოყენება სრულად შეესაბამებოდეს ამ პოლიტიკის ძირითად პრინციპებს და ყველა სხვა მოქმედ სამართლებრივ და შიდა რეგულაციას. ყოველთვის უნდა იქნას გამოყენებული საღი პროფესიული გადაწყვეტილება. თუ არსებობს რაიმე გაურკვევლობა შემოთავაზებული გამოყენების დაშვებასთან ან შესაბამისობასთან დაკავშირებით, გაგრძელებამდე სავალდებულოა AI Navigator-თან კონსულტაცია.

7. ხელოვნური ინტელექტის სისტემების შესყიდვა და დამტკიცება

7.1. ყველა ახალი ხელოვნური ინტელექტის სისტემამ გამოყენებამდე უნდა გაიაროს შემდეგი დამტკიცების პროცესი:

- (i) **მოთხოვნის წარდგენა:** პირმა უნდა წარადგინოს მოთხოვნა ოფიციალური ფორმის გამოყენებით (იხ. დანართი B).
- (ii) **საწყისი შეფასება:** ხელოვნური ინტელექტის ნავიგატორი განიხილავს მოთხოვნას ძირითადი კრიტერიუმების მიხედვით.
- (iii) **რისკის შეფასება:** AIGC-მ უნდა ჩაატაროს დეტალური რისკის შეფასება (იხ. მე-8 ნაწილი).
- (iv) **გადაწყვეტილება:** AIGC დაამტკიცებს, უარყოფს ან დაამოწმებს ხელოვნური ინტელექტის სისტემას განსაზღვრული შეზღუდვებით.
- (v) **რეგისტრაცია:** დამტკიცებული ხელოვნური ინტელექტის სისტემა უნდა აღირიცხოს რეესტრში.

8. რისკის შეფასება და სისტემის კლასიფიკაცია

8.1. თითოეული ხელოვნური ინტელექტის სისტემა კლასიფიცირდება ხელოვნური ინტელექტის აქტის შესახებ კანონის ოთხდონიანი რისკის მოდელის მიხედვით:

მიუღებელი რისკი: ხელოვნური ინტელექტის სისტემები, რომლებიც აკრძალულია კანონით და ამ პოლიტიკით (მაგ., სოციალური ქსელების შეფასება).

მაღალი რისკი: ხელოვნური ინტელექტის სისტემა, რომელიც აკმაყოფილებს ხელოვნური ინტელექტის შესახებ კანონის მე-6 მუხლში დადგენილ კრიტერიუმებს, მათ შორის, მაგრამ არა მხოლოდ, ხელოვნური ინტელექტის სისტემები, რომლებმაც შეიძლება მნიშვნელოვნად იმოქმედონ ინდივიდების ჯანმრთელობაზე, უსაფრთხოებაზე ან ფუნდამენტურ უფლებებზე (მაგ., კრიტიკული ინფრასტრუქტურის მართვაში გამოყენებული ხელოვნური ინტელექტის სისტემა ან სამუშაო განაცხადების ფილტრაციისთვის გამოყენებული ხელოვნური ინტელექტის სისტემა). ასეთი სისტემები ექვემდებარება უფრო მკაცრ წესებს, მათ შორის სავალდებულო რეგისტრაციას, მკაცრ ტესტირებას და უწყვეტ მონიტორინგს.

შეზღუდული რისკი: ეს კატეგორია მოიცავს ხელოვნური ინტელექტის სისტემებს, რომლებიც მანიპულირების ან მოტყუების კონკრეტულ რისკებს წარმოადგენენ და შესაბამისად, ექვემდებარებიან გამჭვირვალობის სპეციფიკურ ვალდებულებებს. ეს ვალდებულებები განსხვავდება ხელოვნური ინტელექტის სისტემის ტიპისა და მისი გამოყენების მიხედვით:

