

Providing a Structural Methodology for Measuring and Analyzing the Quality of Theses and Dissertations in the Country

Mohammad Javad Ershadi*

PhD in Industrial Engineering; Associate Professor; Information Technology Department; Iranian Research Institute for Information Science and Technology (IranDoc); Tehran, Iran;
Email: Ershadi@irandoc.ac.ir

Masoumeh Nabizadeh

MSc in Industrial Engineer, Industrial Engineering Department; Islamic Azad University; Science and Research Branch; Tehran, Iran Email: nabizadeh797@gmail.com

Received: 03, Aug. 2020 Accepted: 20, Apr. 2021

Abstract: Today, measuring data quality is one of the most important strategies in improving data-driven business processes. Any correct decision to improve systems in this category of organizations depends on an appropriate analysis of data quality. Research data and especially theses/dissertations of graduates of the whole country based on this principle from various aspects of data quality need to be reviewed and evaluated. In the process of registering, the quality control of metadata is one of the most important parts that examines the information items of documents (such as the name of the researcher as well as supervisors and advisors, abstract, index, etc.). In the current situation, incompatible documents are identified during the quality control process and after entering in the system (as text), the relevant document is systematically returned to the researcher. Lack of a standard framework for identified non-conformities besides absence of an appropriate partitioning, make statistical analysis of metadata quality problems difficult, and also make it impossible to analyze the root of observed errors. Therefore, in this study, the incompatible structure observed after standardization has been used experimentally in two-month periods and the results have been presented. Incompatible in thesis defense date, the title page (Persian and English) and the existence of white pages in the file have been among the most important reasons for returning documents to users. Also, the analysis of all data showed that 59% of the incompatibles were related to the attached files and 41% to the information recorded in the system. Finally, based on the performed analysis, guidelines have been provided for the users of the system in order to improve the quality of data, according to specialized areas.

Keywords: Data Quality, Information Quality, Quality Control, Iranian Research Institute for Information Science and Technology (IranDoc), Theses / Dissertations Registration System, Incompatible Analysis

* Corresponding Author

**Iranian Journal of
Information
Processing and
Management**

**Iranian Research Institute
for Information Science and Technology
(IranDoc)**

ISSN 2251-8223

eISSN 2251-8231

Indexed by SCOPUS, ISC, & LISTA

Vol. 37 | No. 3 | pp. 667-694

Spring 2022

<https://www.doi.org/10.35050/JIPM010.2022.256>



ارائه یک روش ساختارمند در اندازه گیری و تحلیل کیفیت اطلاعات پایان نامه ها و رساله های درون کشور

محمدجواد ارشادی

دکتری مهندسی صنایع؛ دانشیار؛ پژوهشگاه علوم
و فناوری اطلاعات ایران (ایرانداک)؛ تهران، ایران؛
Ershadi@irandoc.ac.ir

معصومه نبی زاده

کارشناسی ارشد مهندسی صنایع؛ گروه مهندسی صنایع؛
دانشکده فنی-مهندسی؛ دانشگاه آزاد اسلامی؛
واحد علوم و تحقیقات؛ تهران، ایران؛
nabizadeh797@gmail.com



مقاله برای اصلاح به مدت ۴ ماه نزد پدیدآوران بوده است.

پذیرش: ۱۴۰۰/۰۱/۳۱

دریافت: ۱۳۹۹/۰۵/۱۳

نشریه علمی | رتبه بین المللی
پژوهشگاه علوم و فناوری اطلاعات ایران
(ایرانداک)

شاپا (جایی) ۸۲۲۳-۲۲۵۱

شاپا (الکترونیکی) ۸۲۳۱-۲۲۵۱

نمایه در SCOPUS و ISC، LISTA و

jipm.irandoc.ac.ir

دوره ۳۷ | شماره ۳ | صص ۶۶۷-۶۹۴

بهار ۱۴۰۱

<https://www.doi.org/10.35050/JIPM010.2022.256>

JIPM010.2022.256



چکیده: امروزه، اندازه گیری کیفیت داده یکی از مهم ترین راهبردها در بهبود فرایندهای کسب و کارهای داده محور به حساب می آید. هر گونه تصمیم درست برای بهبود سیستم ها در این دسته از سازمان ها به یک تحلیل مناسب از کیفیت داده ها وابسته است. داده های پژوهشی و به ویژه، پایان نامه / رساله ها (پارساهای دانش آموختگان کل کشور بر پایه همین اصل از جنبه های گوناگون کیفیت داده نیازمند بررسی و ارزیابی است. در فرایند ثبت «پارساهای کنترل کیفیت فراداده ها یکی از مهم ترین بخش هاست که به بررسی اقلام اطلاعاتی مدارک (مانند نام پژوهشگر، استادان راهنما و مشاور، چکیده، فهرست و ...) می پردازد. در وضعیت موجود، مدارک نا همخوان در فرایند کنترل کیفیت شناسایی شده و مدرک مربوطه پس از درج در سامانه (به صورت متن) به صورت سیستمی به پژوهشگر برگردانده می شود. استاندارد نبودن نا همخوان های شناسایی شده و نبود دسته بندی مناسب برای مدارک نا همخوان به اطلاع رسانی سلیقه ای به پژوهشگران منجر شده، تحلیل های آماری از مشکلات کیفی فراداده ها را با دشواری روبه رو ساخته، و تحلیل ریشه ای خطاهای مشاهده شده را ناممکن می سازد. از این رو، در این پژوهش ساختار نا همخوان های مشاهده شده پس از استاندارد شدن به صورت آزمایشی در دوره های دو ماهه استفاده شده و نتایج آن ارائه شده است. نا همخوانی در تاریخ

دفاع، صفحه عنوان (فارسی و انگلیسی) و وجود صفحه‌های سفید در فایل «پارسا» از جمله مهم‌ترین دلایل بازگرداندن مدارک به کاربران بوده است. همچنین، تحلیل همه داده‌ها نشان داد که ۵۹ درصد از ناهمخوانی‌ها به فایل‌های ضمیمه‌شده و ۴۱ درصد به اطلاعات ثبت‌شده در سامانه مربوط می‌شود. در نهایت، بر پایه تحلیل‌های انجام‌شده، رهنمودهایی به‌منظور بهبود کیفیت داده‌ها به تفکیک حوزه‌های تخصصی برای کاربران سامانه ارائه شده است.

کلیدواژه‌ها: کیفیت داده، کیفیت اطلاعات، کنترل کیفیت، پژوهشگاه علوم و فناوری اطلاعات ایران (ایرانداک)، سامانه ثبت پایان‌نامه/ رساله، تحلیل ناهمخوان

۱. مقدمه

در دنیای امروز، کیفیت داده و اطلاعات یکی از مهم‌ترین دغدغه‌های سازمان به حساب آمده و به حوزه‌ای بسیار تأثیرگذار در سامانه‌های اطلاعاتی تبدیل شده است. رشد شدید مخزن‌های داده و دسترسی مستقیم مدیران و کاربران به منابع مختلف داده از یک سو و نیاز کاربران به اطلاعات باکیفیت، ضرورت توجه به مقوله کیفیت داده و آگاهی نسبت به آن را افزایش داده است. در بررسی‌های انجام‌شده، پژوهشگران کیفیت داده را یکی از ۶ بعد اصلی در سامانه‌های اطلاعاتی مدیریت (MIS)^۱ معرفی کردند. سامانه‌های اطلاعاتی تحقیقاتی (RIS)^۲ به‌عنوان یکی از شناخته‌شده‌ترین سامانه‌های اطلاعاتی مدیریت در حوزه پژوهش به ثبت، کنترل، و اشاعه پژوهش‌ها پرداخته و امروزه، مدیریت و سیاست‌گذاری حوزه پژوهش بدون آن قابل تصور نیست. از این رو، کیفیت داده‌ها و اطلاعات پژوهشی در یک RIS و همچنین، ارزیابی و اندازه‌گیری آن یک راهبرد مدیریتی به حساب می‌آید. سامانه «گنج» یکی از کلیدی‌ترین سامانه‌ها در اشاعه برون‌دادهای پژوهش در کشور به حساب می‌آید. به همین دلیل، بهبود کیفیت نتایج جست‌وجوی این سامانه، یکی از راهبردهای اساسی در ارتقای عملکرد آن است. ریشه مشکلات کیفی مشاهده‌شده در نتایج به‌دست‌آمده از فرایند جست‌وجو و همچنین، افزایش رضایت کاربران سامانه «گنج» به کارکرد صحیح و باکیفیت سامانه ثبت وابسته است که در آغاز کلان فرایند ثبت، سازماندهی و اشاعه اطلاعات قرار دارد. بر پایه طرح کیفیت سامانه ثبت، ناهمخوانی‌های مشاهده‌شده در مدارک شناسایی شده و به پژوهشگر اطلاع‌رسانی می‌شود. اداره ثبت و فراهم‌آوری اطلاعات در «پژوهشگاه علوم و فناوری اطلاعات ایران» مسئولیت شناسایی

1. management information systems (MIS)

2. research information systems (RIS)

ناهمخوانی‌ها و اقدامات اصلاحی یا اصلاح بر روی مدارک ناهمخوان شناسایی شده را بر عهده دارد. در وضعیت موجود، موارد ناهمخوان مشاهده شده توسط کارکنان اداره ثبت و فراهم‌آوری اطلاعات در فرایند ثبت مدارک «پارسا» به صورت ساختارنیافته است. بدین معنا که هر یک از کارکنان شاغل در فرایند، پس از مشاهده هرگونه ناهمخوانی در مدارک ثبت شده توسط دانش‌آموختگان، بر پایه دانش خود و شناختی که از وضعیت مطلوب مدارک دارند، ناهمخوانی‌ها را در سامانه ثبت در فضای اختصاص داده شده به این کار ثبت می‌کنند. از سوی دیگر، در ساختار موجود سامانه ثبت هیچ گونه دسته‌بندی مشخصی برای ثبت موارد ناهمخوانی در مدارک «پارسا» وجود ندارد. این وضعیت، مشکلات زیر را از جنبه‌های فرایندی و مدیریتی به همراه دارد.

