



Examiner's Report

Principal Examiner Feedback

Summer 2018

Pearson Edexcel GCSE

In Psychology (5PS02)

Paper 1 Social and Biological Psychology

Edexcel and BTEC Qualifications

Edexcel and BTEC qualifications are awarded by Pearson, the UK's largest awarding body. We provide a wide range of qualifications including academic, vocational, occupational and specific programmes for employers. For further information visit our qualifications websites at www.edexcel.com or www.btec.co.uk. Alternatively, you can get in touch with us using the details on our contact us page at www.edexcel.com/contactus.

Pearson: helping people progress, everywhere

Pearson aspires to be the world's leading learning company. Our aim is to help everyone progress in their lives through education. We believe in every kind of learning, for all kinds of people, wherever they are in the world. We've been involved in education for over 150 years, and by working across 70 countries, in 100 languages, we have built an international reputation for our commitment to high standards and raising achievement through innovation in education. Find out more about how we can help you and your students at: www.pearson.com/uk

Summer 2018

Publications Code 5PS02_01_1806_ER

All the material in this publication is copyright

© Pearson Education Ltd 2018

General Comments

This year saw an improvement overall in candidate performance, largely due to a better standard of responses to the essay question.

Responses to the extended response question were of a higher standard to 2017, with those using the correct study gaining greater credit than for the essay in the previous year. There was better balance in responses overall between AO1 and AO2.

In terms of specific detail, there was a similar mix of very specific responses and generic responses that have been witnessed in previous series. Candidates will always gain more marks with statements that use specific detail in examinations compared to generic, vague, statements.

With regard to application, it was again a similar trend to previous series. With some excellent AO2 application to the stimulus (with some lovely diagrams for question Q5(a)!); and other generic, rote learned, responses with no application at all. Where a scenario is provided, candidates are reminded to explicitly apply their ideas to the scenario given to them.

The remainder of this Examiner's Report will focus on each individual question which may be useful in helping prepare candidates for the new specification 9-1 GCSE examinations.

Based on their performance on this paper candidate are offered the following advice:

- Balance the content for extended open questions. Where AO1 and AO2 are assessed in a 10-mark question, half of the response should be AO1 and the other half AO2.
- Use specific detail in all responses, rather than using vague, generic, statements.
- Where a scenario is given, apply all ideas explicitly to the context.
- Read the question carefully, and respond to the question rather than giving a pre-learnt answer that they hope will fit what is being asked.

Comments on Individual Questions

Question 1(a)

The vast majority of candidates suggested a suitable example of censorship, with most gaining the mark available. The most common responses were the Watershed, age ratings, bleeping out swear words, removing/blocking out nudity. Common errors included examples of what to censor (e.g. aggression), rather than a type of censorship.

Question 1(b)

Candidates were awarded up to two marks for the one argument in favour of media censorship. Most commonly candidates scored 1 mark for a basic response, with only the best responses being able to elaborate appropriately to gain the second mark. Common candidate answers included protection of young people, trying to stop/prevent children copying inappropriate behaviour.

Question 1(c)

Candidates were awarded up to three marks for description of the procedure of their content analysis in relation to the scenario provided. In general, candidates performed well on this question with a lot scoring at least two of the three available marks. Common candidate answers included choosing a list of aggressive behaviours with an example, tallying, use of a range of programmes, or asking others to also do the analysis. Weaker responses were either brief or gave vague ideas of how to do the content analysis.

Question 1(d)

The vast majority scored 1 mark for choosing the correct response from the list.

Question 2(a)

Candidates were awarded up to four marks for explanation of how Social Learning Theory (SLT) could explain Jason's behaviour in the scenario. This question discriminated effectively, with regards to candidate ability levels, with a spread of marks over the range. The best responses had good application of SLT to the scenario given. Common candidate answers included identification of a

role model, similar characteristics to the role model, reward / vicarious reinforcement, modelling the character within the game and being more aggressive. A common error included not giving enough points. Candidates need to provide 4 points for 4 marks, but some only gave 2 or 3 points, therefore, self-limiting their achievable marks.

Question 2(b)

Candidates were awarded up to two marks for one strength of Social Learning Theory as an explanation of aggression. Candidates really struggled with this question, often giving AO1 description or brief strengths that were not enough for a mark. Where a study was used and gained a mark, there was generally not enough to link back to SLT for the second mark. Common candidate answers included using Bandura, or Anderson and Dill. Common errors included using statements that SLT supported nurture, use of AO1 only (can explain why people learn aggression), or naming a study but with no findings.

Question 2(c)

Candidates were awarded up to two marks for description of the role of hormones as a cause for aggression. Candidates performed well overall, generally scoring at least 1 mark. Common candidate answers included testosterone, castration of animals, or injecting testosterone.

Question 2(d)

The vast majority scored 1 mark for choosing the correct response from the list.

Question 2(e)(i)

Candidates were awarded one mark for an accurate definition of 'nature'. The vast majority scored 1 mark, with common candidate answers including what we are born with, genetics, internal, hormones, brain. Common errors included nurture ideas (e.g. environment, upbringing, surroundings).