- **პირდაპირი ურთიერთქმედების სისტემები (მაგ., ჩატბოტები):** როდესაც იყენებთ სისტემებს, რომლებიც შექმნილია ინდივიდებთან პირდაპირი ურთიერთქმედებისთვის, როგორცაა ჩატბოტები შიდა ადამიანური რესურსების მხარდაჭერისთვის ან გარე მომხმარებელთა მომსახურებისთვის, მომხმარებლებს მკაფიოდ უნდა ეცნობოს, რომ მათ აქვთ კომუნიკაცია ხელოვნურ ინტელექტთან.
- **სინთეტიკური მედია (სურათები, აუდიო, ვიდეო):** ხელოვნური ინტელექტის სისტემებს, რომლებიც გამოიყენება რეალისტური სურათების, აუდიოს ან ვიდეო კონტენტის (მაგ., Deepfakes) გენერირების ან მანიპულირებისთვის, უნდა ჰქონდეთ მკაფიოდ აღნიშნული, როგორც „ხელოვნური ინტელექტით გენერირებული“ ან „ხელოვნური ინტელექტით მანიპულირებული“, თუ კონტენტი აშკარად არის ხელოვნური ან შემოქმედებითი და არ არის განკუთვნილი რეალობის წარმოსადგენად.
- **ზოგადი დანიშნულების ხელოვნური ინტელექტი (GPAI) ტექსტის გენერირებისთვის:** ისეთი ინსტრუმენტების ძირითადი მოდელები, როგორცაა ChatGPT, Gemini და Copilot, განიხილება, როგორც ზოგადი დანიშნულების ხელოვნური ინტელექტი (GPAI) და ექვემდებარება მათი პროვაიდერების გამჭვირვალობის მოთხოვნებს. ჩვენი, როგორც მომხმარებლების (განმთავსებლების) მიზნებისთვის, როდესაც ეს ინსტრუმენტები გამოიყენება დამხმარე ფუნქციებისთვის, როგორცაა ტექსტის შეჯამება, ელექტრონული ფოსტის შედგენა ან იდეების გენერირება, გენერირებულ ტექსტს არ სჭირდება იქნეს მონიშნული როგორც ხელოვნური ინტელექტის მიერ შექმნილი, იმ მოსაზრებით, რომ გამოყენებამდე ის უნდა განიხილოს და დააკორექტიროს ადამიანმა. ამ შემთხვევებში, ხელოვნური ინტელექტი ფუნქციონირებს როგორც ასისტენტი/დამხმარე და შინაარსზე საბოლოო პასუხისმგებლობა მთლიანად პერსონალს ეკისრება.

მინიმალური რისკი: აღნიშნული კატეგორია ვრცელდება ხელოვნური ინტელექტის სისტემებზე, რომლებიც უმნიშვნელო რისკს წარმოადგენენ ინდივიდების უფლებების,

თავისუფლებების ან უსაფრთხოებისთვის. სისტემა კლასიფიცირდება, როგორც მინიმალური რისკის მქონე მხოლოდ იმ შემთხვევაში, თუ ის აკმაყოფილებს ყველა შემდეგ პირობას:

- ის არ იღებს გადაწყვეტილებებს ან რაიმე განსაზღვრებას, რომლებიც გავლენას ახდენს ფიზიკური პირების სამართლებრივ უფლებებზე;
- ის არ ამუშავებს სენსიტიური ტიპის პერსონალურ მონაცემებს; და
- მისი პოტენციური გაუმართაობა ან არასწორი მიღებული მონაცემები/ინფორმაცია მატერიალურ უსაფრთხოებაზე გავლენას არ მოახდენს.

ხელოვნური ინტელექტის სისტემა შეიძლება გადაკვალიფიცირდეს უფრო მაღალი რისკის კატეგორიაში (შეზღუდული ან მაღალი რისკი), თუ მისი დანიშნულება შეცვლილია პერსონალური, ბიომეტრიული ან უსაფრთხოებისთვის კრიტიკული მონაცემების დამუშავების ჩათვლით. მიუხედავად იმისა, რომ ხელოვნური ინტელექტის შესახებ კანონით კონკრეტულად არ რეგულირდება, რაიმე მინიმალური რისკის მქონე ხელოვნური ინტელექტის სისტემის გამოყენება უნდა შეესაბამებოდეს ამ პოლიტიკაში დადგენილ ზოგად პრინციპებს.

მინიმალური რისკის მქონე ხელოვნური ინტელექტის სისტემების მაგალითები: ხელოვნური ინტელექტით მართული სპამის ფილტრები, ამინდის პროგნოზის მოდელები, ინვენტარიზაციის მართვის სისტემები, რომლებიც იყენებენ მარტივ ხელოვნურ ინტელექტს მოთხოვნის პროგნოზირებისთვის, გრამატიკისა და მართლწერის შემოწმების დახმარე საშუალებები, რომლებიც ჩადებულია ტექსტის დამუშავების ან ელექტრონული ფოსტის კლიენტებში, ავტომატური თარგმანის ინსტრუმენტები (მაგ. ვებსაიტის ავტომატური თარგმანი), რომლებიც გამოიყენება სწრაფად გაგებისთვის (საჯაროდ ხელმისაწვდომი ტექსტის თარგმნა უმნიშვნელო რისკს მხოლოდ მაშინ წარმოადგენს, როდესაც საქმე არ ეხება მგრძნობიარე /სენსიტიურ მონაცემებს).