◇ وضعیت موجود به دلیل ناهمخوانی با استانداردهای یک مدرک امکان تحلیل مناسب از علت بازگرداندن مدارک را به دست نمی‌دهد. این موضوع در حالی است که در صورت تحلیل درست از علت ناهمخوانی مدارک و شناسایی فراوانی هر علت می‌توان راهکارهای بهبود فرایند و سامانه ثبت را به صورت کاراتر و اثربخش‌تر ارائه کرد؛

◇ در وضعیت کنونی، دانش موجود در اداره ثبت و فراهم‌آوری اطلاعات در رابطه با مشکلات مدارک «پارسا»ها ثبت و اشاعه داده نمی‌شود. همچنین، در صورت مشاهده مشکلات جدید امکان استاندارد کردن آن و قراردادن در سامانه فراهم نیست؛

◇ امکان آموزش کامل نفرات جدید (در صورت نیاز) و فراهم کردن چارچوبی برای توسعه ظرفیت سامانه در کنترل تعداد بیشتری مدرک فراهم نیست. این بدان معناست که در وضعیت موجود در صورت افزایش ظرفیت سامانه و کمک گرفتن از نفرات جدید لازم است هر فرد جدید در کنار سایر افراد قرار گیرد و فقط در صورتی که فرد آموزش دهنده از دقت و حوصله فراوان برخوردار باشد، می‌تواند دانش خود را در رابطه با مشکلات کیفی مدارک و علت‌های بازگرداندن آن‌ها، به فرد جدید منتقل کند. این روش سیستمی نبوده و امکان خطا و فراموش شدن برخی علت‌ها را به همراه دارد.

افزون بر موارد بالا، استاندارد کردن و تحلیل بر روی علت‌های ناهمخوانی در سامانه ثبت، یکپارچه کردن فرایندهای ثبت و نمایه‌سازی را در بلندمدت با راحتی بیشتری مواجه خواهد کرد و توسعه کلان سامانه‌ها را نیز به همراه خواهد داشت. همچنین، تحلیل‌های

دقیق‌تر از وضعیت دانشگاه‌های مختلف در رابطه با کیفیت مدارک ثبت‌شده توسط آن‌ها را آسان خواهد کرد. این پژوهش بنا دارد با توجه به مزیت‌های فراوان استانداردسازی علت‌ها در بازیابی مدارک «پارسا»، به ارائه چارچوبی استاندارد برای ناهمخوانی‌های مشاهده‌شده در مدارک (پارساها)ی دانش‌آموختگان کل کشور پردازد و تحلیل مناسبی از چگونگی ناهمخوانی‌ها و میزان آن‌ها ارائه دهد.

۲. پیشینه پژوهش

پژوهش‌های پیشین بر این حقیقت تأکید دارند که کیفیت داده، حوزه‌های مربوط به انواع داده‌های نیمه‌ساختاریافته، ساختارنیافته و چندرسانه‌ای را به‌طور کلی مورد توجه قرار می‌دهد. پژوهش‌های انجام‌شده متمرکز بر روی داده‌های نیمه‌ساختاریافته و ساختارنیافته در حوزه کیفیت داده به ارتباط تاریخی محکم بین کیفیت داده و طرح و ساختار پایگاه داده اشاره دارند. حتی رویکردهای یکپارچه با اینکه دارای سوگیری هستند، اما تمرکز اساسی آن‌ها بر مجموعه‌ای از داده‌های ساختاریافته است که ارائه‌دهنده بیشترین منابع اطلاعاتی در سازمان‌ها هستند (Avenali et al. 2008). امروزه، کاربرد داده‌های نیمه‌ساختاریافته و فاقد ساختار به‌عنوان منابع سازمانی بسیار رایج‌تر است. مدیریت دانش، حفاظت وب و سیستم‌های اطلاعات جغرافیایی دربرگیرنده حوزه‌های تحقیقاتی مهمی در چارچوب کیفیت داده هستند. افزون بر آن، تکنیک‌های کیفیت داده برای داده‌های نیمه‌ساختاریافته و فاقد ساختار نیز مورد بررسی قرار گرفته‌اند (Rahman et al. 2018). بهبود تکنیک‌های کیفیت داده برای داده‌های فاقد ساختار و نیمه‌ساختاریافته در این حوزه‌ها نیازمند درجه بالاتری از ارتباط میان تخصص‌های مختلف است که البته، ممکن است بهبود کیفیت داده را طولانی سازد (Taleb, Serhani & Dssouli 2018).

با توجه به ابعاد گوناگون کیفیت اطلاعات، نگرانی از کیفیت و قابلیت اعتماد اطلاعات نه تنها در بین جست‌وجوگران آنلاین و متخصصان اطلاعات، بلکه در بین عموم مردم رو به افزایش است (Russell-Rose, Chamberlain, & Azzopardi 2018). متخصصان اطلاعات بر پایه این اصل، به مشتریان در فهم محدودیت‌های منابع داده کمک می‌کنند و به پرورش تفکر انتقادی در قبال داده‌ها می‌پردازند، و همچنین دقت و اعتبار اطلاعات جمع‌آوری‌شده را ارزیابی می‌کنند. از این رو امروزه، بهبود کیفیت، به روشنی موضوعی مهم و حیاتی بوده و لازم است به‌عنوان یک استراتژی کلیدی سازمان‌ها مورد توجه قرار

گیرد. در تمام زنجیره تولید داده از تولیدکننده پایگاه داده گرفته تا مصرف‌کننده نهایی آن ایجاد همکاری و یک تعهد همه‌جانبه در سیستم برای بهبود پیوسته کیفیت ضروری و اجتناب‌ناپذیر است.

ارتباط بین کیفیت داده و کیفیت فرایند به دلیل ارتباط و تنوع ویژگی‌های فرایندهای کسب‌وکار در سازمان‌ها، بخش عمده‌ای از تحقیقات را دربرمی‌گیرد. تأثیرات متفاوت کیفیت داده‌ها در سه سطح مشخص به نام‌های عملیاتی، تاکتیکی و استراتژیک مورد بررسی قرار گرفته است (Batini et al. 2009). کیفیت داده و ارتباط آن با کیفیت سرویس‌ها، محصولات، عملیات کسب‌وکار و رفتار مصرف‌کننده در قالب اصطلاحات عمومی زیادی در این پژوهش مورد بررسی قرار گرفته است. مدل‌های استاندارد و رایج در کیفیت داده مانند مدل Wang & Strong (1996) و مدل انستیتو انبار داده^۱ در سال ۲۰۰۶ و همچنین، چارچوب ارائه‌شده توسط Vaziri, Mohsenzadeh, & Habibi (2017) و نیز Khosroanjom et al. (2011) پژوهش‌هایی هستند که در حوزه کیفیت داده جنبه‌های گوناگونی را مورد توجه قرار داده‌اند^۲. همچنین، امروزه، مدل‌هایی نیز مانند مدل ارائه‌شده توسط (Kwon, Lee & Shin 2014) برای کیفیت کلان‌داده‌ها در سازمان ارائه شده است. چارچوب اصلی دیگری که در پیشینه پژوهش وجود دارد، دسته‌بندی گسترده‌ای است که در پژوهش «ارشادی، جلالی‌منش و نصیری» (۱۳۹۸) برای انواع ناهمخوانی‌ها در سامانه ثبت صورت گرفته است. افزون بر موارد گفته‌شده، دسته دیگری از پژوهش‌های پیشین به اندازه‌گیری کیفیت داده به‌عنوان یک هدف نگرین‌سته و گزارش‌های گوناگونی را در حوزه اندازه‌گیری جنبه‌های آن ارائه کرده‌اند. حوزه مشتریان و اندازه‌گیری داده‌های آن یکی از این پژوهش‌هاست که در حال حاضر مورد توجه برخی از پژوهشگران (Peltier, Zahay 2006) مدلی برای اندازه‌گیری آن ارائه کردند. این مدل در پژوهش‌های تازه‌تر با تمرکز بر فراداده‌ها توجه افزون‌تری یافته است (Nikiforova 2020). چارچوب مشابهی در داده‌های ثبت اسناد و املاک به‌منظور ارزیابی کیفیت پایگاه داده ارائه شده و ابعادی مانند به‌روز بودن، دقت و هزینه داده‌های بی‌کیفیت در آن مورد توجه قرار گرفته است (Heinrich, Klier 2009).

1. Data warehouse institute

۲. خوانندگان برای مطالعه کامل انواع مدل‌های استاندارد کیفیت داده می‌توانند به کتاب (ارشادی و ارشادی ۱۳۹۸) مراجعه کنند.

اگرچه در داده‌های سامانه‌های اطلاعاتی تحقیقاتی برخی پژوهش‌ها در قالب نمودار کنترل به پایش آن‌ها پرداخته (اشتریان اصفهانی، ارشادی و عزیزی ۱۳۹۹) یا تحلیل ریشه ارائه کرده‌اند (ارشادی و همکاران ۱۳۹۵)، اندازه‌گیری و تحلیل این داده‌ها در قالب یک روش ساختارمند که بتواند برای بهبود فرایندهای تولید و پردازش داده به تصمیم‌های مدیریتی منجر شود، مورد توجه قرار نگرفته است. همچنین، ویژگی متمایزکننده‌ای که ارزیابی کیفیت داده‌ها در سامانه‌های اطلاعاتی تحقیقاتی نسبت به سایر سامانه‌ها دارند، این است که کنترل‌ها و نتایج آن‌ها در قالب متن بوده و ساختار مشخص و از پیش تعیین‌شده‌ای ندارند. این کار دسته‌بندی ویژگی‌های کیفی در قالب ساختارهای استاندارد مانند دقت، کامل بودن و به‌روز بودن را دشوارتر ساخته و لزوم ارائه یک ساختار بومی برای ارزیابی را تقویت می‌کند. برای نمونه، کارشناس کنترل کیفیت یک فایل متنی پایان‌نامه نواقص مشاهده شده در فایل را به‌صورت چند عبارت در سامانه درج می‌کند که نداشت آن به ابعاد استاندارد کیفیت داده را بسیار دشوار می‌سازد.