Question 2(e)(ii)

Candidates were awarded one mark for an accurate definition of 'nurture'. The vast majority scored 1 mark, with common candidate answers including the

environment, upbringing, surroundings. Common errors included nature ideas (e.g. genetics, internal, hormones, brain).

Question 2(f)

Candidates were awarded up to two marks for explanation of a similarity between the social, and biological, explanations for aggression. Candidates struggled with this question and the majority did not gain credit. Where they did gain credit, common candidate answers included the idea that both cannot prove causality, or that both are incomplete.

Question 3

This essay was for 10 marks and was assessed using the levels based mark schemes (not points-based marking). The question was **describe** and **evaluate** so was looking for both knowledge, and understanding, of the study (AO1); and evaluation of the study (AO2). This question discriminated between candidate ability levels well, with candidates overall performing better than in the 2017 essay. Some candidates left this blank or used a different study. Other common errors included brief, vague, evaluation, or inaccurate AO1.

Question 4(a)

The vast majority scored 1 mark for choosing the correct response from the list.

Question 4(b)

Candidates were awarded one mark for identification of one qualification required to be a clinical psychologist. The vast majority scored 1 mark, with common candidate answers including a degree, PhD or Master's degree in Psychology. Common errors included a doctorate (not subject specific), A-level Psychology, or GCSE Psychology.

Question 4(c)

Candidates were awarded one mark for each role of a clinical psychologist they explained (to a max. of 2 marks). Overall candidates performed well, scoring at least one mark. Common candidate answers included treating patients,

diagnosing patients, research, listening to a patient issues. Common errors included using the same role twice (repetition).

Question 5(a)

Candidates were awarded up to four marks for explanation of how Classical Conditioning (CC) could explain why Mika developed a fear of letter boxes from the scenario. As this is AO2 **application** to the scenario was required. This question discriminated effectively between candidate ability levels with a spread of marks over the range. The best responses had good application of CC to the scenario given. Common errors included mixing up the UCS and NS; or use of a diagram only, with no explanation, which limited marks.

Question 5(b)

Candidates were awarded up to two marks for explanation of how generalisation could be related to Mika. This question discriminated effectively with a spread of marks over the range. Common candidate answers included the fear being transferred to something similar, or use of appropriate examples. Common errors included thinking it is about generalisability.

Question 5(c)

The vast majority scored 1 mark for choosing the correct response from the list.

Question 6(a)(i)

Candidates were awarded one mark for a suitable open question. The vast majority scored 1 mark, with the most common error of using a closed question.

Question 6(a)(ii)

Candidates were awarded one mark for a suitable closed question. The vast majority scored 1 mark, with the most common error of using an open question.

Question 6(a)(iii)

Candidates were awarded one mark for a suitable ranked scale question. The vast majority scored 1 mark; with the most common error use of a closed, but not ranked scale, question.

Question 6(b)

Candidates were awarded up to two marks for explanation of one strength of using closed questions to study phobias. This question discriminated effectively with a spread of marks awarded evenly across the marks available. Common accurate candidate answers included: quantitative so easier to analyse; quantitative so more objective; cannot be interpreted in other ways by other researchers. Common errors included the idea that closed questions produced qualitative data.

Question 6(c)

The vast majority scored 1 mark for choosing the correct response from the list.

Question 6(d)

Candidates were awarded up to two marks for explanation of how Demi could give her participants standardised instructions. Candidate performance was spread across the mark range with the question discriminating between abilities effectively. Common candidate answers included the idea that she could read a set of instructions to participants, print off a set of instructions and give them to all, check each participant knows what they are doing and instruct them the same way, or give them a brief with the same instructions on.

Question 6(e)

Candidates were awarded up to two marks for suggesting how Demi could avoid social desirability in her study. Candidate performance was spread across the mark range with the question discriminating effectively. Common candidate answers included that she could tell them it's confidential, ask them to tell the truth, or deceive them.

Question 6(f)

Candidates were awarded one mark for an accurate definition of response bias. The majority could give an accurate definition, with common candidate answers including: answer in a way they think the researcher wants; always answer in a particular way e.g. middle response; change their answer to something socially desirable. The most common error was for candidates to give tautological responses (e.g. a response that's biased).

Question 7(a)

Candidates were awarded one mark for a suitable aim of the case study. The majority could give an accurate definition and gain the mark available.

Question 7(b)

Candidates were awarded up to two marks for description of the findings of the case study. Performance was spread over the mark range, with common candidate answers including generalised his fear to things similar to a rabbit, SD treated the phobia of rabbits, was able to tolerate other animals more afterwards. The most common error was to use the findings of Little Albert (Watson and Rayner, 1920).

Question 7(c)

Candidates were awarded one mark for an appropriate suggestion of why Cover-Jones used the name Little Peter instead of the child's real name. Performance was split with slightly more offering an appropriate suggestion. The most common candidate answers included to protect his identity, stop him being found by the media, stop any ongoing trauma by being traced. The most common errors included using single word answers (e.g. 'confidentiality', 'privacy'), or generic responses not linked to Little Peter.