9. მონაცემთა დაცვა, ინტელექტუალური საკუთრება და უსაფრთხოება

კონფიდენციალურობა გეგმის შესაბამისად: მონაცემთა დაცვის პრინციპი პროექტის/გეგმის შესაბამისად და შეუსრულებლობა უნდა იყოს ინტეგრირებული ყველა ხელოვნური ინტელექტის პროექტში თავიდანვე.

მონაცემთა დაცვის გავლენის შეფასება (იხ. დანართი C): მონაცემთა დაცვაზე ზემოქმედების შეფასება (იხ. დანართი C) სავალდებულოა ნებისმიერი ხელოვნური ინტელექტის სისტემისთვის, რომელიც, სავარაუდოდ, მაღალ რისკს შეუქმნის ფიზიკური პირების უფლებებსა და თავისუფლებებს.

ხელოვნური ინტელექტის მიერ გენერირებული კონტენტის ინტელექტუალური საკუთრება: პერსონალმა უნდა იცოდეს, რომ ხელოვნური ინტელექტის მიერ გენერირებული კონტენტის სამართლებრივი სტატუსი რთულია და ცვალებადია. თუ დამტკიცებული ხელოვნური ინტელექტის სისტემის პირობებში სხვაგვარად არ არის მითითებული, ხელოვნური ინტელექტის მიერ გენერირებული შედეგები შეიძლება არ ექვემდებარებოდეს საავტორო უფლებების დაცვას ან შეიძლება ექვემდებარებოდეს მესამე მხარის უფლებებს. დასაქმების პროცესში ხელოვნური ინტელექტის გამოყენებით შექმნილი ყველა კონტენტი ჯგუფის სამუშაო პროდუქტად ითვლება, მაგრამ პერსონალმა არ უნდა ჩათვალოს, რომ ჯგუფს აქვს ექსკლუზიური საავტორო უფლებები. ხელოვნური ინტელექტის მიერ გენერირებული

კონტენტის საზოგადოებისთვის განკუთვნილ მასალებში ან კომერციულ პროდუქტებში გამოყენება მოითხოვს წინასწარ კონსულტაციას ხელოვნური ინტელექტის ნავიგატორთან

კონფიდენციალური ინფორმაციისა და ინტელექტუალური საკუთრების დაცვა: პერსონალმა განსაკუთრებული სიფრთხილე უნდა გამოიჩინოს ჯგუფის მნიშვნელოვანი/პირადი/კონფიდენციალური ინფორმაციის დასაცავად. ეს მოიცავს, მაგრამ არ შემოიფარგლება მხოლოდ სავაჭრო საიდუმლოებებით, ფასის მიმართ მგრძნობიარე ინფორმაციით, ნოუ-ჰაუთი, ინფორმაციის საიდუმლოების შესახებ შეთანხმებებით დაცული ინფორმაციით, საავტორო უფლებებითა და სამრეწველო საკუთრების უფლებებით. მე-6 პუნქტის თანახმად, ასეთი ინფორმაცია არასდროს არ უნდა შეიტანოთ საჯარო ხელოვნური ინტელექტის სისტემაში.

უსაფრთხოება: ყველა ხელოვნური ინტელექტის სისტემა უნდა შეესაბამებოდეს ჩვენს მიერ დადგენილ კიბერუსაფრთხოების სტანდარტებს.

10. მონიტორინგი, ანგარიშგება და მაჩვენებელთა /საზომი სისტემა/მეტრიკა

ჩვენ გავაკონტროლებთ ჩვენი ხელოვნური ინტელექტის სისტემების მუშაობას, სიზუსტეს და უსაფრთხოებას.

მაღალი რისკის მქონე ხელოვნური ინტელექტის სისტემებისთვის სავალდებულოა AIGC-სთვის რეგულარული ანგარიშები შესრულების შესახებ.