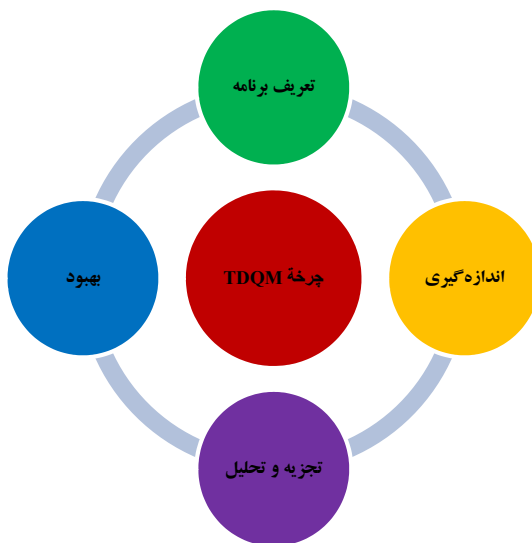
بنابراین، به‌منظور ارائه یک تحلیل آماری مناسب از کیفیت مدارک علمی یا پژوهشی، یک ساختار درست و از پیش تعیین‌شده نیاز است که قابلیت تکرارپذیری نیز داشته باشد. افزون بر این، چارچوب ارائه‌شده در این پژوهش می‌تواند پایه و اساس پژوهش‌های کیفیت داده در سایر دسته‌های پایگاه داده مانند مجلات، کتاب‌ها و کنفرانس‌ها نیز قرار گیرد. این است که در این مقاله ناهمخوانی‌ها شناسایی و طبقه‌بندی شده و تحلیل‌های آماری بر روی نتایج این ناهمخوانی‌ها و رهنمودهای اجرایی ارائه خواهد شد. در بخش بعد، روش پژوهش مورد توجه قرار می‌گیرد.

۳. روش پژوهش

همان‌گونه که در بخش پیشین اشاره شد، ارائه یک چارچوب مشخص و مدون برای ارزیابی و تحلیل کیفیت مدارک پایان‌نامه‌ها/رساله‌ها برای ایجاد توسعه و بهبود کیفیت داده‌های این حوزه کاری حیاتی است. رویکردی که در این پژوهش در ادامه بسیاری از پژوهش‌های پیشین برای ارزیابی و بهبود کیفیت داده استفاده شده، بر چارچوب مدیریت کیفیت فراگیر داده^۱ (TDQM) مبتنی است که در صورت پیاده‌سازی درست می‌تواند بهبود مستمر کیفیت داده را در سازمان به‌دنبال داشته باشد (Wang & Madnick 1990). چارچوب

1. total data quality management (TDQM)

TDQM بهبود کیفیت مستمر داده‌ها را با دنبال کردن مراحل چهارگانه برنامه‌ریزی، اندازه‌گیری، تجزیه و تحلیل، و بهبود پشتیبانی می‌کند (شکل ۱). این چارچوب پایه و اساس بسیاری از برنامه‌های بهبود کیفیت داده قرار گرفته و امروزه نیز در بسیاری از پژوهش‌ها مورد استفاده است (Cahyono & Sucahyo 2020; Liu et al. 2020).



شکل ۱. چرخه TDQM (Wang & Strong 1996)

هر یک از مراحل اصلی TDQM (بر پایه شکل ۱) شامل فعالیت‌های گوناگونی است که در ادامه، به مهم‌ترین آن‌ها اشاره خواهد شد.

♦ تعریف برنامه

تعریف ابعاد کیفیت داده از دیدگاه مصرف‌کننده یا کاربر داده یک گام کلیدی در مدیریت کیفیت جامع داده به حساب می‌آید. در این مرحله، چارچوب اصلی ارزیابی کیفیت داده، شاخص‌ها و چگونگی ارزیابی آن‌ها مشخص می‌شود.

♦ اندازه‌گیری

در این مرحله بر پایه چارچوب و گام‌های مرحله پیشین، کیفیت داده‌ها اندازه‌گیری می‌شود. شاخص‌های تعیین‌شده در مرحله تعریف برنامه، به صورت عملیاتی و در بازه زمانی مشخص‌شده، در این مرحله اندازه‌گیری خواهند شد.

♦ تجزیه و تحلیل

تحلیل نتایج اندازه‌گیری کیفیت داده در گام سوم صورت خواهد پذیرفت. میزان تفاوت‌ها میان شاخص‌های کیفیت داده و جایگاه از قبل تعریف‌شده آن‌ها به کمک تکنیک‌های مناسب در این مرحله تعیین خواهد شد.

♦ بهبود

سرانجام، در گام چهارم، برنامه‌های بهبود کیفیت داده مدون خواهند شد. در این مرحله، اقدامات صورت گرفته از دو بعد حائز اهمیت است: ۱) تغییر ارزش داده‌ها، و ۲) تغییر فرایندهای تولید داده.

این پژوهش بر پایه رویکرد TDQM در چهار گام و در قالب چارچوب زیر انجام خواهد شد. در گام اول، چارچوبی برای اندازه‌گیری کیفیت فراداده‌ها تعیین خواهد شد. در این گام، پس از تعیین چارچوب استاندارد ارزیابی کیفیت مدارک در سامانه ثبت «پارسا»ی درون کشور، ساختار ارائه‌شده توسط خبرگان حوزه کیفیت داده‌های علم و فناوری و روایی آن ارزیابی خواهد شد. سپس، در گام دوم، بر پایه چارچوب تعیین‌شده در گام اول، کیفیت داده‌ها و فراداده‌ها در سامانه ثبت اندازه‌گیری می‌شود. در گام سوم، نتایج به‌دست آمده در گام دوم، مورد تجزیه و تحلیل قرار خواهند گرفت. و در نهایت، در گام چهارم، پیشنهادهایی برای بهبود کیفیت داده‌ها بر پایه تحلیل‌های انجام‌شده ارائه خواهد شد.

شکل ۲، روشی را که این پژوهش بر پایه آن انجام شده به صورت خلاصه نمایش می‌دهد.



شکل ۲. گام‌های پژوهش برای اندازه‌گیری و تحلیل کیفیت مدارک «پارسا»ی درون کشور

در ادامه این پژوهش، در بخش چهارم، نتایج پژوهش در قالب چهار گام پیش‌گفته نشان داده خواهد شد.

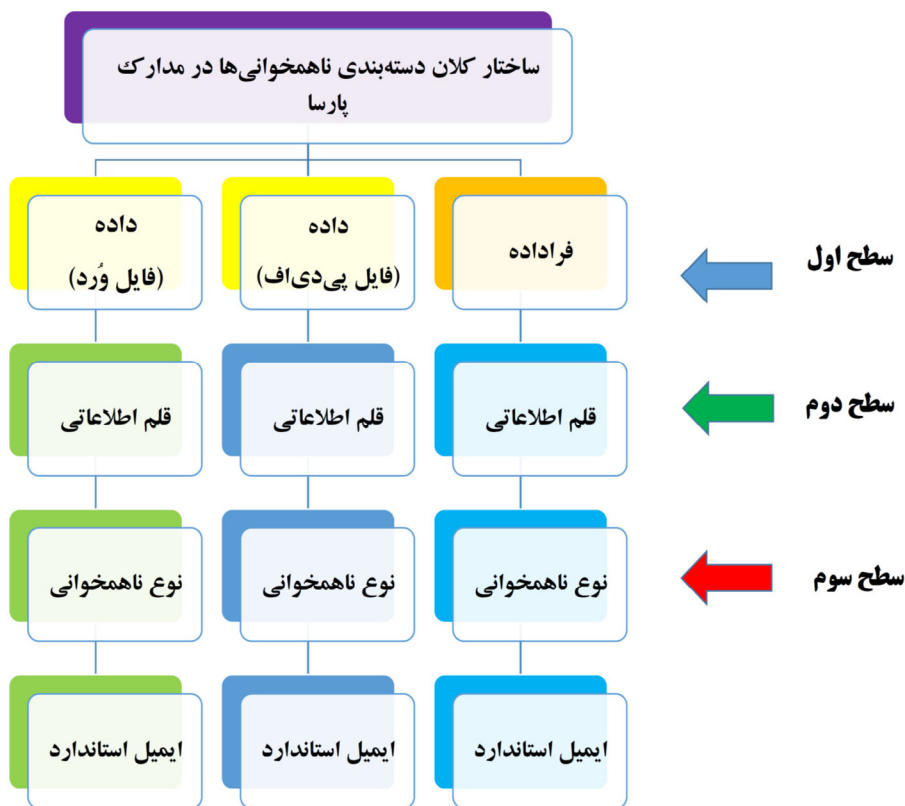
۴. نتایج

در این بخش بر پایه چهار گام ارائه‌شده در بخش سوم، در قالب چهار زیربخش و متناظر با هر گام، به ارائه نتایج خواهیم پرداخت.

۴-۱. ارائه چارچوب کلان ساختار دسته‌بندی موارد ناهمخوان

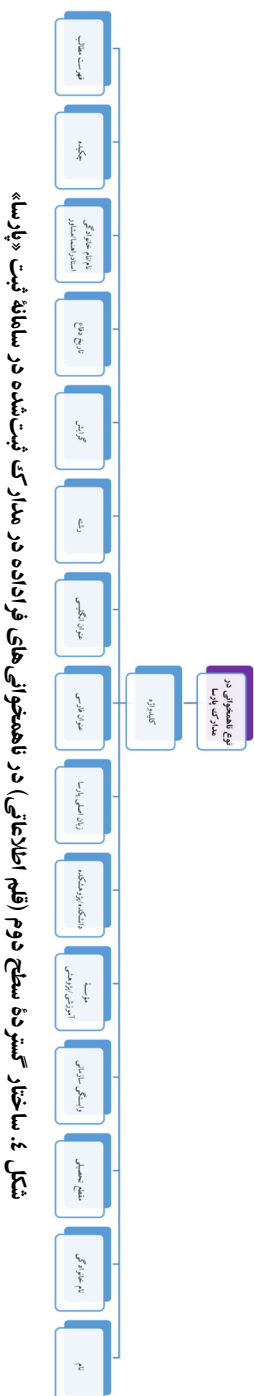
در وضعیت کنونی، کارکنان مدیریت ثبت و فراهم‌آوری اطلاعات ابتدا فراداده ثبت‌شده توسط کاربر (دانشجو) در سامانه ثبت را کنترل می‌کنند. سپس، فایل پی‌دی‌اف و سرانجام فایل ورد کنترل می‌شود. بر این اساس، ناهمخوانی‌های مشاهده‌شده در مدارک «پارسا» نیز به این سه دسته کلی تقسیم‌بندی می‌شود. به بیان دیگر، موارد ناهمخوانی مشاهده‌شده در سامانه ثبت یا به اطلاعات درج‌شده در سامانه توسط کاربر (فراداده) برمی‌گردد (دسته اول) یا به فایل ارسالی ورد (دسته دوم) و یا به فایل ارسالی پی‌دی‌اف (دسته سوم) مربوط می‌شود. از این رو، ناهمخوانی‌ها در سه دسته اصلی تقسیم‌بندی شده‌اند. در هر یک از سه گروه فراداده، داده-فایل ورد و داده-فایل پی‌دی‌اف موارد