Question 7(d)

Candidates were awarded up to four marks for evaluation of systematic desensitisation. Generally candidates struggled with this question, as they did not read the question carefully. As such, a large proportion gave generic, rote-learned evaluation of the study, rather than the therapy used in the case study. The most common correct candidate answers included the idea about systematic

desensitisation leading to less distress than flooding, how it has control over the hierarchy, that patients can gradually move up the hierarchy, or that they are exposed to fear so could be more distressing.

Question 8(a)

Candidates were awarded up to two marks for explanation of how Stephan's genetics may have caused his criminal behaviour. Candidates most commonly achieved one mark for a basic explanation, with the best getting full marks for a more elaborated response. The most common correct answers included the XYY gene, or a criminal gene inherited from family. The most common error included use of learning theories (e.g. SLT).

Question 8(b)

Candidates were awarded up to two marks for description of how self-fulfilling prophecy could explain Stephan's criminal behaviour. Candidates most commonly achieved one mark for a basic explanation, with the best getting full marks for a more elaborated response. The most common correct answers included the idea that expectation led to the behaviour, labelled as violent so fulfilled this, treated differently due to the label. The most common errors included use of Social Learning Theory, or no reference to criminality.

Question 8(c)

Candidates were awarded one mark for identification of a suitable social explanation in Q8(c)(i) and a further three marks in Q8(c)(ii) for description of the social explanation of criminal behaviour. Candidates performed well overall with the majority scoring at least 2 marks of the 4 available. The most common candidate answers included child rearing, maternal deprivation, social learning, family size, parenting style. The most common errors included using a biological explanation.

Question 8(d)

Candidates were awarded one mark for each difference (to a max. of 2 marks). Candidates most commonly scored one mark for a single difference, when the question clearly asked for two differences. The most common correct answers included the use of nature vs. nurture, and the idea of determined vs. changeable.

Question 9(a)

Candidates were awarded one mark for a suitable aim of the study. The majority of candidates were able to gain a mark for the correct aim, with the most common answer being whether attractiveness affected jury decision making. The most common error was to give the aim of a different study.

Question 9(b)

Candidates were awarded up to three marks for describing the procedure of the study. Candidates were split between either doing very well and likely getting full marks, or zero for describing the procedure of the wrong study. Candidates less commonly achieved 1 or 2 marks. The most common candidate answers included attractive or unattractive photo, type of crime, attractiveness was judged, asked to give prison sentence.

Question 9(c)

Candidates were awarded up to two marks for one strength of the study. Candidates struggled with this and most commonly did not receive any credit. Where candidates did receive credit, the most common candidate answers included the idea that a control group was used, independent measures were used, and that it was a controlled study. The most common errors included using the wrong study, or giving a brief response that limited candidate performance.

Question 9(d)

Candidates were awarded one mark for outlining each weakness of the study (to a max. of 2 marks). Candidates performed slightly better on this question than Q9(c) but the most common mark was still zero, generally for using the wrong study. The most common correct candidate answers included the idea that it was not realistic as a jury would see more than a photo, and also that it was not realistic as a judge would pass sentence and a jury only a verdict of guilty or not guilty.

Question 10(a)

The vast majority scored 1 mark for choosing the correct response from the list.

Question 10(b)

The vast majority scored 1 mark for choosing the correct response from the list.

Question 10(c)

Candidates were awarded one mark for stating each ethical issue Jemimah should consider for her study (to a max. of 2 marks). Candidates struggled with the application and demand of this question, most commonly scoring zero marks. The most common errors included giving single word answers (e.g. 'confidentiality', 'informed consent'), or having no link to Jemimah's study. The most common rewardable answers focussed on confidentiality, informed consent, privacy, protection.

Question 10(d)

Candidates were awarded up to two marks for writing an experimental hypothesis for her study. Generally candidates struggled with this question, with only a minority getting full marks. The most common responses focused on the black/white woman given more time in jail, black/white woman given more months in jail. The most common errors included not fully operationalising the DV (no. of months in jail) or giving an aim and not a hypothesis.

Question 10(e)

Candidates were awarded up to two marks for describing how she could make her results more generalisable. Candidate performance varied evenly across the marks available with the question discriminating effectively. Common candidate answers included using more jurors, using jurors from different cultures, or using photos of men as well as women. Common errors included no link to Jemimah's study, or only giving one very brief suggestion which limited the achievable marks.

Question 10(f)

Candidates were awarded up to two marks for explaining one conclusion from Jemimah's results. Candidates tended to struggle with this question as they gave categorical responses that are not consistent with the data given to them in the stimulus material. The most common accepted responses included the conclusion that race does not affect jury decision making, race only very slightly affects jury decision making, or that black suspects may get very slightly longer in prison for the same crime.

Question 10(g)

Candidates were awarded one mark for a suitable definition of reliability. Candidates were evenly split with around half accurately defining reliability. Common candidate answers included consistency over time, or when a study is repeated it finds the same thing each time. Common errors included definitions of validity, or other inaccurate definitions.

Question 10(h)

The vast majority scored 1 mark for choosing the correct response from the list.