ჩვენ გავზომავთ არა მხოლოდ ტექნიკურ პარამეტრებს, არამედ გენერირებულ ბიზნეს ღირებულებასაც (მაგ., დაზოგილი დრო, პროცესების გაუმჯობესება).

11. ინციდენტებზე რეაგირება და ცვლილებების მართვა

ინციდენტები: ხელოვნური ინტელექტის სისტემასთან დაკავშირებული ნებისმიერი ინციდენტი (მაგ., მონაცემთა გაჟონვა, მნიშვნელოვანი ზემოქმედების მქონე არასწორი მიღებული ინფორმაცია, დისკრიმინაციის სავარაუდო შემთხვევა) დაუყოვნებლივ უნდა ეცნობოს ხელოვნური ინტელექტის ნავიგატორს.

ცვლილებები: ხელოვნური ინტელექტის სისტემაში ნებისმიერი მნიშვნელოვანი ცვლილება (მაგ., ახალი ვერსია, მისი დანიშნულებისამებრ გამოყენების ცვლილება) მოითხოვს AIGC-ის მიერ რისკის ახალ შეფასებას. მნიშვნელოვანი ცვლილება მოიცავს, მაგრამ არ შემოიფარგლება მხოლოდ: სისტემის დანიშნულებისამებრ ცვლილებას, მისი ძირითადი მოდელის ახალი ძირითადი ვერსიის მოთავსებას ან მისი შემავალი ან გამომავალი მონაცემების არსებით ცვლილებას.

12. პერსონალის ტრენინგი და ხელოვნური ინტელექტის წიგნიერი შესაძლებლობა

სავალდებულო ტრენინგი: ყველა თანამშრომელს მოეთხოვება გაიაროს საწყისი ტრენინგი ხელოვნური ინტელექტის გამოყენებასთან დაკავშირებით, რომელიც მოიცავს აღნიშნულ პოლიტიკას, მონაცემთა უსაფრთხოებას, რისკების შესახებ ცნობიერებას (მათ შორის დისკრიმინაციას და ჰალუცინაციებს) და დამტკიცებული ხელოვნური ინტელექტის სისტემების სათანადო გამოყენებას. დამატებითი ტრენინგი ჩატარდება საჭიროებისამებრ, მაგალითად, სამართლებრივი რეგულაციების ან ჩვენი ხელოვნური ინტელექტის სისტემების მნიშვნელოვანი განახლებების შემდეგ.

კვალიფიკაციის ამაღლება: უზრუნველყოფს პერსონალს, რომელიც ინტენსიურად მუშაობს ხელოვნურ ინტელექტთან, ფოკუსირებულია ეფექტურ სტიმულსა და დამტკიცებული ხელოვნური ინტელექტის სისტემების მოწინავე ფუნქციებზე.

ხელოვნური ინტელექტის პორტალი: SharePoint-ის შიდა საიტი, რომელიც ხელოვნურ ინტელექტთან დაკავშირებული ყველა საკითხის ცენტრალური რესურსის ცენტრს წარმოადგენს. ის შეიცავს რეესტრს, მომხმარებლის სახელმძღვანელოებს, სასწავლო მასალებს, პოლიტიკის დოკუმენტებს და საკონტაქტო ინფორმაციას.

13. მესამე მხარისა და მომწოდებლების მართვა/მენეჯმენტი

ნებისმიერი საწარმოს ხელოვნური ინტელექტის სისტემის შეძენამდე, გამყიდველმა უნდა გაიაროს საფუძვლიანი სათანადო შემოწმების პროცესი. ამ შეფასებისას შეფასდება გამყიდველის უსაფრთხოების მდგომარეობა, შესაბამისობის სერტიფიკატები, რეპუტაცია და გამჭვირვალობა მისი ხელოვნური ინტელექტის მოდელებთან დაკავშირებით.

დისკუსიისთვის დახურული მოთხოვნა/რომელიც ცვლილებას არ ექვემდებარება: ჩვენ უნდა გვქონდეს ხელმოწერილი მონაცემთა დამუშავების დანართის (DPA) ხელშეკრულება ნებისმიერ მომწოდებელთან, რომლის ხელოვნური ინტელექტის სისტემაც დაამუშავებს ჩვენს არასაჯარო მონაცემებს (დონეები 2-4). ეს DPA უნდა აუკრძალოს მომწოდებელს ჩვენი მონაცემების გამოყენება თავისი მოდელების სწავლების მიზნით.