ناهمخوانی در اقسام اطلاعاتی ثبت شده توسط کاربر قابل دسته‌بندی خواهد بود. شکل ۳، دسته‌بندی توضیح داده شده را به تصویر می‌کشد. فراداده شامل اطلاعات ثبت شده در سامانه ثبت توسط کاربران سامانه است و شامل قلم‌های اطلاعاتی مختلفی مانند نام، نام خانوادگی، نام استاد راهنما و ... است. سطح دوم، همان گونه که در شکل ۳، نشان داده شده، ناظر به قلم‌های اطلاعاتی است. در هر قلم اطلاعاتی ممکن است خطاها و اشتباه‌هایی توسط کاربر رخ دهد. به عبارتی، نوع ناهمخوانی مشاهده شده به قلم اطلاعاتی وابسته است (سطح سوم). برای نمونه، کاربر ممکن است در بخش اطلاعات کتابشناختی سامانه ثبت (فراداده) نام خود را اشتباه تایپ کند یا پسوند نام خود را درج نکند. پس، این ناهمخوانی در سطح اول به فراداده مربوط می‌شود. همچنین، قلم اطلاعاتی «نام»، سطح دوم این ناهمخوانی را تشکیل می‌دهد. سطح سوم، «عدم درج پسوند» نوع ناهمخوانی خواهد بود. در ادامه، ایمیلی نیز که به کاربر ارسال می‌شود، نشان‌دهنده نوع ناهمخوانی مشاهده شده است. بر این اساس، آخرین بخش از ساختار ناهمخوانی‌ها به ایمیل ارسالی به کاربر اختصاص داده شده است. به بیان دیگر، برای استاندارد شدن فرایند آگاهی‌رسانی به دانشجوی ثبت‌کننده «پارسا»، پس از انتخاب نوع ناهمخوانی، متن استاندارد برای کاربر ارسال خواهد شد تا بتواند مدرک خود را بر پایه توضیحات ایمیل ویرایش کند.

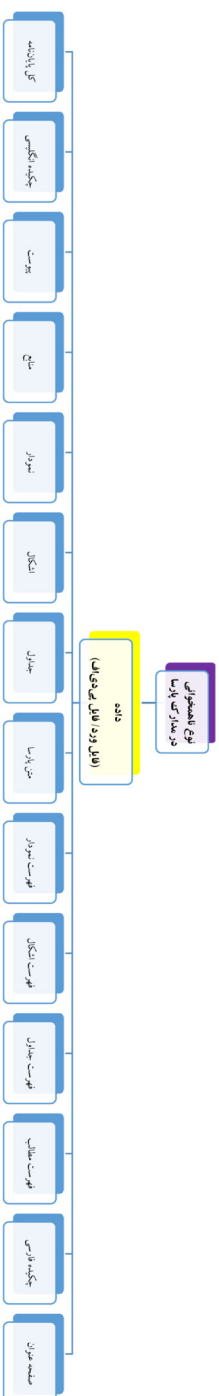


شکل ۳. ساختار کلان موارد ناهمخوان در مدارک ثبت‌شده در سامانه ثبت «پارسا»

همان‌گونه که در ساختار شکل ۳، مشاهده می‌شود، در سطح اول، ساختار یک ناهمخوانی تعیین‌کننده داده یا فراداده مدرک «پارسا» مورد بررسی قرار گرفته است. در این سطح فقط مشخص می‌شود که یک ناهمخوانی در دسته فراداده یا داده (فایل‌های وُرد یا پی‌دی‌اف) قرار دارد. در سطح دوم، از ناهمخوانی‌های مشاهده‌شده در فراداده‌های «پارسا» در سامانه ثبت، قلم اطلاعاتی پایه و اساس دسته‌بندی خواهد بود. بنابراین، در سطح بعدی باید اقلام اطلاعاتی مختلف سامانه ثبت مشخص شود. شکل ۴، دسته‌بندی شکل‌گسترده قلم اطلاعاتی را در فراداده‌های «پارسا» نمایش می‌دهد. به‌صورت مشابه، در داده‌های ثبت‌شده در سامانه ثبت (فایل‌های وُرد و پی‌دی‌اف) نیز سطح دوم، دسته‌بندی ناهمخوان‌ها شامل قلم‌های اطلاعاتی مختلف موجود در این داده‌هاست. شکل ۵، ساختار گسترده این ناهمخوانی‌ها را نمایش می‌دهد.

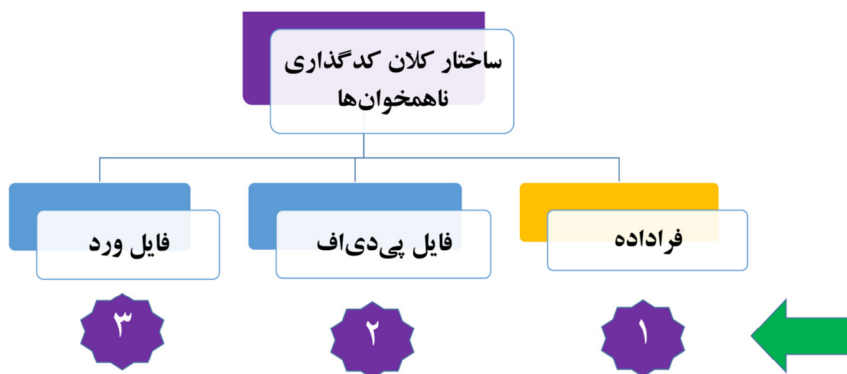


شکل ۵. ساختار گسترده سطح دوم (فهم اطلاعاتی) در ناهمخوانی‌های داده در مدارک ثبت شده در سامانه ثبت «پارسا»

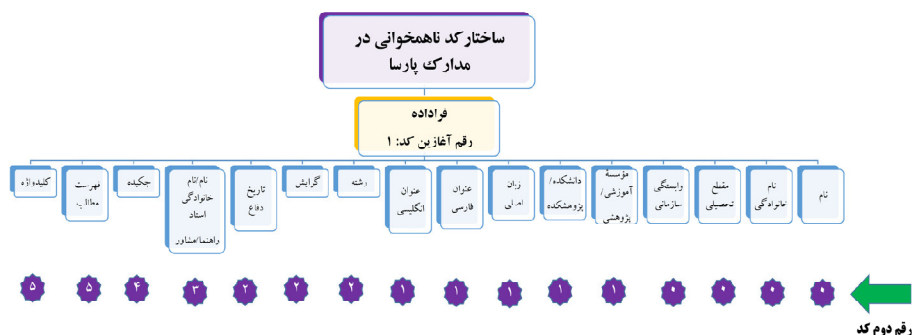


ساختار کدگذاری ناهمخوانی ها

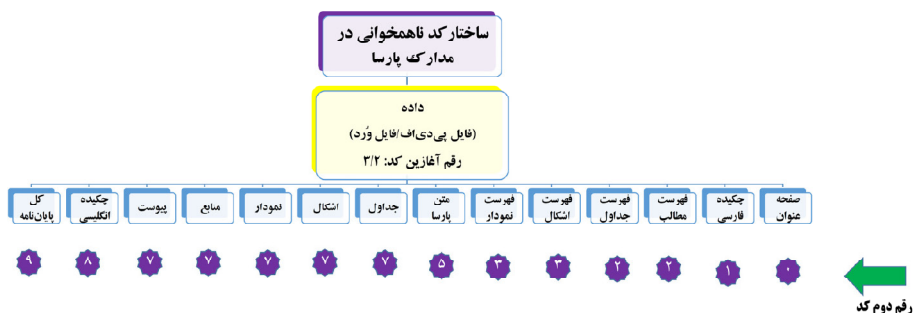
به منظور بهبود دسترسی کارشناسان کنترل کننده سامانه ثبت به ساختار استاندارد موارد ناهمخوانی، کدهای داده شده به هر ناهمخوانی به شکلی استاندارد ارائه شده است تا اختصاص کد به راحتی صورت پذیرد. همان گونه که در شکل ۳، نیز قابل مشاهده است، ناهمخوانی ها در سه دسته فراداده، فایل ورد و فایل پی دی اف ارائه شده است. به همین دلیل، کدهای این سه دسته به ترتیب با عدد ۱، ۲ و ۳ آغاز شده اند. ساختار کلی این کدگذاری در شکل ۶، قابل مشاهده است. از سوی دیگر، رقم دوم به قلم اطلاعاتی مربوط می شود و سعی شده است که ساختار مناسبی نیز برای رقم دوم ارائه شود. شکل های ۷ و ۸ ساختار رقم دوم ناهمخوانی ها را به ترتیب در مورد فراداده (رقم ابتدایی ۱) و فایل پی دی اف (رقم ابتدایی ۲) و فایل ورد (رقم ابتدایی ۳) ارائه کرده است.



شکل ۶. ساختار کلان کدگذاری موارد ناهمخوان در مدارک «پارسا»



شکل ۷. ساختار کدگذاری رقم دوم در موارد ناهمخوانی در فراداده های سامانه ثبت «پارسا»



شکل ۸. ساختار کدگذاری رقم دوم در موارد ناهمخوانی در فایل‌های پی‌دی‌اف و ورد سامانه ثبت «پارسا»

پس از تعیین رقم‌های اول و دوم از ساختار کد ناهمخوانی، رقم سوم به صورت پشت سر هم از یک آغاز شده و ادامه می‌یابد. با توجه به این که در تحلیل انجام شده بر روی کیفیت داده‌ها تمرکز بر دو رقم نخستین است، از ارائه ساختار رقم سوم (با توجه به محدودیت حجم مقاله) چشم‌پوشی شد.