14. პოლიტიკის მიმოხილვა და აუდიტი

აღნიშნული პოლიტიკა წარმოადგენს აქტიურ დოკუმენტს. AIGC-მ უნდა გადახედოს და განაახლოს იგი მინიმუმ წელიწადში ერთხელ, ან უფრო ხშირად მნიშვნელოვანი ტექნოლოგიური ან საკანონმდებლო ცვლილებების საპასუხოდ. შიდა აუდიტი პერიოდულად ამოწმებს ამ პოლიტიკის დაცვას.

15. სანქციები შეუსრულებლობისთვის

ამ პოლიტიკის დარღვევა ჩაითვლება დასაქმების მოვალეობების სერიოზულ დარღვევად და შეიძლება დაექვემდებაროს დისციპლინურ ზომებს შიდა რეგულაციებისა და მოქმედი კანონმდებლობის შესაბამისად, დასაქმების ან კონტრაქტის შეწყვეტის ჩათვლით.

16. დანართები (ხელმისაწვდომია ხელოვნური ინტელექტის პორტალზე)

- **დანართი A:** დამტკიცებული ხელოვნური ინტელექტის სისტემების რეესტრი
- **დანართი B:** ახალი ხელოვნური ინტელექტის სისტემის დამტკიცების მოთხოვნის ფორმა
- **დანართი C:** მონაცემთა დაცვაზე ზემოქმედების შეფასების (DPIA) შაბლონი
- **დანართი D:** საკონტაქტო ინფორმაცია

დანართი A: დამტკიცებული ხელოვნური ინტელექტის სისტემების რეესტრი

ბოლო განახლება: 2025 წლის 1 ოქტომბრიდან | მხარდაჭერილია: AI Navigator-ის მიერ

აღნიშნული რეესტრი უზრუნველყოფს ჯგუფში გამოსაყენებლად შეფასებული და დამტკიცებული ხელოვნური ინტელექტის სისტემების ზუსტ სიას. ნებისმიერი ხელოვნური ინტელექტის სისტემა, რომელიც არ არის წარმოდგენილი ამ სიაში, აკრძალულად ითვლება, გარდა იმ შემთხვევისა როდესაც ხელოვნური ინტელექტის მმართველობის კომიტეტის (AIGC) მიერ არ არის გაცემული მკაფიო თანხმობა.

ხელოვნური სისტემის დასახელება	კატეგორია	რისკის კლასიფიკაცია	დამოწმებული გამოყენების კლასის მაგალითები	მაქსიმუმი დაშვებული მონაცემთა კლასიფიკაციის	სტატუსი
Adobe acrobat ხელოვნური ინტელექტის ასისტენტი	საჯარო ხელოვნური ინტელექტი	შეზღუდული რისკი	ხელოვნური ინტელექტი სწრაფად უკეთებს ანალიზს თქვენს დოკუმენტაციას საკვანძო/ძირითადი ინფორმაციის ამოღებისთვის, ნიმუშების განსაზღვრისთვის და კომპლექსური პასუხების გასაცემად რომელიც პირდაპირ კავშირშია ტექსტში მოცემულ წყაროებთან.	დონე 3 (კონფიდენციალური) შეზღუდვებით	დამოწმებული
Celsia ხელოვნური ინტელექტის ასისტენტი/დამხმარე	საქმიანობა ხელოვნური ინტელექტი	შეზღუდული რისკი	ფილტრავს შესაბამის ინფორმაციას შიდა ჯგუფის დოკუმენტებიდან, აწარმოებს მოკლე შეჯამებას და პასუხობს ESG - თან დაკავშირებულ კითხვებს.	დონე 4 (მკაცრად კონფიდენციალური)	დამოწმებული

Chat GPT (თავისუფალი/პლიუს ვერსიები)	საჯარო ხელოვნური ინტელექტი	შეზღუდული რისკი	ელფოსტის შედგენა, დოკუმენტების შეჯამება, იდეების გენერირება ან ტექსტის თარგმნა	დონე 3 (კონფიდენციალური) შეზღუდვებით	დამოწმებული
Deepl Pro	საქმიანობა ხელოვნური ინტელექტი	შეზღუდული რისკი	დოკუმენტაციის გადათარგმნა, რომელიც მოიცავს მკაცრად კონფიდენციალურ ინფორმაციას.	დონე 4 (მკაცრად კონფიდენციალური)	დამოწმებული
Google Gemini	საჯარო ხელოვნური ინტელექტი	შეზღუდული რისკი	ელ ფოსტის შედგენა, დოკუმენტების შეჯამება, იდეების გენერირება ან ტექსტის თარგმნა.	მე-3 დონე (კონფიდენციალური) შეზღუდვებით	დამტკიცებული
Google Translate	საჯარო ხელოვნური ინტელექტი	შეზღუდული რისკი	მხოლოდ არა სენსიტიური საჯარო ინფორმაციის თარგმნა	) დონე 2 (შიდა)	დამტკიცებული
Microsoft Copilot for M365	საქმიანობა ხელოვნური ინტელექტი	შეზღუდული რისკი	ელფოსტის შედგენა, დოკუმენტებისა და შეხვედრების შეჯამება, მონაცემთა ანალიზი Microsoft 365 აპლიკაციებში.	დონე 4 (მკაცრად კონფიდენციალური)	დამტკიცებული
NotebookLM	საჯარო ხელოვნური ინტელექტი	შეზღუდული რისკი	კონკრეტული პროექტის დოკუმენტები, კვლევები და შეხვედრების ჩანაწერები გარდაისახოს ინტერაქტიულ ექსპერტად, რომლის მეშვეობითაც შეძლებთ მოითხოვოთ რეზიუმეები,	დონე 3 (კონფიდენციალური)	დამტკიცებული

			ფაქტობრივი პასუხები და ახალი ინფორმაცია, რომელიც მხოლოდ თქვენს პირად წყაროებზეა დაფუძნებული.		
--	--	--	--	--	--

დანართი B: ახალი ხელოვნური ინტელექტის სისტემის დამტკიცების მოთხოვნის ფორმა

გთხოვთ, სრულად შეავსოთ ეს ფორმა და წარუდგინოთ AI Navigator-ს ჯგუფის მიზნებისთვის გამოსაყენებელი ნებისმიერი ახალი AI სისტემისთვის.

ნაწილი 1: ინფორმაციის მომთხოვნა

- სრული დასახელება
- დეპარტამენტი
- სამუშაოს დასახელება
- მოთხოვნის თარიღი:

ნაწილი 2: ხელოვნური ინტელექტის სისტემის დეტალები

- ხელოვნური ინტელექტის სისტემის დასახელება
- მომწოდებელი/
- სისტემის ვებგვერდის ბმული:
- ბმული კონფიდენციალურობის პოლიტიკაზე / მომსახურების პირობებზე:

ნაწილი 3: დანიშნულებისამებრ გამოყენება

- **ბიზნეს მიზანი:** (გთხოვთ, აღწეროთ ბიზნეს პრობლემა, რომლის გადაჭრასაც ცდილობთ და როგორ დაგეხმარებათ ეს ხელოვნური ინტელექტის სისტემა. რა არის ძირითადი მიზანი?)
- **კონკრეტული დავალებები:** (ჩამოთვალეთ კონკრეტული დავალებები, რომელთა შესრულებასაც აპირებთ ხელოვნური ინტელექტის სისტემით, მაგ., „იურიდიული კონტრაქტების თარგმნა“, „მარკეტინგული ასლის გენერირება“, „მომხმარებელთა უკუკავშირის ანალიზი“)

ნაწილი 4: მონაცემთა გამოყენება

- **რა არის მონაცემთა კლასიფიკაციის უმაღლესი დონე, რომლის დამუშავებასაც აპირებთ ამ ხელოვნური ინტელექტის სისტემით?** (იხილეთ ხელოვნური ინტელექტის პოლიტიკის 2.2 ნაწილი)
 - დონე 1: საჯარო
 - დონე 2 - შიდა
 - დონე 3 - კონფიდენციალური
 - დონე 4 - მკაცრად კონფიდენციალური
- **დაამუშავებს თუ არა ხელოვნური ინტელექტის სისტემა რაიმე პერსონალურ მონაცემს GDPR-ით განსაზღვრული წესით?**
 - კი

- არა
- გამოყენებული იქნება თუ არა ხელოვნური ინტელექტის სისტემა ისეთი გადაწყვეტილებების მისაღებად, რომლებსაც მნიშვნელოვანი გავლენა აქვთ ინდივიდებზე (მაგ., დაქირავება, შესრულების შეფასება)?
 - კი
 - არა

ნაწილი 5: საწყისი რისკის შეფასება

- წარმოადგენს თუ არა გამყიდველი მონაცემთა დამუშავების დანართს (DPA), რომელიც აშკარად კრძალავს ჩვენი მონაცემების გამოყენებას მათი ან ნებისმიერი მესამე მხარის მოდელის ტრენინგისთვის?
 - კი
 - არა
 - არ ვიცი
- თქვენი ინფორმაციით, არსებობს თუ არა რაიმე ცნობილი უსაფრთხოების სუსტი მხარეები/წერტილები ან ეთიკური საკითხები, რომლებიც დაკავშირებულია ამ ხელოვნური ინტელექტის სისტემასთან? (გთხოვთ, მოგვაწოდოთ დეტალები არსებობის შემთხვევაში).