برای بررسی روایی چارچوب تهیه شده، کارگروهی از خبرگان حوزه کیفیت داده/فراداده تشکیل شد. معاون اطلاعات علم و فناوری ایران، مدیر سازماندهی و تحلیل اطلاعات، مدیر ثبت و فراهم‌آوری اطلاعات و سه تن از کارشناسان اصلی شاغل در کنترل مدارک «پارسا» این کارگروه را تشکیل دادند. هر یک از اعضای این کارگروه از جنبه‌های گوناگون فنی و کاربردی پیش‌نویس تهیه شده را مورد ارزیابی و بررسی قرار دادند. سادگی استفاده برای کارشناسان سامانه ثبت، راحتی فهم آن‌ها برای دانشجوی کاربر سامانه ثبت، قابلیت تحلیل و جمع‌بندی و همچنین، کامل و جامع بودن کدها و ساختار ارائه شده از معیارهای اصلی در بررسی روایی ساختار ارائه شده بود. در پایان، پس از نهایی شدن ساختار و کدها، گام دوم پژوهش یعنی جاری‌سازی این ساختار آغاز شد که در ادامه، به ارائه آن پرداخته خواهد شد.

۴-۲. ارزیابی کیفیت داده‌ها و فراداده‌های «پارسا»ها بر پایه چارچوب تدوین شده

در این گام، چارچوب استاندارد ارائه شده در بخش پیشین در سطح عملیاتی توسط کارشناسان کنترل مورد استفاده قرار گرفت. در این راستا، راهبر سامانه ثبت بر پایه کدگذاری انجام شده ناهمخوانی‌های استاندارد شده را در سامانه تعریف نمود تا کارشناسان بتوانند در زمان ارزیابی هر مدرک «پارسا» از ساختار تازه استفاده کنند. ساختار

اعتبارسنجی شده ناهمخوانی‌ها به صورت آزمایشی در دوره دوماهه استفاده شد. در انتهای مرحله ارزیابی هر مدرک «پارسا»، در صورتی که آن مدرک ناهمخوان ارزیابی می‌شد، یک یا چند کد از کدهای استاندارد شده در گام دوم، انتخاب شده و در سامانه ثبت شد. این کدها در هر بازه زمانی مورد نظر راهبر سامانه امکان گزارش‌گیری داشته و می‌توانند بستر مناسبی برای تحلیل نتایج در گام‌های دیگر پژوهش را فراهم آورند. در گام سوم این پژوهش تحلیل‌هایی از نتایج به دست آمده در گام دوم ارائه خواهد شد.

۳-۴. تجزیه و تحلیل نتایج ارزیابی کیفیت

در این مرحله، پس از دریافت اطلاعات مدارک بازگردانده شده (۱۰,۰۰۰ مدرک) و کدهایی که برای هر مدرک بازگردانده شده انتخاب شده (ساختار این کد در بخش سوم نشان داده شد)، ابتدا فراوانی هر کد به دست آمد و سپس، درصد مشاهده هر کد از تقسیم تعداد مشاهده هر کد به کل تکرارهای همه کدها به دست آمد. برای افزایش دقت در تحلیل نتایج ارزیابی کیفیت داده‌ها، نتایج در این بخش در دو زیربخش جداگانه ارائه خواهد شد. در زیربخش اول، تحلیلی کلان برای همه داده‌ها و فراداده‌ها -مستقل از این که یک مدرک در چه حوزه تخصصی مانند فنی-مهندسی یا علوم انسانی باشد- ارائه خواهد شد. سپس، از آنجا که حوزه تخصصی دانش‌آموختگان تأثیر به‌سزایی در کیفیت مدارک ثبت شده دارد، تحلیلی دیگر متناسب با حوزه تخصصی مدارک ارائه می‌شود. یکی از مهم‌ترین مزیت‌های این تحلیل آن است که در سامانه ثبت می‌توان رهنمودهایی بر پایه حوزه تخصصی کاربران هنگام ورود اطلاعات ارائه کرد. برای نمونه، راهنمای ارائه شده به یک دانشجوی فنی-مهندسی می‌تواند نسبت به یک دانشجوی علوم پایه که از نرم‌افزار Latex برای نوشتن متن «پارسا»ی خود استفاده می‌کند، متفاوت باشد.

در ادامه و در زیربخش ۴-۳-۱، نتایج تحلیل کلان بر روی مدارک ارائه خواهد شد.

۳-۴-۱. نتایج تحلیل ارزیابی کیفیت در همه مدارک ناهمخوان

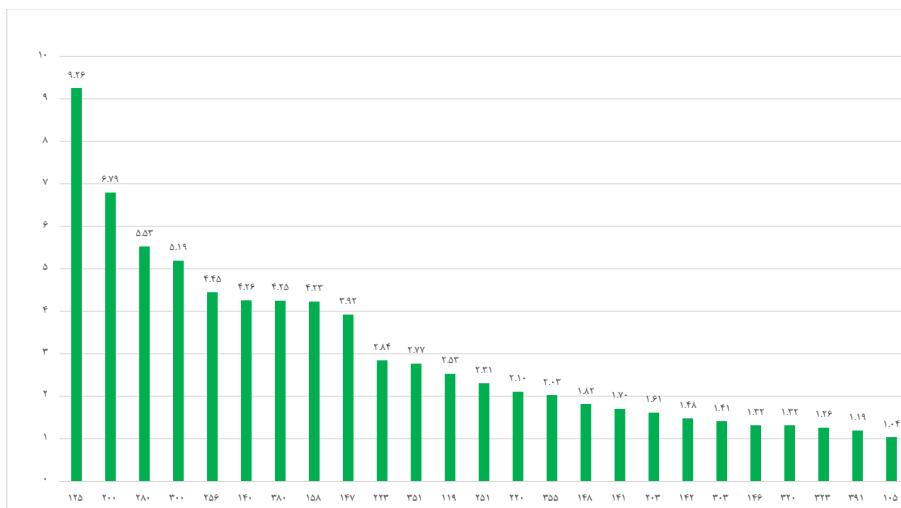
در این بخش کدهای مربوط به داده و فراداده به صورت کامل در کنار هم قرار گرفتند و هیچ‌گونه جداسازی صورت نگرفت. همچنین، دسته‌بندی ویژه‌ای بر روی مدارک (مانند حوزه تخصصی) انجام نشد. در جدول ۱، فراوانی مشاهده تمام کدها به تفکیک هر کد مشاهده می‌شود. نمودار ۱، مقایسه بهتری از اطلاعات دریافت شده از جدول ۱، را

به دست می‌دهد. گفتنی است، برای این که کدهای پُر تکرار قابلیت مشاهده پذیری بهتری داشته باشند و بتوان به شکل مناسب مقایسه انجام داد، در این نمودار فقط کدهایی که بیشتر از یک درصد از همه کدها سهم داشتند، ارائه شده‌اند.

جدول ۱. فراوانی و درصد مشاهده علت‌های ناهمخوانی در همه مدارک

ردیف	نام کد	فراوانی	درصد مشاهده ردیف	نام کد	فراوانی	درصد مشاهده
۱	۱۲۵	۹۶۳	۹/۲۶	۱۱۵	۵۰	۰/۴۸
۲	۲۰۰	۷۰۷	۶/۷۹	۱۳۰	۴۸	۰/۴۶
۳	۲۸۰	۵۷۵	۵/۵۳	۱۰۱	۴۷	۰/۴۵
۴	۳۰۰	۵۴۰	۵/۱۹	۳۵۴	۴۶	۰/۴۴
۵	۲۵۶	۴۶۳	۴/۴۵	۲۲۵	۴۵	۰/۴۳
۶	۱۴۰	۴۴۳	۴/۲۶	۳۵۳	۴۴	۰/۴۲
۷	۳۸۰	۴۴۲	۴/۲۵	۱۲۸	۴۳	۰/۴۱
۸	۱۵۸	۴۴۰	۴/۲۳	۱۳۵	۴۲	۰/۴۰
۹	۱۴۷	۴۰۸	۳/۹۲	۳۲۸	۴۲	۰/۴۰
۱۰	۲۲۳	۲۹۶	۲/۸۴	۱۵۵	۴۱	۰/۳۹
۱۱	۳۵۱	۲۸۸	۲/۷۷	۲۸۴	۴۱	۰/۳۹
۱۲	۱۱۹	۲۶۳	۲/۵۳	۲۳۰	۳۷	۰/۳۶
۱۳	۲۵۱	۲۴۰	۲/۳۱	۳۳۳	۳۶	۰/۳۵
۱۴	۲۲۰	۲۱۹	۲/۱۰	۱۴۳	۳۴	۰/۳۳
۱۵	۳۵۵	۲۱۱	۲/۰۳	۲۵۵	۳۴	۰/۳۳
۱۶	۱۴۸	۱۸۹	۱/۸۲	۱۲۹	۳۳	۰/۳۲
۱۷	۱۴۱	۱۷۷	۱/۷۰	۱۲۱	۳۲	۰/۳۱
۱۸	۲۰۳	۱۶۸	۱/۶۱	۲۱۰	۳۲	۰/۳۱
۱۹	۱۴۲	۱۵۴	۱/۴۸	۳۵۲	۳۲	۰/۳۱
۲۰	۳۰۳	۱۴۷	۱/۴۱	۱۵۶	۳۱	۰/۳۰
۲۱	۱۴۶	۱۳۷	۱/۳۲	۳۲۵	۳۰	۰/۲۹
۲۲	۳۲۰	۱۳۷	۱/۳۲	۲۸۷	۲۹	۰/۲۸
۲۳	۳۲۳	۱۳۱	۱/۲۶	۱۲۴	۲۸	۰/۲۷

ردیف	نام کد	فراوانی	درصد مشاهده ردیف	نام کد	فراوانی	درصد مشاهده
۲۴	۳۹۱	۱۲۴	۱/۱۹	۱۲۶	۲۵	۰/۲۴
۲۵	۱۰۵	۱۰۸	۱/۰۴	۲۰۱	۲۵	۰/۲۴
۲۶	۱۳۶	۸۸	۰/۸۵	۲۹۲	۲۵	۰/۲۴
۲۷	۱۳۳	۸۱	۰/۷۸	۳۳۰	۲۵	۰/۲۴
۲۸	۱۱۷	۸۰	۰/۷۷	۲۵۷	۲۳	۰/۲۲
۲۹	۱۱۶	۷۵	۰/۷۲	۳۰۱	۲۳	۰/۲۲
۳۰	۲۲۸	۷۰	۰/۶۷	۲۱۳	۲۱	۰/۲۰
۳۱	۲۹۱	۶۸	۰/۶۵	۲۹۴	۲۱	۰/۲۰
۳۲	۲۹۳	۶۱	۰/۵۹	۳۸۴	۲۱	۰/۲۰
۳۳	۳۹۳	۶۰	۰/۵۸	۱۱۱	۱۹	۰/۱۸
۳۴	۲۳۳	۵۹	۰/۵۷	۱۰۲	۱۸	۰/۱۷
۳۵	۳۹۴	۵۸	۰/۵۶	۱۲۲	۱۸	۰/۱۷
۳۶	۲۵۴	۵۷	۰/۵۵	۳۸۷	۱۸	۰/۱۷
۳۷	۲۵۲	۵۶	۰/۵۴	۲۳۸	۱۷	۰/۱۶
۳۸	۲۵۳	۵۶	۰/۵۴	۱۰۴	۱۶	۰/۱۵
۳۹	۱۲۳	۵۵	۰/۵۳	۱۵۳	۱۶	۰/۱۵
۴۰	۱۱۵	۵۰	۰/۴۸	۱۵۷	۱۵	۰/۱۴



نمودار ۱. درصد مشاهده کدهای مختلف برای ناهمخوانی‌ها در همه مدارک

نتایج ارائه شده در نمودار ۱، نشان می‌دهد که بیشترین ناهمخوانی به ناهمخوانی در کیفیت قرارداد (اطلاعات ثبت شده در سامانه توسط کاربر) مربوط می‌شود (کد ۱۲۵). به بیان دیگر، مشکل تاریخ دفاع و ناهمخوانی آن با صفحه عنوان «پارسا» در بیش از ۹ درصد مدارک بازگردانده دیده شده است. در ۵ کد اصلی استفاده شده در مدارک ناهمخوان به جز کد اول که به تاریخ دفاع کاربر مربوط می‌شود، ۴ کد دیگر در مورد داده (فایل‌های پی‌دی‌اف و ورد) است. کد ۲۰۰ (رتبه دوم)، ناهمخوانی صفحه عنوان فارسی در فایل پی‌دی‌اف با اطلاعات ثبت شده در سامانه، در ۶/۶۷ درصد از مدارک بازگردانده دیده شده است. کد ۲۸۰ (رتبه سوم)، به صفحه عنوان انگلیسی مربوط می‌شود. ناهمخوانی این صفحه با اطلاعات درج شده در سامانه در ۳/۵۳ درصد از مدارک ناهمخوان دیده شده است. کدهای ۳۰۰ و ۲۵۶، مشکلات صفحه عنوان فارسی در فایل ورد و همچنین، وجود صفحه‌های سفید در پایان‌نامه، رتبه‌های چهارم و پنجم هستند که به ترتیب ۵/۱۹ و ۴/۴۵ درصد از مدارک بازگردانده شده دارای این ناهمخوانی‌ها بوده‌اند.

۴-۳-۲. نتایج تحلیل ارزیابی کیفیت در مدارک ناهمخوان به تفکیک حوزه تخصصی

در این بخش به تحلیل علت‌های بازگرداندن مدارک «پارسا» به تفکیک حوزه‌های تخصصی مانند علوم پزشکی، فنی-مهندسی و ... پرداخته خواهد شد. در ادامه، بر پایه تحلیل‌های انجام گرفته، برخی نتایج کلیدی به تفکیک گروه تخصصی ارائه خواهد شد

(جدول ۲).

جدول ۲. مقایسه درصد مشاهده کدهای پُر تکرار به تفکیک حوزه تخصصی

حوزه تخصصی	دامپزشکی	علوم انسانی	علوم پایه	فنی مهندسی	علوم پزشکی	کشاورزی	هنر
کدهای	کد	درصد کد	درصد کد	درصد کد	درصد کد	درصد کد	درصد کد
پُر تکرار	۲۵۶	۱۰	۱۲۵	۹	۳۹۱	۱۰	۱۲۵
دارای	۱۲۵	۸	۲۰۰	۷	۲۵۶	۶	۲۸۰
ناهمخوانی	۱۵۸	۷	۲۸۰	۶	۱۲۵	۶	۲۰۰
	۲۰۰	۷	۳۰۰	۵	۲۰۰	۴	۳۰۰
	۳۰۰	۷	۱۴۰	۵	۱۵۸	۵	۳۵۱

از میان ۱۰۴۰۶ کد استفاده‌شده در مدارک بازگردانده‌شده به کاربران ۹۳ عدد (۱ درصد) به حوزه علوم پزشکی، ۲۱۵۴ عدد (۲۱ درصد) به حوزه فنی-مهندسی، ۶۰ عدد (۱ درصد) به حوزه دامپزشکی، ۵۷۴۷ عدد (۵۵ درصد) به حوزه علوم انسانی، ۱۰۳۹ عدد (۱۰ درصد) به حوزه علوم پایه، ۵۹۳ عدد (۶ درصد) به حوزه کشاورزی، و ۷۲۰ عدد (۷ درصد) به حوزه هنر مربوط است.

در حوزه دامپزشکی، کدهای ۲۵۶، ۱۲۵، ۱۵۸، ۲۰۰ و ۳۰۰ بیشترین فراوانی را در بین کدهای انتخاب‌شده داشتند. صفحه‌های سفید در فایل پی‌دی‌اف، ناهمخوانی تاریخ دفاع ثبت‌شده با صفحه عنوان، ثبت بیش از یک کلیدواژه در یک فیلد و نیز ناهمخوانی صفحه عنوان فارسی در فایل‌های پی‌دی‌اف و وُرد بیشترین دلیل‌ها در ناهمخوانی‌ها بودند که به ترتیب، ۱۰، ۸، ۷، ۷ و ۷ درصد از مدارک ناهمخوان دارای این نوع ناهمخوانی‌ها بودند. در حوزه علوم انسانی، کدهای ۱۲۵، ۲۰۰، ۲۸۰، ۳۰۰ و ۱۴۰ بیشترین فراوانی را در بین کدهای انتخاب‌شده داشتند. از این بین، ناهمخوانی تاریخ دفاع ثبت‌شده با صفحه عنوان بیشترین فراوانی را داشت. ناهمخوانی صفحه عنوان فارسی در فایل پی‌دی‌اف، ناهمخوانی صفحه عنوان انگلیسی در فایل پی‌دی‌اف، ناهمخوانی صفحه عنوان فارسی در فایل وُرد و در نهایت، بی‌دقتی در ثبت چکیده به ترتیب، بیشترین دلیل‌ها در ناهمخوانی‌ها بودند که به ترتیب ۹، ۷، ۶، ۵ و ۵ درصد از مدارک ناهمخوان دارای این نوع ناهمخوانی‌ها بودند.

در حوزه علوم پایه، کدهای ۳۹۱، ۲۵۶، ۱۲۵، ۲۰۰ و ۱۵۸ بیشترین فراوانی را در بین کدهای انتخاب‌شده داشتند. نوشتن پایان‌نامه با فرمتی دیگر، وجود صفحه‌های سفید

در فایل پی‌دی‌اف، ناهمخوانی تاریخ دفاع ثبت شده با صفحه عنوان، ناهمخوانی میان اطلاعات ثبت شده در سامانه و صفحه عنوان فارسی در فایل پی‌دی‌اف و در نهایت، درج بیش از یک کلیدواژه در یک فیلد بیشترین دلیل‌ها در ناهمخوانی‌ها بودند که به ترتیب ۱۰، ۶، ۶ و ۵ درصد از مدارک ناهمخوان دارای این نوع ناهمخوانی بودند.

در حوزه فنی-مهندسی، کدهای ۱۲۵، ۲۰۰، ۳۰۰، ۲۸۰ و ۲۵۶ بیشترین فراوانی را در بین کدهای انتخاب شده داشتند. ناهمخوانی تاریخ دفاع ثبت شده با صفحه عنوان، ناهمخوانی میان اطلاعات ثبت شده در سامانه و صفحه عنوان فارسی در فایل پی‌دی‌اف، ناهمخوانی میان اطلاعات ثبت شده در سامانه و صفحه عنوان فارسی در فایل ورد، ناهمخوانی صفحه عنوان انگلیسی در فایل پی‌دی‌اف و در نهایت، وجود صفحات سفید و سپس، ناهمخوانی تاریخ دفاع ثبت شده با صفحه عنوان بیشترین دلیل‌ها در ناهمخوانی‌ها بودند که به ترتیب، ۱۱، ۷، ۶ و ۵ درصد از مدارک ناهمخوان دارای این نوع ناهمخوانی بودند.

در حوزه علوم پزشکی، کدهای ۲۰۰، ۲۵۶، ۱۴۰، ۱۱۹ و ۱۲۵ بیشترین فراوانی را در بین کدهای انتخاب شده داشتند. در این حوزه بر خلاف حوزه‌های دیگر، کدهای مربوط به فراداده سهم بیشتری در ۵ کد پُر تکرار دارند. ناهمخوانی صفحه عنوان فارسی در فایل پی‌دی‌اف، وجود صفحه‌های سفید در فایل پی‌دی‌اف، وجود کلمه «چکیده»/ خلاصه» در ابتدای چکیده ثبت شده در سامانه، ناهمخوانی عنوان انگلیسی پایان‌نامه/ رساله در سامانه با برگ عنوان، بیشترین دلیل‌ها در ناهمخوانی‌ها بودند که به ترتیب، ۸، ۷، ۵ و ۴ درصد از مدارک ناهمخوان دارای این نوع ناهمخوانی بودند.