მხოლოდ AIGC-ის შიდა გამოყენებისთვის

- მოთხოვნის ID
- მოთხოვნილი მონაცემები:
- ხელოვნური ინტელექტის ნავიგატორის შეფასება
- გადაწყვეტილება: დამტკიცებული/უარყოფილია
- გადაწყვეტილების თარიღი
- დასაბუთება / შეზღუდვები:

დანართი C: ხელოვნური ინტელექტის სისტემის მონაცემთა დაცვის გავლენის შეფასების (DPIA) შაბლონი

აღნიშნული შაბლონი უნდა შეიცავს ნებისმიერი ხელოვნური ინტელექტის სისტემისთვის, რომელიც კლასიფიცირებულია, როგორც მაღალი რისკის მქონე ან ნებისმიერი ხელოვნური ინტელექტის სისტემისთვის, რომელიც ამუშავებს მეოთხე დონის მონაცემებს (მკაცრად კონფიდენციალური), განსაკუთრებით მაშინ, როდესაც პერსონალური მონაცემები მუშავდება დიდი მასშტაბით ან დამუშავება, სავარაუდოდ, მაღალ რისკს შეუქმნის ფიზიკური პირების უფლებებსა და თავისუფლებებს.

ხელოვნური ინტელექტის სისტემის მონაცემთა დაცვის გავლენის შეფასების რეკომენდაცია : უნიკალური ID

ხელოვნური ინტელექტის სისტემის დასახელება: (ხელოვნური ინტელექტის სისტემის დასახელება)

ხელოვნური ინტელექტის სისტემის მფლობელი : სახელი

შეფასების თარიღი: თარიღი

ნაწილი 1: დამუშავების აღწერა

- დამუშავების ბუნება/არსი ხასიათი და ფარგლები: აღწერეთ ხელოვნური ინტელექტის სისტემის ფუნქციონალურობა, მის მიერ დამუშავებული მონაცემების არსი, მონაცემთა მოცულობა და შემავალი მონაცემების კატეგორიები
- **დამუშავების მიზანი:** რა არის ჯგუფის, მომხმარებლების ან თანამშრომლებისთვის მოსალოდნელი შედეგები და სარგებელი?
- დამუშავების კონტექსტი: აღწერეთ მონაცემთა სუბიექტებთან/საგანთან ურთიერთობა და მათი გონივრული მოლოდინები. დეტალურად მიუთითეთ მონაცემთა წყაროები და ნებისმიერი მონაცემების გადინება/ მესამე მხარეებისთვის.

ნაწილი 2: აუცილებლობისა და პროპორციულობის შეფასება

- **კანონიერი საფუძველი:** რა არის პერსონალური მონაცემების დამუშავების კანონიერი საფუძველი GDPR-ის მე-6 მუხლის შესაბამისად?
- **მიზნობრივი შეზღუდვა:** მონაცემები მუშავდება თუ არა მკაცრად განსაზღვრული, ცალსახა და ლეგიტიმური მიზნებისთვის?
- **მონაცემთა მინიმოზიზაცია:** არის თუ არა დამუშავებული მონაცემები ადეკვატური, რელევანტური და ლიმიტირებული იმით, რაც აუცილებელია იმ მიზნებთან მიმართებაში, რისთვისაც ისინი მუშავდება?