در حوزه کشاورزی، کدهای ۱۵۸، ۱۲۵، ۲۵۶، ۲۰۰ و ۱۴۰ بیشترین فراوانی را در بین کدهای انتخاب شده داشتند. درج بیش از یک کلیدواژه در یک فیلد، ناهمخوانی تاریخ دفاع با اطلاعات صفحه عنوان پایان‌نامه/ رساله، وجود صفحه‌های سفید در فایل پی‌دی‌اف، ناهمخوانی صفحه عنوان فارسی در فایل پی‌دی‌اف با اطلاعات سامانه، و در نهایت، وجود کلمه «چکیده»/ خلاصه» در ابتدای چکیده ثبت شده در سامانه بیشترین دلیل‌ها در ناهمخوانی‌ها بودند که به ترتیب، ۹، ۷، ۷ و ۴ درصد از مدارک ناهمخوان دارای این نوع ناهمخوانی بودند.

در حوزه هنر، کدهای ۱۲۵، ۲۸۰، ۲۰۰، ۳۰۰ و ۳۵۱ بیشترین فراوانی را در بین کدهای انتخاب شده داشتند. ناهمخوانی تاریخ دفاع با اطلاعات صفحه عنوان پایان‌نامه/ رساله با صفحه عنوان بیشترین فراوانی را داشت. ناهمخوانی میان صفحه عنوان انگلیسی فایل

پی دی اف با اطلاعات سامانه، ناهمخوانی صفحه عنوان فارسی در فایل پی دی اف با اطلاعات سامانه، ناهمخوانی صفحه عنوان فارسی در فایل وُرد با اطلاعات سامانه در اولویت های بعدی هستند. در نهایت، ناقص بودن متن فایل وُرد، بیشترین دلیل ها در ناهمخوانی ها بودند که به ترتیب، ۱۰، ۷، ۷، ۵ و ۵ درصد از مدارک ناهمخوان دارای این نوع ناهمخوانی بودند. همچنین، در ۱۰۴۰۶ بار استفاده از کدها، ۶۱۰۵ بار از کدهای داده (مشکل در فایل ها) و ۴۳۰۱ بار از کدهای فراداده (اطلاعات ثبت شده در سامانه) استفاده شده است. به بیان دیگر، ۵۹ درصد از ناهمخوانی ها به فایل های ضمیمه شده و ۴۱ درصد به اطلاعات ثبت شده در سامانه مربوط می شود. همان گونه که پیش از این نیز اشاره شد، ناهمخوانی تاریخ دفاع بیشترین مشکل مشاهده شده است. همچنین، بی دقتی در ثبت چکیده (کد ۱۴۰) و درج بیش از یک کلیدواژه در یک فیلد (کد ۱۵۸) دارای بیشترین تکرار در ناهمخوانی های مشاهده شده هستند. در ادامه، برخی نتایج کلیدی بر پایه تحلیل های انجام گرفته به تفکیک گروه تخصصی ارائه خواهد شد (جدول ۲).

۴-۴. ارائه راهکارهای بهبود

همان گونه که در بخش سوم نیز اشاره شد، گام پایانی این پژوهش ارائه راهکارهای بهبود است. به بیان دیگر در گام چهارم، می توان بر پایه تحلیل های گام سوم، راهکارهای ارتقای کیفیت را پیشنهاد داد. از این رو، در این بخش بر مبنای ارزیابی و تحلیل انجام شده بر روی کدهای پُر تکرار می توان اقدامات اصلاحی زیر را مورد توجه قرار داد.

- ◇ ارائه راهنمایی های مناسب به کاربران در سامانه هنگام درج تاریخ دفاع، به منظور کاهش مشکل ناهمخوانی تاریخ دفاع ثبت شده در سامانه با صفحه عنوان؛
- ◇ ارائه فرمت پیشنهادی مناسب در سامانه برای صفحه های عنوان فارسی / انگلیسی هنگام ارسال فایل ها توسط کاربر؛
- ◇ فراهم کردن امکان تهیه فایل پی دی اف از فایل وُرد آماده شده توسط کاربر هنگام ارسال فایل ها؛
- ◇ یادآوری نبودن صفحه های سفید در فایل پی دی اف و ارائه راهنمایی های مناسب در این مورد هنگام ارسال فایل ها در سامانه ثبت؛
- ◇ فراهم کردن زیرساخت مناسب در سامانه ثبت به منظور پیشگیری از ثبت همزمان چند کلیدواژه در یک فیلد توسط کاربران؛

- ◇ از آنجا که یکی از ناهمخوانی‌های مشاهده‌شده وجود عبارت «چکیده/ خلاصه» در سامانه ثبت هنگام ثبت اطلاعات پارساست، می‌توان در این مورد راهنمایی‌های لازم را در سامانه به کاربران ارائه کرد؛
- ◇ از آنجا که در حوزه‌های تخصصی مختلف کدهای پُر تکرار با یکدیگر متفاوت هستند، راهکارها و راهنمایی‌ها را می‌توان متناسب با حوزه تخصصی کاربر ارائه کرد. برای نمونه، بخش قابل توجهی از کاربران حوزه علوم پایه، فایل‌های دیگر، به‌جز فایل پی‌دی‌اف را بارگذاری کرده‌اند که در این زمینه می‌توان رهنمودهای لازم را به دانشکده‌های مرتبط ارائه کرد. این ناهمخوانی در دیگر حوزه‌های تخصصی دیده نشده است.

از سوی دیگر، با توجه به اینکه این پژوهش در یک بازه زمانی مشخص انجام شده است، ایجاد رویه برای فرایند کنترل کیفیت در استفاده از کدهای ناهمخوانی، ایجاد فرهنگ کاری مناسب، بازنگری شیوه‌نامه‌های ارزیابی عملکرد کارکنان مدیریت ثبت و فراهم‌آوری اطلاعات، و همچنین، بهینه‌سازی کدهای ناهمخوانی در زمان نهایی‌سازی اجرای کدها از جمله اقدامات کلیدی آتی هستند.

۵. بحث

در فرایند ثبت «پارسا»ها استاندارد نبودن ناهمخوانی‌های شناسایی‌شده و نبود دسته‌بندی مناسب برای مدارک ناهمخوان، به اطلاع‌رسانی سلیقه‌ای پژوهشگران منجر شده، تحلیل‌های آماری از مشکلات کیفی فراداده‌ها را با دشواری روبه‌رو ساخته، و تحلیل ریشه‌ای خطاهای مشاهده‌شده را ناممکن می‌سازد. از این رو، در این پژوهش ساختار ناهمخوان‌های مشاهده‌شده در سامانه ثبت پس از استاندارد شدن، در قالب ساختاری درختی در سامانه ثبت قرار داده شده است. همان‌گونه که پیش از این نیز اشاره شد، هدف غایی این پژوهش بهبود کیفیت داده (یا فراداده)های «پارسا»های درون کشور است. به‌منظور ارائه راهکارهای بهبود، دو نوع استراتژی کلی قابل پیگیری است که در پژوهش‌های پیشین، مانند Batini et al. (2014)، Glowalla et al. (2019)، Günther et al. (2009) et al. نیز به آن‌ها به تفصیل اشاره شده است.

۱. استراتژی‌های داده‌محور^۱

۲. استراتژی‌های فرایندمحور^۲

استراتژی‌های داده‌محور، کیفیت داده‌ها را با اصلاح مستقیم ارزش داده‌ها بهبود می‌بخشد. برای نمونه، کیفیت داده‌های منسوخ‌شده با به‌روزرسانی آن‌ها از یک پایگاه داده دیگر و جایگزینی آن‌ها با داده‌های به‌روزشده بهبود می‌یابد. استراتژی‌های فرایندمحور نیز کیفیت را توسط طراحی مجدد فرایندهای تولید و تغییر داده بهبود می‌بخشند.

در ارائه راهکارهای بهبود در این پژوهش برخی راهکارهای ناظر بهبود داده شده و برخی دیگر فرایندهای تولید و پردازش را مورد توجه قرار داده‌اند. راهکارهایی مانند ارائه راهنما به کاربر در فرایند ثبت «پارسا» و یا در دسته‌بندی فرایندمحور قرار می‌گیرند. راهکارهای فرایندمحور در بهبود کیفیت داده در پژوهش‌های (Falge, Otto & Azeroual et al. (2020), Österle (2012), Michelberger, Mutschler & Reichert (2011) مورد توجه قرار گرفته است. از سوی دیگر، معرفی فرمت مشخص برای فایل‌های «پارسا» و برطرف کردن مشکلات کیفی فایل‌های پی‌دی‌اف در دسته راهکارهای داده‌محور قرار می‌گیرند.

در همین راستا، (Azeroual et al. (2020) در پژوهش خود راهکارهایی را برای مواجه شدن با ناهمخوانی در مدیریت داده‌های پژوهشی مورد توجه قرار داده است. چنین راهکارهایی در پژوهش‌های (Sidi et al. (2012) و نیز (Batini et al. (2009) ارائه شده است.

۶. نتیجه‌گیری و ارائه پیشنهادها

در این پژوهش بر پایه رویکرد TDQM چارچوبی برای ارزیابی و تحلیل و در نهایت، بهبود کیفیت داده‌ها و فراداده‌های «پارسا» ارائه شد. پس از پیاده‌سازی این چارچوب در سامانه ثبت «پارسا»‌های درون کشور نتایج به‌صورت جامع مورد تحلیل قرار گرفت و راهکارهای بهبود کیفیت ارائه شد.

طیف وسیع تخصص‌های گوناگون در پایان‌نامه‌ها و رساله‌ها و دیدگاه‌های گوناگونی که پژوهشگران هر یک از این حوزه‌ها دارند، موجب می‌شود که تحلیل کیفیت اطلاعات ثبت‌شده به آسانی امکان‌پذیر نشود. همان‌گونه که نتایج این پژوهش نشان داد، کیفیت اطلاعات ثبت‌شده (در سطح فراداده) و اطلاعات موجود در فایل‌های بارگذاری‌شده «پارسا»‌ها

در هر حوزه متفاوت است و این نشان می‌دهد که لازم است برای بهبود کیفیت اطلاعات «پارسا»ها در سطح کلان، راهکارهایی را برای هر حوزه تخصصی تعریف نمود. در پژوهش‌های آتی با استفاده از تکنیک‌های داده کاوی بر روی نتایج به‌دست آمده از کدهای ناهمخوان، می‌توان تحلیل‌های کامل‌تری بر روی مشکلات کیفی مدارک، به تفکیک دانشگاه، رشته یا گرایش یا سایر اقلام اطلاعاتی انجام داد. از آنجا که کدهای ناهمخوانی بر پایه استفاده از تکنیک‌ها و رویکردهای کیفیت داده نگاشته شده‌اند، ایجاد نگاشت میان این کدها و همچنین، ابعاد کیفیت داده مانند کامل بودن، دسترس‌پذیری و ارزیابی وضعیت کدها از رویکرد کیفیت داده در پژوهش‌های آتی قابل توجه خواهد بود.

از سوی دیگر، همبستگی میان کدهای ناهمخوانی از جمله مواردی است که در این گزارش مورد بررسی قرار نگرفته است. پیشنهاد می‌شود این شاخص نیز در آینده مورد بررسی قرار گیرد. نتایج کاربردی این کار از این جهت مورد توجه است که کارکنان کنترل‌کننده مدرک در صورت مشاهده یک کد می‌توانند پیش‌بینی کنند که ناهمخوانی دیگری نیز با احتمال وقوع بالا امکان رخداد دارد و در فرایند کنترل کیفیت آن ناهمخوانی را نیز مورد توجه قرار دهند.

فهرست منابع

- ارشادی، محمدجواد، تقی رجبی، فرهاد شیرانی، و نسا رضایی. ۱۳۹۵. کاربرد تکنیک تحلیل ریشه در حل مشکلات کیفی سامانه‌های اطلاعاتی تحقیقاتی: مطالعه موردی سامانه اشاعه اطلاعات پایان‌نامه‌ها/ رساله‌های دانش‌آموختگان داخل کشور (گنج). مدیریت اطلاعات ۱ (۱): ۷۵-۸۹.
- ارشادی، محمدجواد، و محمد مهدی ارشادی. ۱۳۹۸. مدل‌ها و روش‌های مدیریت کیفیت داده: نگرشی کاربردی بر داده‌های علم و فناوری. تهران: پژوهشگاه علوم و فناوری اطلاعات ایران.
- ارشادی، محمدجواد، عمار جلالی‌منش، و جلال‌الدین نصیری. ۱۳۹۸. طراحی مدل کیفیت فراداده: مورد کاوی سامانه ثبت پایان‌نامه/ رساله در پژوهشگاه علوم و فناوری اطلاعات ایران. پژوهشنامه پردازش و مدیریت اطلاعات ۳۴ (۴): ۱۴۹۹-۱۵۲۸.
- اشتریان اصفهانی، آیناز، محمدجواد ارشادی، و امیر عزیزی. ۱۳۹۹. پایش شاخص‌های کیفیت داده‌های پژوهشی به کمک نمودار کنترل مشاهده‌های انفرادی با دامنه متحرک. پژوهشنامه پردازش و مدیریت اطلاعات ۳۵ (۴): ۹۳۳-۹۵۷.

References

- Avenali, A., C. Batini, P. Bertolazzi, & P. Missier. 2008. Brokering infrastructure for minimum cost data procurement based on quality–quantity models. *Decision Support Systems* 45 (1): 95-109.
- Azeroual, O., G. Saake, M. Abuosba, & J. Schöpfel. 2020. Data Quality as a Critical Success Factor for User Acceptance of Research Information Systems. *Data* 5 (2): 35.
- Batini, C., F. Cabitza, C. Cappiello, & C. Francalanci. 2008. A comprehensive data quality methodology for web and structured data. *International Journal of Innovative Computing and Applications* 1 (3): 205-218.
- Batini, C., C. Cappiello, C. Francalanci, & A. Maurino. 2009. Methodologies for data quality assessment and improvement. *ACM computing surveys (CSUR)* 41 (3): 1-52.
- Cahyono, S. H., & Y. G. Sucahyo. 2020. Pengukuran Kualitas Data Menggunakan Framework Total Data Quality Management (TDQM): Studi Kasus Sistem Informasi Beasiswa Universitas Indonesia Data Quality Assessment Using the TDQM Framework: A Case Study of University of Indonesia (UI) Scholarship Information System. *Jurnal IPTEK-KOM (Jurnal Ilmu Pengetahuan dan Teknologi Komunikasi)* 22 (2): 193-206.
- Ershadi, M. J., R. Aiasi, & S. Kazemi. 2018. Root cause analysis in quality problem solving of research information systems: a case study. *International Journal of Productivity and Quality Management* 24 (2) 284-299.
- Ershadi, M. J., & D. Omidzadeh. 2018. Customer validation using hybrid logistic regression and credit scoring model: a case study. *Calitatea* 19 (167): 59-62.
- Falge, C., B. Otto, & H. Österle. 2012. Data quality requirements of collaborative business processes. In *2012 45th Hawaii International Conference on System Sciences* (pp. 4316-4325). IEEE.
- Glowalla, P., P. Balazy, D. Basten, & A. Sunyaev. 2014. Process-driven data quality management-An application of the combined conceptual life cycle model. In *2014 47th Hawaii International Conference on System Sciences* (pp. 4700-4709). IEEE.
- Günther, L. C., E. Colangelo, H. H. Wiendahl, & C. Bauer. 2019. Data quality assessment for improved decision-making: a methodology for small and medium-sized enterprises. *Procedia Manufacturing* 29: 583-591.
- Heinrich, B., M. Klier, & M. Kaiser. 2009. A procedure to develop metrics for currency and its application in CRM. *Journal of Data and Information Quality (JDIQ)* 1 (1): 5.
- Khosroanjom, D., M. Ahmadzade, A. Niknafs, & R. K. Mavi. 2011. Using fuzzy AHP for evaluating the dimensions of data quality. *International Journal of Business Information Systems* 8 (3): 269-285.
- Kwon, O., N. Lee, & B. Shin. 2014. Data quality management, data usage experience and acquisition intention of big data analytics. *International Journal of Information Management* 34 (3): 387-394.
- Liu, Q., G. Feng, X. Zhao, & W. Wang. 2020. Minimizing the data quality problem of information systems: A process-based method. *Decision Support Systems* 137: 113381.
- Michelberger, B., B. Mutschler, & M. Reichert. 2011. Towards process-oriented information logistics: why quality dimensions of process information matter. 107-120
- Nikiforova, A. 2020. Definition and Evaluation of Data Quality: User-Oriented Data Object-Driven Approach to Data Quality Assessment. *Baltic Journal of Modern Computing* 8 (3): 391-432.
- Ochoa, X., & E. Duval. 2006. Quality metrics for learning object metadata. In *EdMedia+ Innovate Learning* (pp. 1004-1011). *Association for the Advancement of Computing in Education (AACE)*.
- Peltier, J. W., D. Zahay, & D. R. Lehmann. 2013. Organizational learning and CRM success: a model for linking organizational practices, customer data quality, and performance. *Journal of interactive marketing* 27 (1): 1-13.

- Petrović, M. 2020. Data quality in customer relationship management (CRM): *Literature review. Strategic Management* 25 (2): 40-47.
- Rahman, M. S., M. Mannan, M. A. Hossain, M. H. Zaman, & H. Hassan. 2018. Tacit knowledge-sharing behavior among the academic staff: Trust, self-efficacy, motivation and Big Five personality traits embedded model. *International Journal of Educational Management* 32 (5): 761-782.
- Russell-Rose, T., J. Chamberlain, & L. Azzopardi. 2018. Information retrieval in the workplace: A comparison of professional search practices. *Information Processing & Management* 54 (6): 1042-1057.
- Sharma, S. 2020. Big Data Analytics for Customer Relationship Management: A Systematic Review and Research Agenda. In *International Conference on Advances in Computing and Data Sciences* (pp. 430-438). Springer, Singapore.
- Sidi, F., P. H. S. Panahy, L. S. Affendey, M. A. Jabar, H. Ibrahim, & A. Mustapha. 2012. Data quality: A survey of data quality dimensions. In *2012 International Conference on Information Retrieval & Knowledge Management* (pp. 300-304). IEEE. Kuala Lumpur. [DOI:10.1109/InfRKM.2012.6204995]
- Taleb, I., M. A. Serhani, & R. Dssouli. 2018. Big data quality assessment model for unstructured data. In *2018 International Conference on Innovations in Information Technology (IIT)* (pp. 69-74). IEEE. AL AIN UAE.
- Vaziri, R., M. Mohsenzadeh, & J. Habibi. 2017. Measuring data quality with weighted metrics. *Total Quality Management & Business Excellence*, 30 (5-6), 708-720.
- Wang, R. Y., & D. M. Strong. 1996. Beyond accuracy: What data quality means to data consumers. *Journal of management information systems* 12 (4): 5-33.
- Wang, R. Y., & E. Mdnick Stuart. 1990. A Polygen Model for Heterogeneous Database Systems: *The Source Tagging Perspective, Proceedings of the 16th International Conference on Very Large Data Bases, San Francisco, CA, United States*, 519-538.

محمدجواد ارشادی

دارای مدرک دکتری تخصصی رشته مهندسی صنایع از دانشگاه علم و صنعت است. ایشان هم‌اکنون دانشیار گروه پژوهشی مدیریت فناوری اطلاعات است.

کاربرد مدیریت و مهندسی کیفیت در علوم و فناوری اطلاعات، کیفیت داده و اطلاعات، کنترل کیفیت آماری، مدیریت کیفیت جامع، بازمهندسی فرایندهای کسب و کار، بهینه‌سازی، الگوریتم‌های فراابتکاری، تجزیه و تحلیل سیستم‌ها و داده کاوی آموزشی از جمله علایق پژوهشی وی است.



معصومه نبی زاده

دارای مدرک تحصیلی کارشناسی ارشد مهندسی صنایع گرایش سیستم و بهره وری از دانشگاه علوم و تحقیقات تهران است. کنترل کیفیت، مدیریت ریسک، امنیت شغلی، مدیریت پروژه و روش های MCDM از جمله علایق پژوهشی وی است.