ნაწილი 3: რისკის შეფასება

- **მონაცემთა საგნების უფლებებისა და თავისუფლებების რისკები:** იდენტიფიცირება პოტენციური რისკების, როგორიცაა:

- უნებართვო წვდომა ან მონაცემთა დარღვევა/გაჟონვა
- არასწორი შედეგები, რაც არასწორ გადაწყვეტილებებს იწვევს.
- დისკრიმინაციული მიკერძოება შედეგებში (მაგ., დაქირავების ან საკრედიტო ქულების შეფასებისას).
- ხელოვნური ინტელექტით მართულ გადაწყვეტილებებში გამჭვირვალობის ან ახსნის ნაკლებობა.
- ანონიმური ან ფსევდონიმური მონაცემების ხელახალი იდენტიფიცირება.
- **ალბათობა და სიმძიმე:** თითოეული იდენტიფიცირებული რისკისთვის შეაფასეთ მისი ალბათობა (დაბალი/საშუალო/მაღალი) და ზემოქმედების პოტენციური სიმძიმე.

ნაწილი 4: რისკების შესამცირებლად გათვალისწინებული ზომები

- **ტექნიკური ზომები:** აღწერეთ არსებული უსაფრთხოების კონტროლის მექანიზმები (მაგ., დაშიფვრა, წვდომის კონტროლი, ფსევდონიმიზაცია).
- **ორგანიზაციული ზომები:** აღწერეთ პროცედურული კონტროლის მექანიზმები (მაგ., ადამიანური ზედამხედველობა, რეგულარული აუდიტი, პერსონალის ტრენინგი, მომწოდებლის მონაცემთა დაცვის სააგენტო).
- **მიკერძოების შემცირება:** აღწერეთ ნაბიჯები რომლებიც გამოყენებულია ალგორითმული მიკერძოების შესამოწმებლად და შესამცირებლად.
- **გამჭვირვალობის ზომები:** როგორ იქნება მონაცემთა საგნების ინფორმირება ხელოვნური ინტელექტის სისტემის გამოყენებისა და მათი უფლებების თაობაზე?

ნაწილი 5: კონსულტაცია და დამტკიცება

- **კონსულტაცია მონაცემთა დაცვის ოფიცერთან:** ჩაიწერეთ/ჩაინიშნეთ მონაცემთა დაცვის ოფიცრის მიერ მოწოდებული რჩევა.
- **ხელოვნური ინტელექტის გენერირებული შინაარსის/კონტენტის გადაწყვეტილება:**
 - დამუშევრით გაგრძელება
 - დამატებითი ზომების შესრულებისთვის საგნის/სუბიექტის დამუშავების გაგრძელება/განხორციელება
 - ნუ გააგრძელებთ დამუშავებას.
 - გაიარეთ კონსულტაცია საზედამხედველო ორგანოსთან.

სისტემიდან გამოსვლა/დამამთავრებელი ნაწილი

- **ხელოვნური ინტელექტის სისტემის მფლობელი** -----

- მონაცემთა დაცვის ოფიცერი
- AI Governance Committee AIGC ხელმძღვანელი

ხელოვნური ინტელექტის მართვის კომისიის ხელმძღვანელი

დანართი D: საკონტაქტო ინფორმაცია

აღნიშნულ პოლიტიკასთან ან ჯგუფში ხელოვნური ინტელექტის გამოყენებასთან დაკავშირებული ნებისმიერი შეკითხვის შემთხვევაში, გთხოვთ, დაუკავშირდეთ შესაბამის პირებს.

ძირითადი საკონტაქტო პირი

- **ხელოვნური ინტელექტის ნავიგატორი**
 - სახელი: Tomáš Brandejský
 - ელექტრონული ფოსტა: t.brandejsky@energo-pro.com
 - პასუხისმგებლობა: პერსონალის პირველი რიგის მხარდაჭერა, ტრენინგების კოორდინაცია და ხელოვნური ინტელექტის სისტემის დამტკიცების პროცესის მართვა.

მმართველობა და ესკალაცია

- **ხელოვნური ინტელექტის მართვის კომისია (AIGC)**
 - თავმჯდომარე: Jakub Fajfr
 - გენერალური კონსული: Christian Blatchford
 - მონაცემთა დაცვის ოფიცერი (DPO) : Marina Radeva
 - კადრების უფროსი : Nikola Šugra
 - ელექტრონული ფოსტა : AIGC@energo-pro.com
 - პასუხისმგებლობები: სტრატეგიული ზედამხედველობა, დამტკიცებული ხელოვნური ინტელექტის სისტემების რეესტრის დამოწმება და რეგულარული განახლება, ხელოვნური ინტელექტის პოლიტიკის დამტკიცება, რისკების მართვა/მენეჯმენტი და მაღალი რისკის მქონე ხელოვნური ინტელექტის სისტემების საბოლოო დამტკიცება.